

IRS stands for Information Retrieval Systems.

Here

Information: Data & facts.

Information is Communicated (or) received knowledge  
conveying a Particular facts.

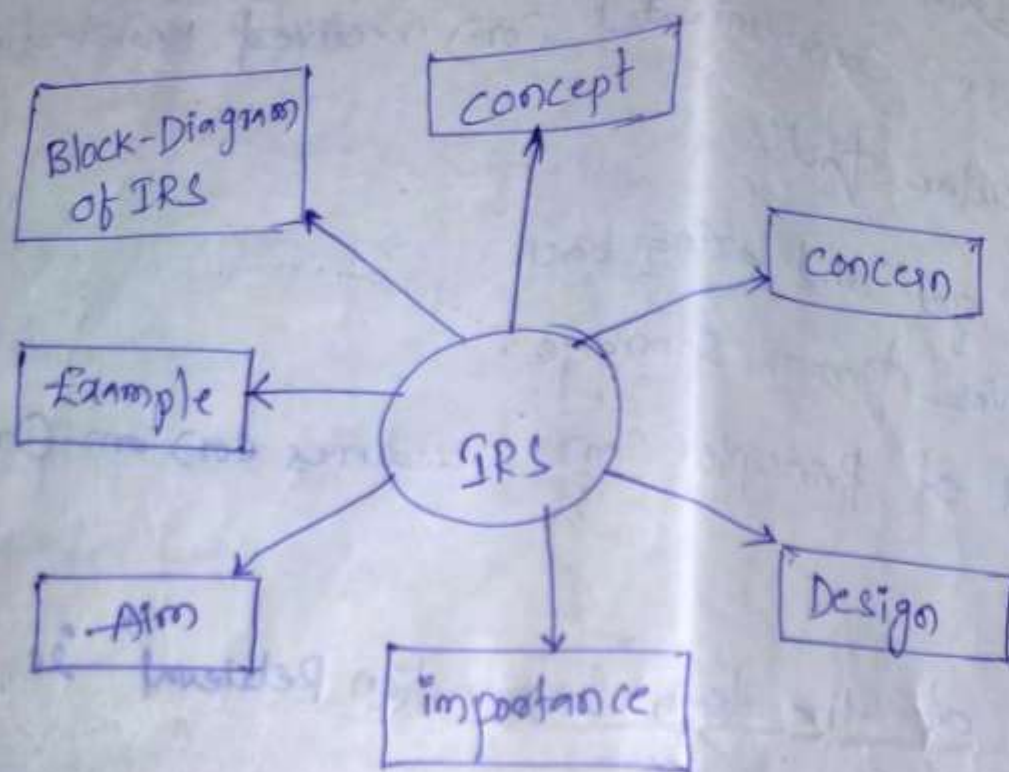
Retrieval :- To get and bring back  
i.e. recover from storage.

System: A set of Principles (or) procedures (or) an Organized  
- schema.

⇒ The meaning of the term "Information Retrieval" is  
"Very broad".

⇒ The term IR (Information Retrieval) was coined by "Kelvin Moore"  
in 1950. It gained popularity in the research communication  
from 1961 onwards.

⇒ Goal of IRS is to provide the information needed to  
satisfy the user's question.





## Introduction to Information Retrieval Systems:

⇒ The concept of Information Retrieval System (IRS) is "self-explanatory" from the "terminological point of view" and refers to a system which retrieves information.

Here, "self-explanatory" means easily understood from the information already given and not needing further explanation.

"Terminological" means technical words and expressions used in a particular subject.

⇒ IRS is concerned with two basic aspects:

- i) How to store information &
- ii) How to retrieve information.

### Design :-

⇒ An IRS is designed to "analyze process" and store sources of information and retrieve when ever required. (cor)

⇒ An IRS is a set of 'rules and procedures' for performing some (or) all of the following operations:



- 4
- \* Indexing (or) Constructing the representation of documents.
  - \* Search Formulation / Constructing the representation of information needs).
  - \* Searching / Matching representations of documents against the representations of need.
  - \* Index language / Generation of rules of representation.

### Importance of IRS :-

- ⇒ Information Retrieval Systems are Very Important.
- Important to make sense of the data.
- ⇒ Imagine how hard it would be to find some information on the Internet without Google (or) other search engines.

### Aim of IRS :-

- ⇒ It provide the "best possible" information from a database.

### Example of IRS :-

- ⇒ The most Obvious example of an IRS is "Google" and the english language has been extended with the term "Google it" means search for something.



## Block Diagram of IRs:-

It consists of 3 components. They are

- \* Input
- \* Processor
- \* Output

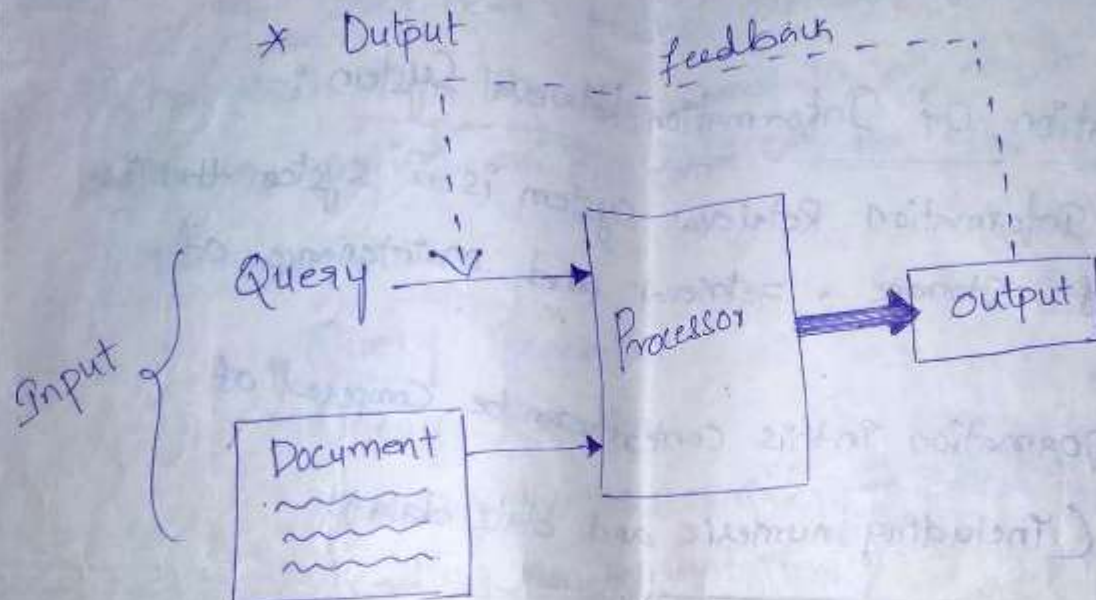


Fig. IR Block Diagram

\* Input :- It stores the only representation of the "Query" (or) "document".

Query: It represents the expression of the user's information need.

Document: It represents the list of extracted words considered to be significant.

\* Processor :- It involves performing "actual retrieval" function and also executing the search strategy in response to a query.

\* output :-

A set of extracted document is called output.

feedback :-

It Improving the subsequent run after sample retrieval.

### (\*) Defination Of Information Retrieval System :-

⇒ An Information Retrieval System is a system that is Capable of storage, retrieval and maintenance of information.

⇒ Information in this context can be composed of

\* text (Including numeric and date data).

\* images

\* audio

\* Video and

\* other multi-media Objects.

→ The following terms discussed in IRS.

User :-

The term 'User' represents an end user of the Information System who has minimum knowledge of computers and technical fields in general.

Item :-

The term 'Item' is used to represent the "Smallest Complete Unit" that is processed and manipulated by the System.





## \* Objectives of Information Retrieval Systems:

The general objective of an IRS is minimizing the overhead of a user by locating needed information.

### \* Overhead: -

The Over-head can be expressed in terms of "time" i.e. how much time the user spends in all of the steps leading to reading an item containing the needed information.

### Example: -

- \* Query Generation
- \* Query Execution
- \* Scanning the result of Query to select item to read,
- \* Reading non-relevant items.

The Success of the Information System is very "Subjective", and it depends on two things

i) What Information is needed.

ii) The willingness of a user to accept overhead.

### \* Relevant Item: -

In Information retrieval System the term "relevant" item is used to represent an item containing the needed information.



### \* User's perspective :-

From a User's perspective "relevant" and "needed" are Synonyms.

### \* System Perspective :-

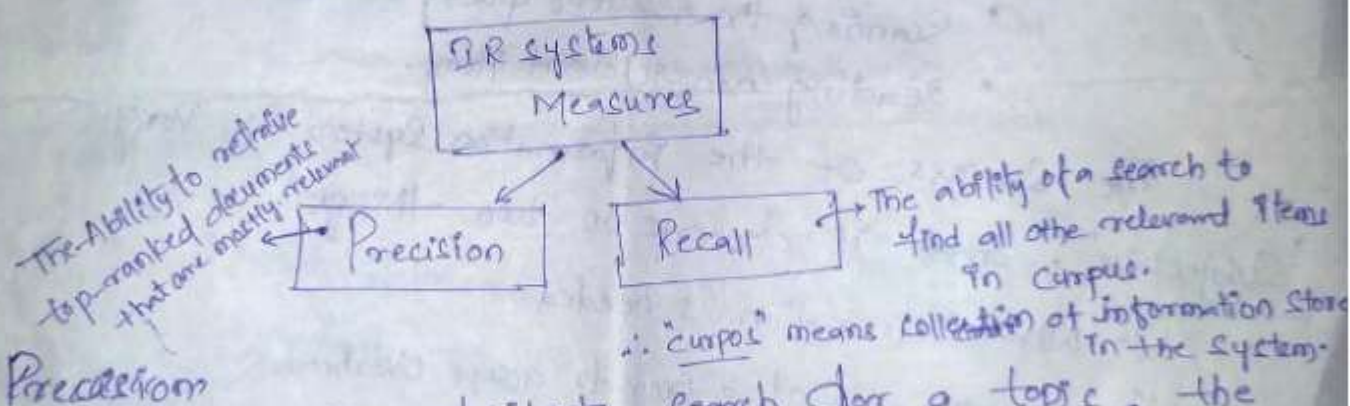
From a System Perspective, Information could be relevant to search statement even though it is not needed / relevant to users.

Eg: User already knew the information.

→ The Objectives of IRS is commonly associated with two major measures:

i) Precision

ii) Recall.



### Precision

→ When a User decides to search for a topic, the total ~~data~~ database is logically divided into 4 segments, as shown below.

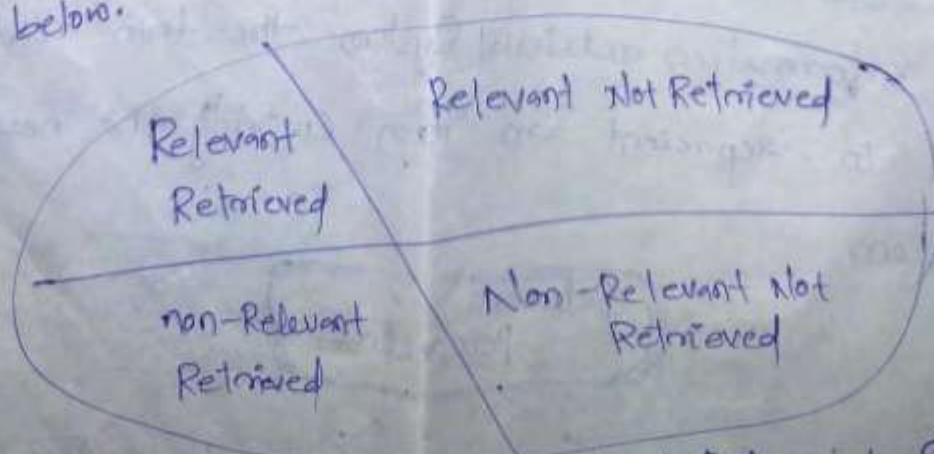


Fig: Effects of Search on total Document Space.



Let us see, what is "Relevant item" and "non-relevant item" (a)

Relevant items:-

It is used to represent an item containing the needed information. - that helps the searcher/user to answering the question.

Non-relevant item:-

Non-relevant items are those items that do not provide any useful information directly.

⇒ The Precision and Recall formally defined as:

$$* \text{ Precision} = \frac{\text{Number - Retrieved - Relevant}}{\text{Number - Total - Retrieved}}$$

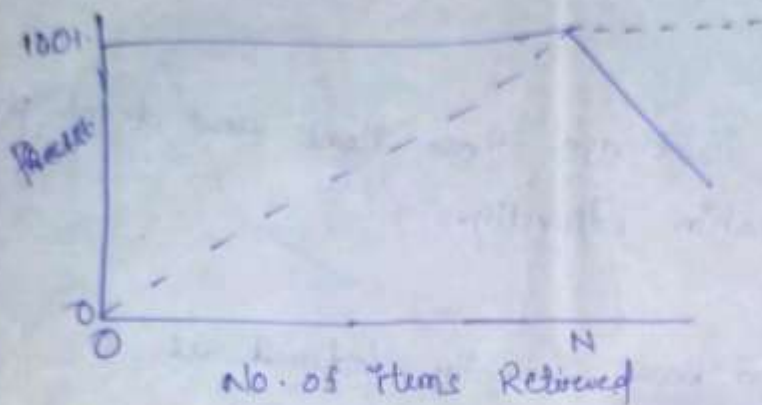
$$* \text{ Recall} = \frac{\text{Number - Retrieved - Relevant}}{\text{Number - Possible - Relevant}}$$

Where,

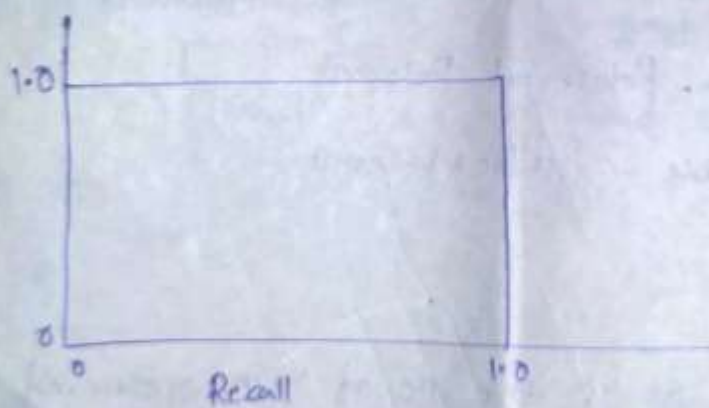
- \* "Number - Retrieved - Relevant" is the no. of items retrieved that are relevant to the "User's" needed search.
- \* "Number - Total - Retrieved" is the total no. of items retrieved from the "Query".
- \* "Number - Possible - Relevant" is the no. of relevant items in the "db" "database".

⇒ In a search of "100%", (or) in 100% search "25%" is Precision (i.e. Precision = 25%) and <sup>the remaining</sup> ~~recall~~ "15%" of the User effort is Overhead on retrieving the non-relevant items.

Let us see the 'Precision' and 'Recall' graph:



(a) Ideal Precision and Recall



(b) Ideal Precision/Recall Graph.

⇒ From the diagram, the basic properties of precision (Solid line) and Recall (dashed line) can be observed.

Here, 'Precision starts off' at 100% and maintains that value as long as relevant items are retrieved.

'Recall starts off' close to zero and increase as long as relevant items are retrieved. (until all possible relevant items have been retrieved)



⇒ Once all "N" relevant items have been retrieved, the only items being retrieved are non-relevant.

⇒ Precision is directly affected by ~~non~~ retrieval of non-relevant items and drops to the number close to zero.

⇒ Recall is not affected by retrieval of non-relevant items and remains at 100%.



Fig. c. Achievable Precision/Recall graph.

⇒ Fig (b) & fig (c) show the Optimal and currently achievable relationships b/w Precision and Recall.

⇒ To understand the implications on fig (c),

⇒ Assume that there are 100 relevant items in the database and from the graph at Precision of .3 (i.e 30%) there is an associated recall of .5 (i.e 50%).

This means there would be 50 relevant items in the hit file from the recall value.

⇒ A precision of 30% means user would likely review 167 items to find the 50 relevant items.

## ⑧ Functional Overview

A total Information Store and Retrieval system is composed of 4 major functional processes:

- ① Item Normalization
- ② Selective Dissemination of Information (i.e. "mail")
- ③ Archival Document Database Search
- ④ Index database Search along with the -Automatic File Build Process that support index files.

### Item Normalization

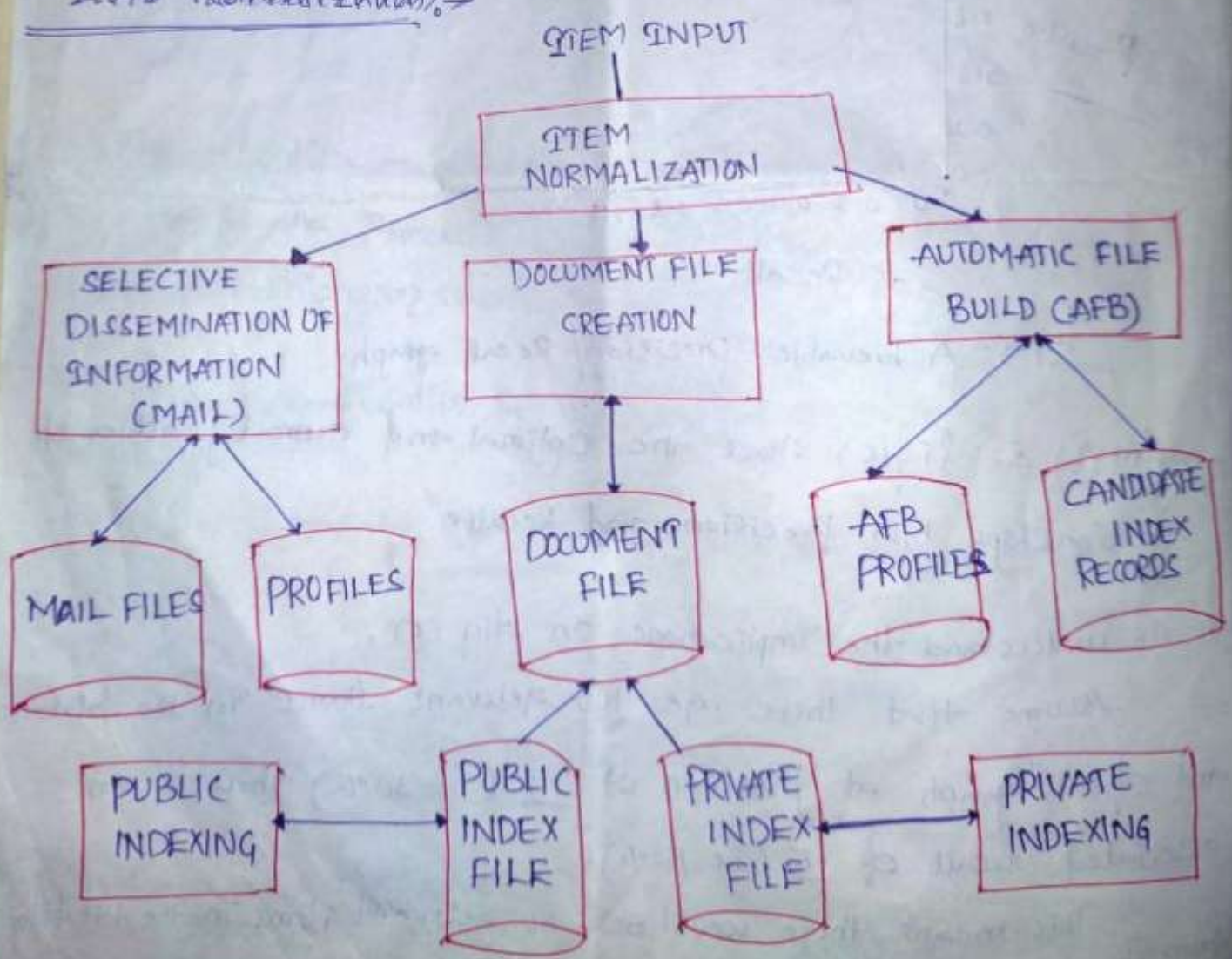


Fig: Total Information Retrieval System.



## ① ITEM NORMALIZATION :-

The first step in any integrated system is to normalize the incoming items to a standard format.

### Item / Item Input :-

The term "Item" is used to represent the "Smallest Complete Unit" that is processed and manipulated by the system.

The definition of item varies how a specific source creates information.

### For Example :-

→ A complete document, such as a book, a newspaper or magazine could be an item. Sometimes the "chapters" and "articles" inside the magazine also consider as an item.

### Normalization :-

It is the process of bringing (or) retrieving something to a Normal Condition (or) state.

⇒ Item Normalization provides "logical restructurizing of the item".

The following "Operations" are performed during the item

- Normalization:

- \* 'Identification' of processing tokens (eg: words).
- \* "characterization" of the tokens.
- \* 'Stemming' of the tokens. (eg: removing word endings).



Now, let us see the text Normalization Process: 14

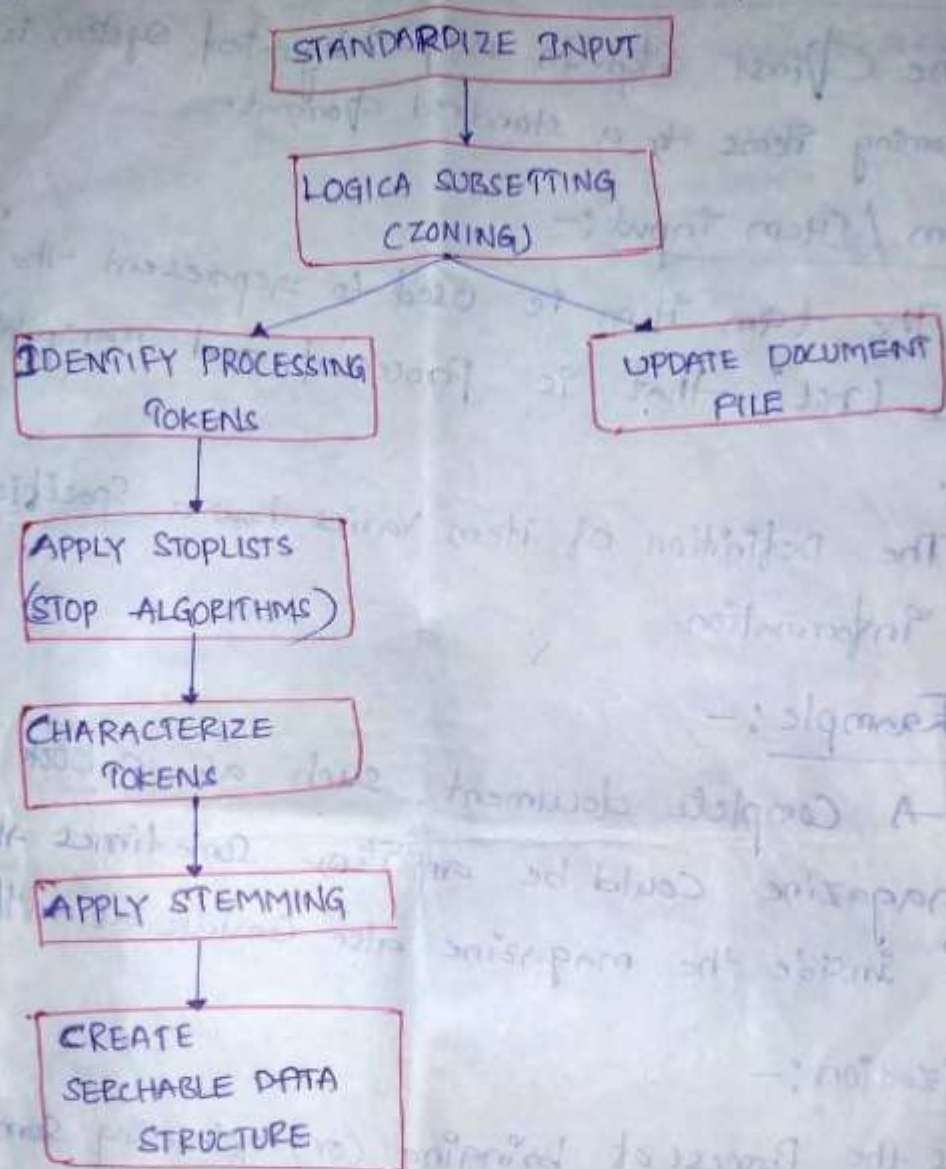


Fig: The Text Normalization Process.

### Standardize Input:

⇒ Standardizing the input takes the different formats of input data and performs the translation so that the result of translation is acceptable by the system.

⇒ A system may have a single format for all the items or allow multiple formats.



for Example :-

Standardization could be translation of foreign languages into "Unicode". Here Unicode is an international standard based upon 16-bit (i.e. bytes) that will be able to represent all languages.

→ Every language has a different internal binary ending for the character in the language. UTF-8  
↳ Unicode Transformation  
format - 8

→ "One standard encoding" that covers English, French, Spanish etc is ISO-Latin.

→ The another encoding for other language groups such as Russian, Japanese, Arabic etc is KOI-7 & KOI-8.

→ Multimedia is an extra dimension to the Normalization process. here KOI-7 → KOI Obmena Informatsiy, 7 bit  
which mean code for "information exchange" 7 bit

In order to normalize the textual input, the multimedia input also need to be standardized.

There are lot of options to standardized Normalization.  
i.e. Input may be Video (or) Audio (or) Image.

### Multimedia Options to the Standardization

Video

Audio

Image

Video :-

→ If the input is "video" then following "digital standards" are ~~being~~ applied.

→ The digital standards are either MPEG-2, MPEG-1, AVI (or) Real Media.

MPEG → Motion Picture Expert Group



① MPEG standards are the most universal standards for higher quality video where RealMedia is the most common standard for lower quality video being used on Internet.

### Audio :-

Audio standards are typically WAV/ (or) Real media (Real Audio).

\* WAV → Wave form Audio file format.

\* Real audio is a proprietary audio format developed by Real Networks in 1995.

Image:- Images vary from JPEG to BMP → Bitmap image file.

### Logical Subsetting, (Zoning):-

JPEG → Joint picture expert Group

→ The second step in the text normalization process is to parse the item into logical sub-divisions called zoning.

→ A typical item is sub-divided into zones and the zones which may overlap and can be hierarchical, such as

Title, Author, Abstract, main text, conclusion and References.

→ In zoning search is restricted to "specific zones" sometimes.

### Index:-

If the user is interested in discussing "Einstein" article then the search should not include the Bibliography, instead of that it include references to articles written by "Einstein".

→ Once search is completed user is efficiently retrieving the needed information.

→ Zoning is performing the two tasks.

i> Identify processing tokens.

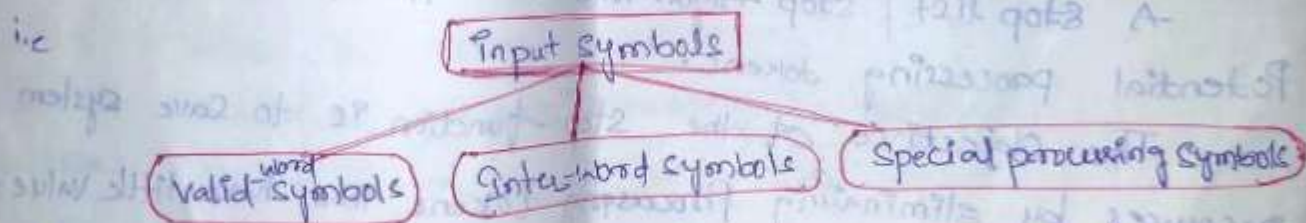
ii> Update Document file.



## Identify Processing Tokens:-

17

- Once standardization and zoning has been completed, let us see two terms "Identify" & "Processing token".
- Information (i.e. words) that are used in the search process need to be "identified" in the system.
  - The term "processing token" is used because a "word" is not the most efficient unit on the search structure.
  - The first step in the identification of a processing token is determining the "word".
  - System determine the words by dividing the input symbols into three classes:
    - i) Valid Word Symbols
    - ii) Inter-word Symbols
    - iii) Special processing symbols.



### Valid Word Symbols:-

A word is defined as a contiguous set of word symbols sharing a common border, bounded by inter-word symbols.

### Inter-Word Symbols:-

In many systems, inter-word symbols are non-searchable and should be carefully selected.

Examples: Semi-colons, blanks, an apostrophe.

### ii) Special processing Symbols:-

There are some symbols that may be required in special processing.

Example: Hyphen (-)

i.e. A Hyphen can be used in many ways,

- \* If it is used at left to the word, means it gives the judgment of the writer. Ex: Bernstein - 84.

- \* If it is used at the end of a line, means indicates the continuation of a word.

- \* If it is used in other place means it links the independent words.

### Apply Stop Lists (Stop Algorithms):-

Next Step,

A stop list / stop algorithm is applied to the list of potential processing tokens.

The objective of the stop function is to save system resources by eliminating processing tokens that have little value to the system.

The best technique to eliminate the majority of these words via stop algorithm is first list out them individually.

Examples of stop algorithms are:

- Stop all the numbers greater than "99999" (this was selected to allow dates to be searchable).

- \* Stop any processing token that has numbers and characters is intermixed.



## Characterize Tokens :-

19

Next step for finalizing the processing tokens is 'Identification of any specific word characteristic'.

The characteristic is used in the system to assist "disambiguation of a particular word".

For example :-

for a word "plane"

- the system understand that it could mean "level or flat" as an adjective.

- \* "aircraft carrier" as a noun.

- \* "The act of ~~smoothing~~ or evening" as a verb.

## Apply Stemming :-

Once the potential processing token has been identified and characterized, most of the systems apply stemming algorithms to normalize the token.

The decision to perform stemming is tradeoff (or like b/w) between "precision" of search and "recall".

For Example :-

The system must keep singular, plural, past tense, possessive etc.

## Create Searchable data structure :-

Once the processing tokens have been finalized, based upon the stemming algorithm, they are used as updates to the searchable data structure.

→ The searchable data structure is the internal representation<sup>20</sup> (i.e. not visible to the user) of items that the user query searches.

## ② SELECTIVE DISSEMINATION OF INFORMATION:

SDI is a concept that was introduced in information science by Hans Peter Luhn in 1958.

SDI is also called as "mail".

The selective dissemination of information (mail) process provides the capability to dynamically compare newly received items in the information system against standing statements of user.

The Mail process is composed of 3 things:

i) Search process

ii) User statements of interest (profiles)

iii) User Mail files.

### i) Search Process:-

After receiving item, it is processed against every user profile.

### ii) Profiles:-

A Profile contains search statements along with list of user mail files.

→ Profiles define all the areas in which a user is interested.

### iii) User Mail files:-

→ When a search statement is satisfied, the item is placed in the "mail files" associated with profile.

→ Items in mail files are typically viewed in the time of receipt order and they are automatically deleted after a specified time period (e.g. after a month etc).



### ③ DOCUMENT DATABASE SEARCH :-

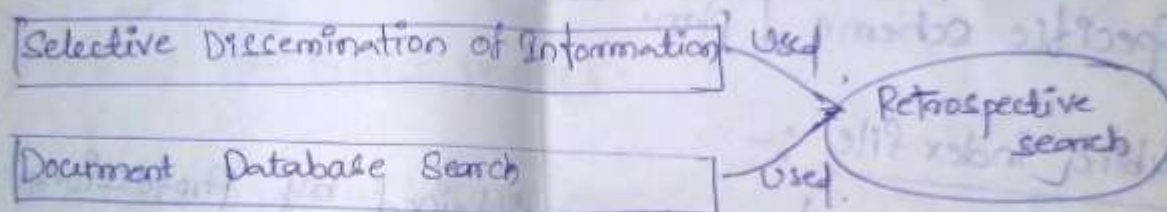
21

→ The document database search process is composed of

- Search process
- Uses entered Queries
- the document database.

→ The document database, which contains all the items been received, processed & stored by the system.

→ The Selective Dissemination of Information and document database search both are used retrospective search.  
i.e



→ The search can be performed in 2 ways

- i.e
1. Perspective Search
  2. Retrospective Search.

#### Perspective Search :

Searching in future and it is time valid.

#### Retrospective Search :-

Searching in past and it is time invalid.

⇒ Document database is very large, it consist of 100's of millions of items.

classified in 2<sup>nd</sup> 22  
Once the Document Database Search is completed then  
it creates the document file.

The Document file is classified into 2 types.

- i) Public Index file (Public indexing)
- ii) Private Index file (Private indexing).

### Indexing :-

It is Originally called as "cataloging".

→ Indexing refers to Organization of data according to specific scheme / plan.

### Public Index file :-

- Public Index files are maintained by Professional Library Services.
- These files have access list that allow any one to search or retrieve data.

~~Private~~ Public index files are used to perform the public indexing.

### Private Index file :-

- Private Index files typically have very limited access.
- Every user can have one or more no. of private index files that leading to a very large no. of files.
- Private Index files are used to perform the private indexing.



#### ④ Automatic File Build (AFB):-

The automatic file build is also called as "Information Extraction".

AFB is classified in 2 types / performed through

i) AFB profiles

ii) Candidate Index Record

#### AFB profiles:-

The index term extraction process are stored in Automatic file Build profiles -

#### Candidate Index Record:-

- When an item is processed it results in creation of Candidate Index Records.

## ① Relationship to Database Management Systems:-

These are 2 Major categories of systems available to process items:

1. IRS
2. DBMS

i.e.

Relationship to

DBMS

→ A "confusion" can be arise, when the "software systems" supporting each of these Applications get confused with the data they are manipulating.

① \* An Information Retrieval System is "software" that has the features and functions required to manipulate "information". items

VS (Versus)

\* A DBMS that is Optimized to handle "structured" data.

② \* Information is "fuzzy text".

The term fuzzy is used to imply the results from the minimal standards (or) Controls on the creators of the text items.

i.e. the minimal standards in terms of Author is.

"The authors is trying to present concepts, ideas and abstractions along with supporting facts."

VS (Versus)

\* The Structured data is well defined data (facts) typically represented by "tables".

i.e. table is associated with "attributes" and attribute contains semantic description.



For example:

→ There is no confusion to what values enter in a specific database record.

⇒ we have different attributes like "employee Name" and "employee salary".

③ \* On the other hand, if two different people generate "an abstract" for the same item, they can be different.

i.e. One abstract may generally discuss the most important topic in an item.

Another abstract may specify the details of many topics.

### (VS) Versus

\* With structured data a user enters a specific request and the result returned with the desired information.

The results are frequently tabulated and presented in a report format for ease of use.

④ \* Search of "information" items in IR & DBMS.

\* A search of information items has a high probability of not finding all the items a user is looking for.

so that <sup>in search</sup> user enter the additional items of interest, the process is called "iterative search". doing something again & again

→ An information retrieval system gives the capabilities to the user in finding the relevant items such as "relevance feedback".

(26)

→ The Result from the Information System search are presented in "relevance ranked Order".

\* The Confusion Comes, when DBMS software is used to store - "Information".

This is easy to implement, but the system lacks the "ranking and relevance" feed back features that are critical to ~~the~~ an Information System.

→ It is also possible to have structured data used in an Information System, (Such as TOPIC).

finally,

from a Practical Standpoint, The Integration of DBMS's and Information Retrieval System is Very Important.

(\*) Digital Libraries and Data Warehouses:-

Two Other Systems frequently described in the context of Information retrieval are "Digital Libraries and Data Warehouses" (Data Marts).

There is significant Overlap between these two Systems

Digital Libraries:-

They are also called as "Online Library" (or) "Internet Library" (or) Digital Repositories.

A digital library is a library in which collections are stored in "digital format".



An electronic (or) Digital Library is a type of Information retrieval System (IRS).

It provides the Architecture to

- \* Model
- \* Map
- \* Integrate &
- \* transform scattered information in digital documents.

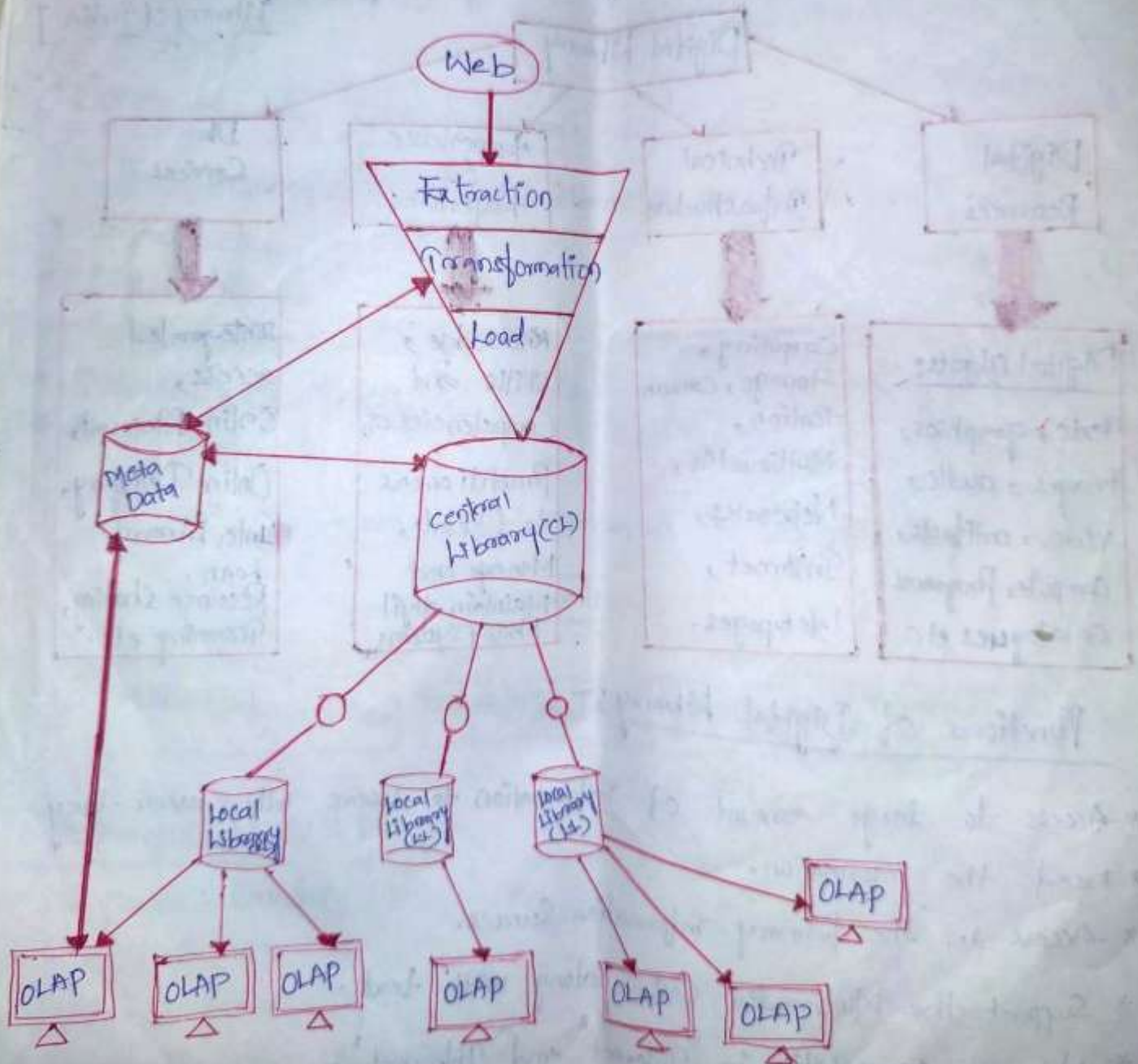


Fig: Digital Library

[ : OLAP → Online Analytical Processing ]

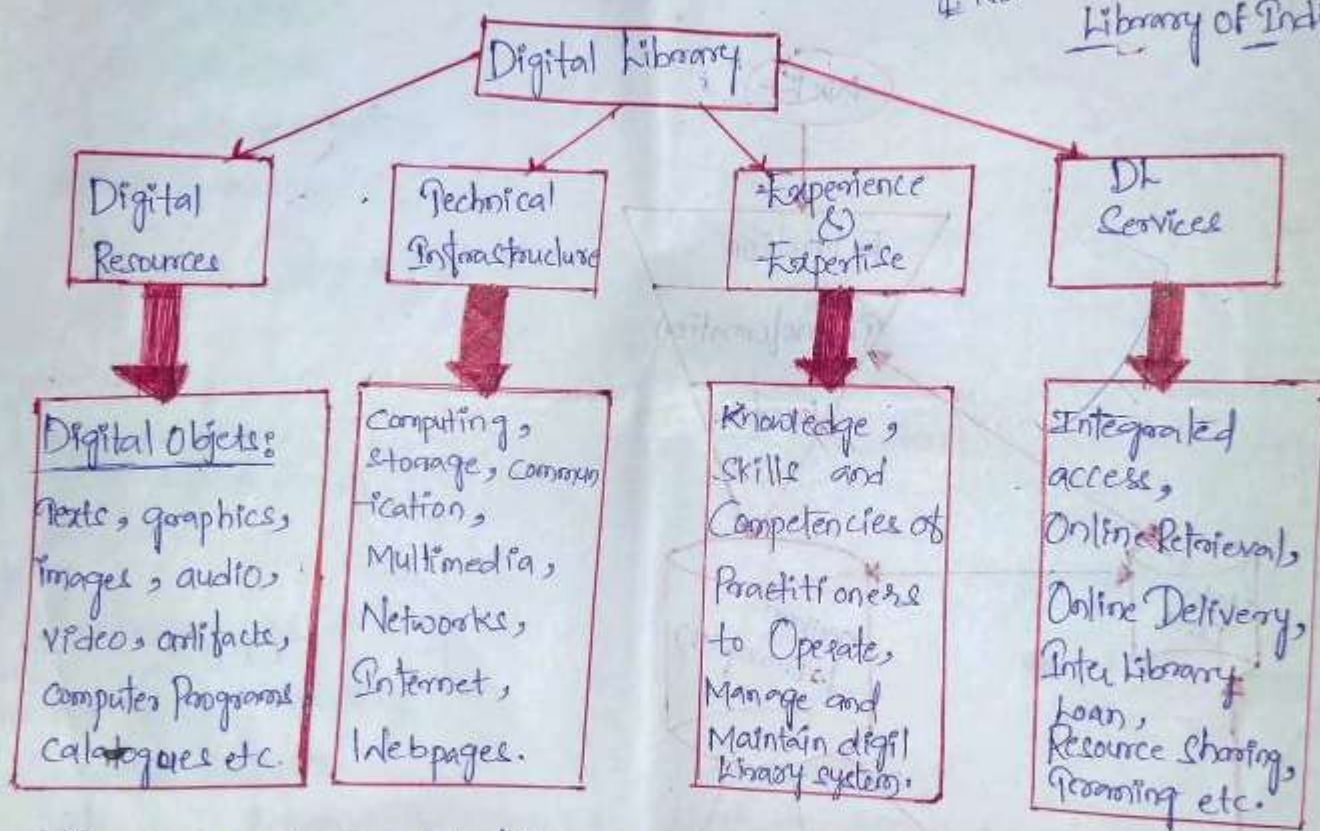


→ The term DL's (Digital Library's) was first popularized by the NASA Digital Libraries Initiative in 1994.

→ Digital Library is not a single entity. It requires technology that link the resources of many collections.

The links between the ~~various~~ digital libraries and their resources are transparent to users shown in following figure.

[NDLA → National Digital Library of India]



### Functions of Digital Library:-

- \* Access to large amount of information to users whenever they need the information.
- \* Access to the primary information sources.
- \* Support the Multimedia content along with text.
- \* Network Accessibility in "Internet" and "Intranet".
- \* Client-Server Architecture
- \* Advanced Search & Retrieval.
- \* Integration With Other Digital Libraries.



## Purpose of Digital Library:-

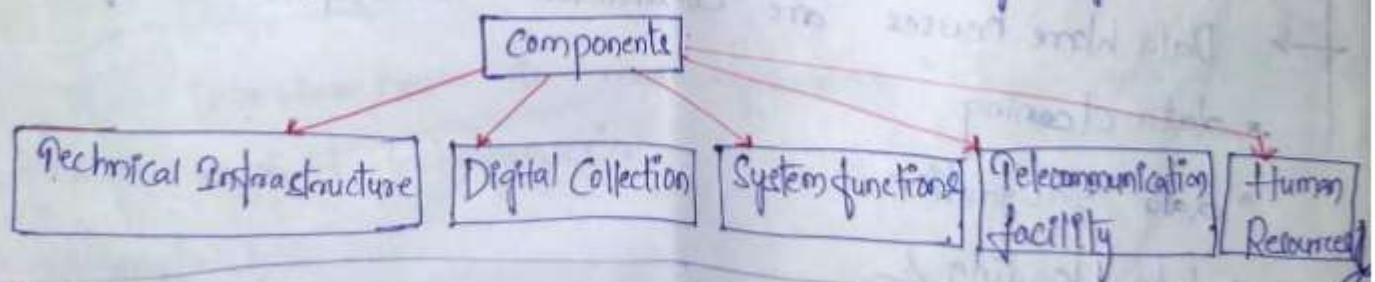
- It expedite the systematic development of procedures to collect, store and organize information in digital form.
- It Promote efficient Delivery of information economically to all Users.
- It Strengthen Communication and collaboration "between" and "among" educational Institutions.
- It takes "leadership Role" in the generation and dissemination of knowledge.

## Components of Digital Libraries:-

The Components of Digital Libraries are:

- \* Infrastructure
- \* Digital Collection
- \* System functions
- \* Telecommunication facility
- \* Human Resources.

You can Represent these Components in following way.



## Advantages of DL's:-

- No physical Boundary
- Multiple Accesses
- Universal Accessibility
- Round-the clock Availability
- Added Values
- Enhanced IR.

### Limitations:

- \* Lack of Screening (or) Validation.
- \* Lack of Preservation of "Best in class".
- \* Lack of Preservation of "fixed Copy".
- \* Lack of Preservation of scientific research.  
i.e. for the record and duplicating scientific research.

### Data Warehouse (or) Data Mart: -

→ A data warehouse is a repository of information collected from multiple sources, stored under a 'Unified Schema' and that are usually resides at a 'Single site', (or)

~~The process of collecting information from multiple sources~~  
~~is called~~.

→ A data warehouse refers to a place where data can be stored for useful mining.

where "mining / data mining" refers to process of extracting useful data from the databases.

→ Data Ware houses are constructed via a process of

- \* data cleaning
- \* data Integration
- \* data loading
- \* periodic data refreshing.

→ The "Lesson" "Data Warehouse" came from the Commercial Sector then Academic Sources.



Goal :-

⇒ The 'goal' of the data warehouse is to provide the critical information to decision makers so that they can answer the future queries.

⇒ The data warehouse consist of the "data" containing information - directory that describes the contents and meaning of data being stored.

Working:

An "Input function" that captures data and moves it into the data warehouse.

↓  
The "Search data & manipulation tools" that allow the user to access & analyze the warehouse data.

↓  
The "Delivery Mechanism" export the data to other - warehouses, data marts and external systems.

⇒ Focus of Data Warehouse:

The Data Warehouse is more focused on "Structured data" & "decision support technologies".

For Example :-

The typical framework for Construction and Use of - data warehouse for all electronics.

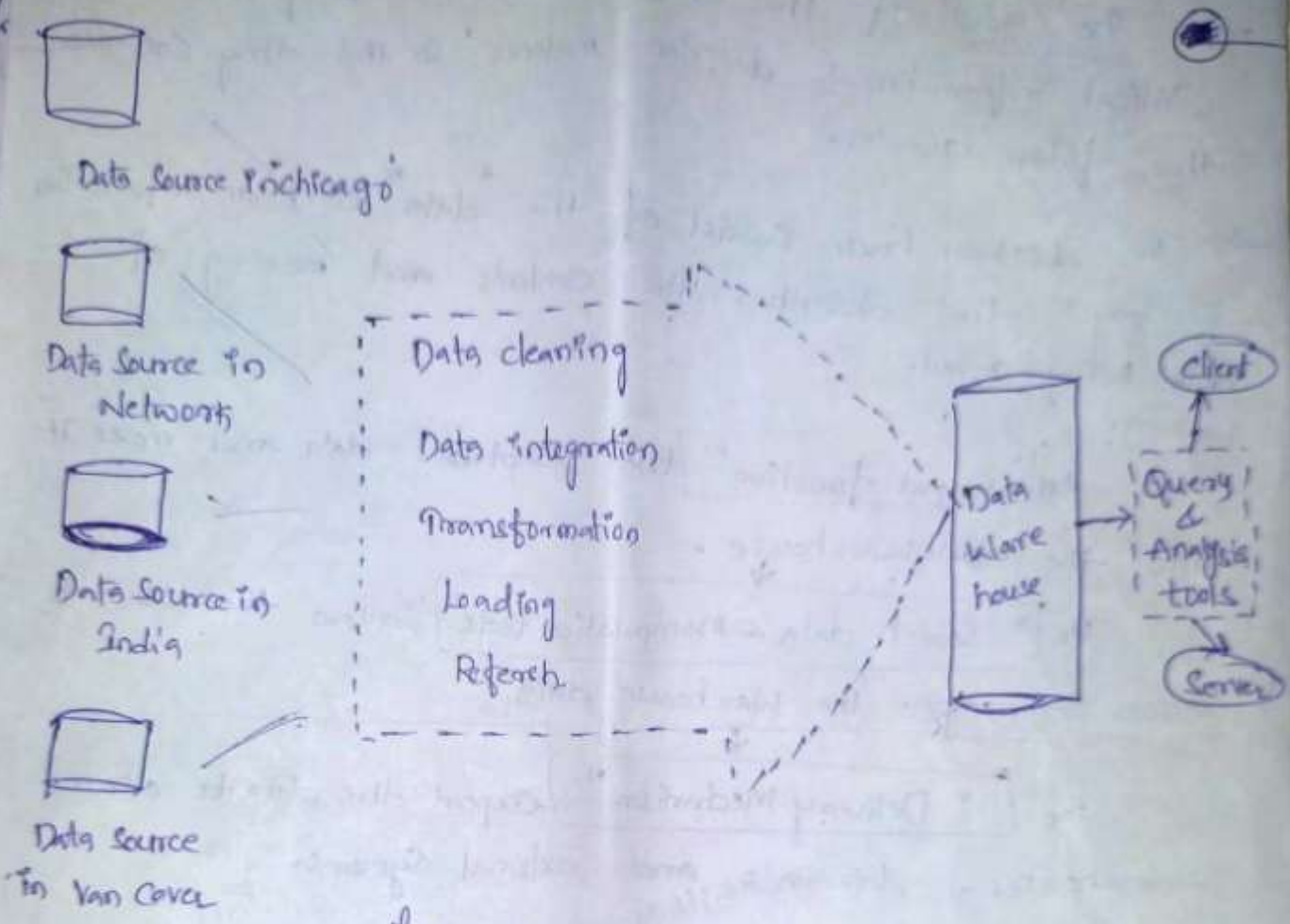


Fig (a) Data-Ware house.

Data mining :-

- The term Data mining is also called "KDD" Knowledge Discovery in Data Base
- KDD is the process of ~~converting~~ discovering valuable information from a collection of data. (or)
- The process of converting the raw data into useful information.



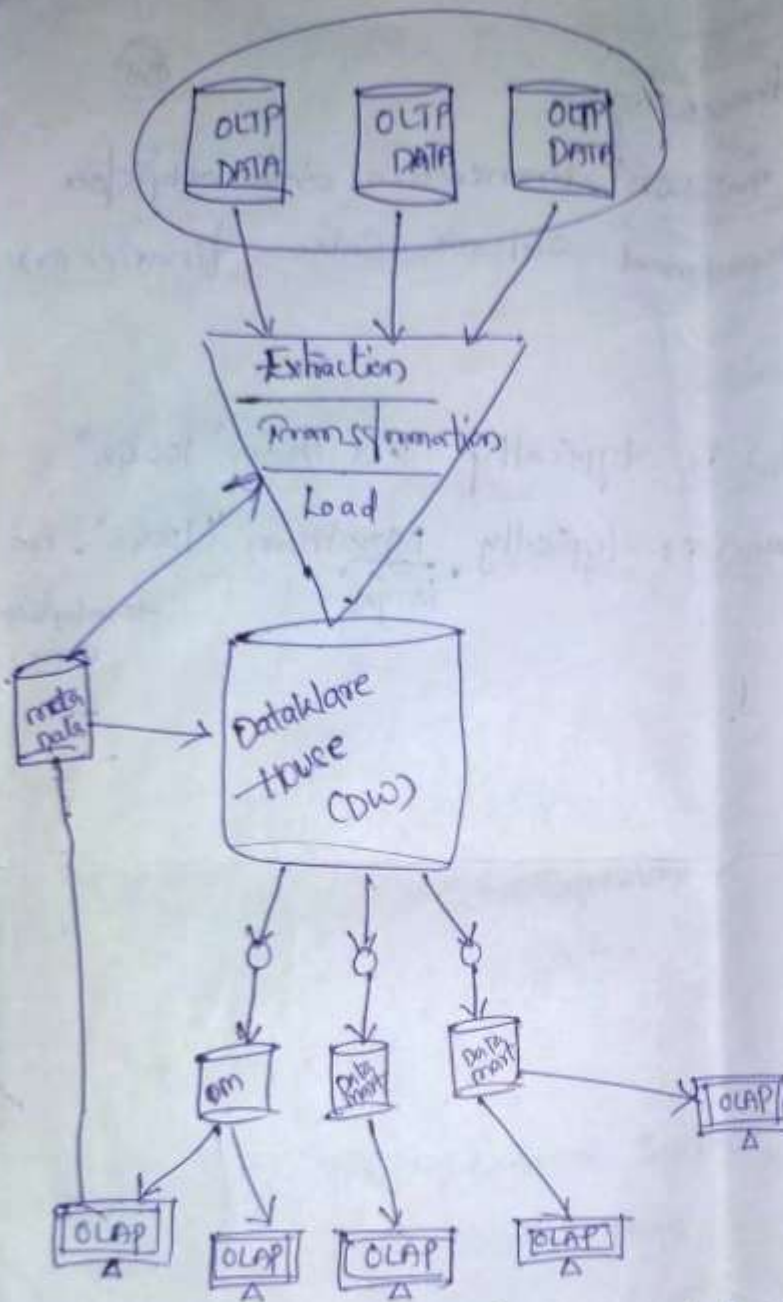


Fig (b) Data Warehouse.

OLAP:-

Online analytical processing

→ This Approach is used to answer Multidimensional Data model.

## \* OLTP :-

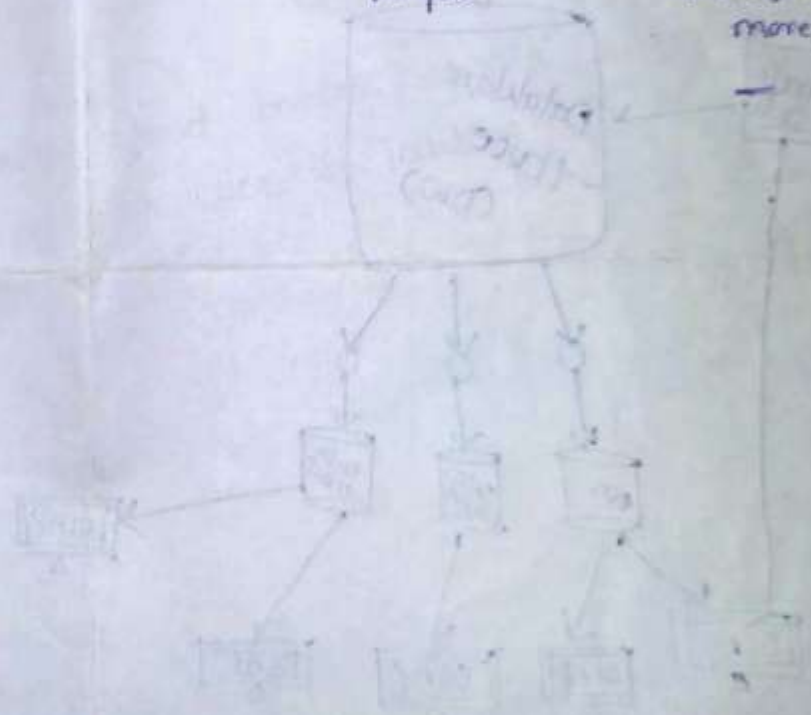
(34)

### Online Transaction Processing

→ It "captures and maintains" transaction data only for a specific department such as Sales, finance (or) HR (Human Resources).

→ The DM (Data Mart) is typically less than "100GB"

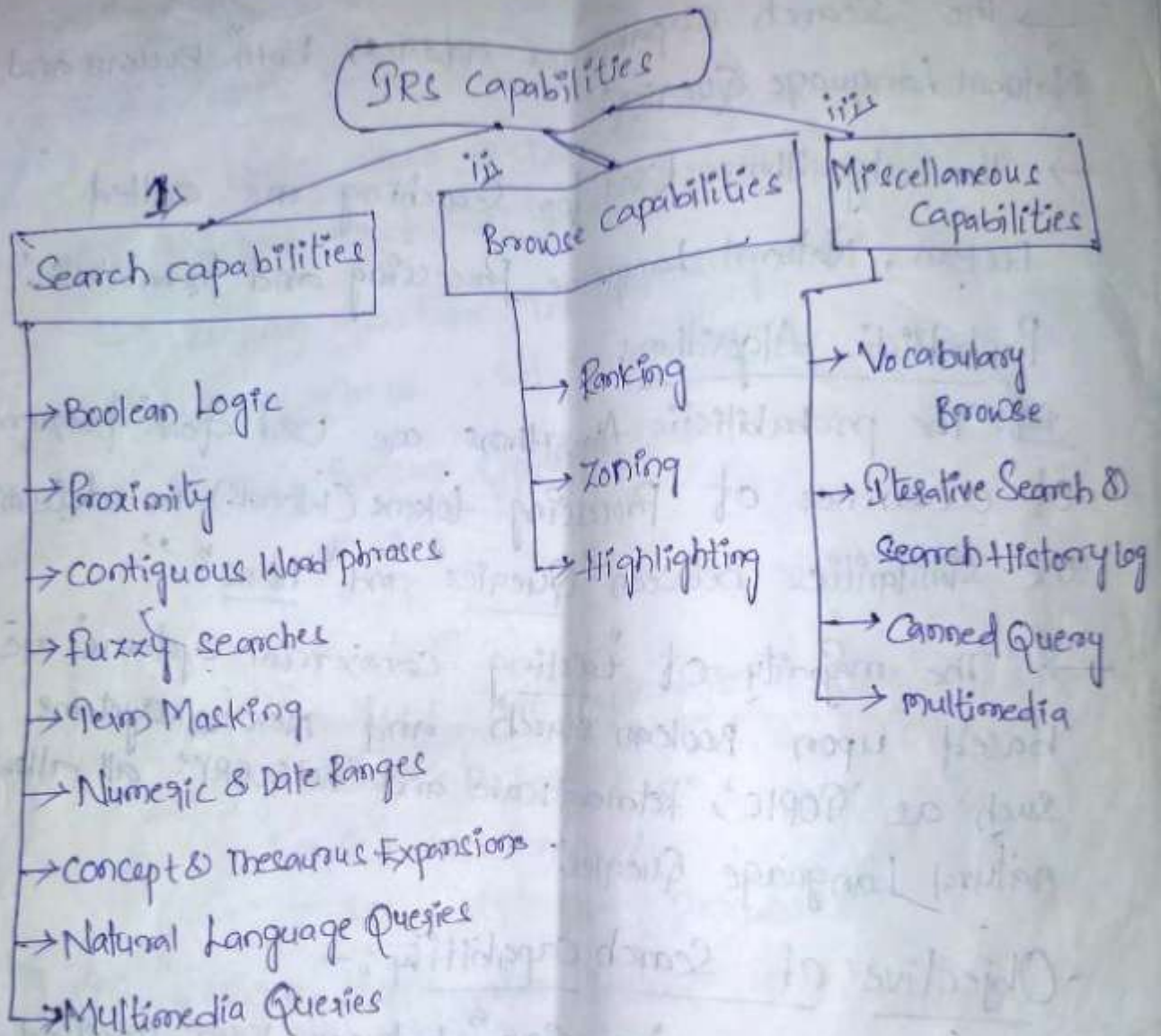
→ The DW (Data Warehouse) is typically more than "100GB". i.e. larger terabytes or more.





## IRS Capabilities

(35)



In Information Retrieval System, search and browse capabilities are play crucial role to assist the user in locating the relevant items.

## → Search Capabilities :-

(36)

→ The Search Capabilities address both "Boolean and Natural Language Queries".

→ The algorithms used for searching are called "Boolean, Natural Language Processing and Probabilistic".

### Probabilistic Algorithm:

Def: The probabilistic Algorithms are used for "frequency of occurrence of processing tokens (words) in determining the similarities between "Queries" and "Items"."

→ The majority of "Existing commercial systems" are based upon Boolean search and newer systems such as "TOPIC", "Retrievalware" and "INQUERY" all allow the natural Language Queries.

### Objective of Search capability :-

It allow the "mapping" between User specified need and the items in the information database as a ~~result~~ <sup>and</sup> information database will answer the need.

→ One of the concept that has been occasionally in commercial systems are "Weighting".

### Weighting :-

The process of holding the location and ranking of relevant items.



IRs provides the different search capabilities (37)  
They are:

### → Boolean Logic :-

- Boolean logic allows a user to logically relate multiple concepts together to define what information is needed.
- The Boolean functions in processing tokens are identified any where within an item.
- The Typical Boolean Operators are :

\* AND

\* OR

\* NOT

These Operations are implemented using

→ Set Intersection

→ Set Union &

→ Set Difference procedures.

→ The Boolean logic also include the "Order of Boolean Operations" and "default precedence order (or) no precedence Order".

### Order of Boolean Operations :-

Placing the portion of search statement in parentheses are used to specify the Order of Boolean Operations.

### default precedence order Operations :-

If the parentheses are not used the system follows a default precedence Ordering of Operations.

Ex:- typically NOT then AND then OR.

A Special type of Boolean Search is called "M of N" logic.  
 i.e. A set of possible search "terms", that are identified and accepted by any item that contains subset of the terms.

For example: -

Find any item containing any two of the following terms: "AA", "BB", "CC"

This can be expanded into a "Boolean Search" that performs an AND operation between all the combinations of two terms and OR operation to results together i.e.



i.e.  $(AA \text{ AND } BB)$  term 1  
 $(AA \text{ AND } CC)$  term 2  
 $(BB \text{ AND } CC)$  term 3

AND Operation b/w all combination.

$(AA \text{ AND } BB) \text{ OR } (AA \text{ AND } CC) \text{ OR } (BB \text{ AND } CC)$  OR operation for result.



## Example: - Use of Boolean Operations: -

### Search statement

\* COMPUTER OR PROCESSOR NOT MAINFRAME

\* COMPUTER OR (PROCESSOR NOT MAINFRAME)

\* COMPUTER AND NOT PROCESSOR OR MAINFRAME

### System Operation

→ select all items discussing "computers" and/or "processors" that do not discuss mainframes.

→ select all items discussing computers and/or the items that discuss processors and do not discuss mainframes.

→ select all items that discuss computer and not processors OR mainframes in the item.

Fig: Use of Boolean Operations.

## ii) Proximity: -

Proximity is used to restrict the distance allowed within an item between two search terms.

The typical format for proximity is:

Term1 within "m" Units of Term2

(or)  
TERM1 with in m units of TERM2

Here "m" is indicating the "distance operator".

→ The distance operator m is an integer number and units are in characters, words, sentences and/or paragraphs.



## Proximity Operators:

- \* Direction Operator → (i.e. before/after)
- \* ADJ (Adjacent) Operator → (forward only direction i.e. one direction)
- \* Distance is set to zero → (within same paragraph/unit)

### \* Direction Operator:-

Some times, the proximity relationship contains <sup>"</sup>Direction Operator<sup>"</sup> indicating the direction (before/after) that the second term must be found within no. of (number of) units specified.

For Example:-

#### Search Statement

- \* "United" within five words of  
"American"

#### System Operation

→ It would "Hit on",  
"United states and American Interest",  
"United Airlines and American Airlines",  
It would not "Hit on",  
"United states of America and the American dream".

### \* ADJ (Adjacent) Operator:-

A special case of ~~prosign~~ proximity operator is the Adjacent ADJ operator, that normally has a distance operator of "one" <sup>direction</sup> i.e. forward only direction.

For Example:-

#### SEARCH STATEMENT

- \* "Venetian" ADJ "Blind"

#### SYSTEM OPERATION

→ It would find the items that mention a "Venetian Blind" on a window but not items discussing a "Blind Venetian".



\* Distance is set to zero :-

- An another special case <sup>where</sup> "distance is set to zero" meaning "with in the same semantic unit".

for Example :-

SEARCH STATEMENT

"Nuclear" within zero paragraphs of "clean-up".

SYSTEM OPERATION

It would find the items that have "nuclear" and "clean-up" in the same paragraph.

iii) Contiguous Word phrases :-

→ In, A Contiguous Word phrase is two or more words that are treated as a "single semantic unit".

for Example :-

"United states of America"

i.e It is "four" words that specify a search term representing single semantic unit (i.e "country").

→ The single semantic unit satisfies all the proximity operators i.e AND, Direction operator & Distance <sup>is set to</sup> zero operator.

→ A CWP is specified in both ways

i.e \* Query

\* a special Search Operator.

for Example :-

A Query could specify "manufacturing" AND "United states of America".

i.e

Which returns any item that contains the word  
"manufacturing" and the contiguous words "United States of America"

→ The Proximity and Boolean Operators are "Binary Operators"  
but Contiguous Word phrases are "N"ary Operators  
Where "N" is the number of words in the CWP.

→ The Contiguous word phrases are called "Literal strings"  
in "WAIS" and ~~Exact~~ phrases in "Retrieval Ware".

Where WAIS:- WAIS is a standard for

- \* indexing

- \* string &

- \* Retrieving the Information from source

- document that can be located anywhere on the Internet.

→ In WAIS multiple Adjacency (Adj) Operators are used to  
define a Literal string.

Ex:- "United" ADJ "states" ADJ "of" ADJ "America".



### iv) Fuzzy search:-

Fuzzy search provide the capability to "locate the spellings of words" that are similar to the entered search term.

→ This function is primarily used to compensate for errors in spellings of words.

→ A fuzzy search on the term "Computes" would automatically include the following the words from the information database:

"Computes"

"computer"

"compute"

"Computer"

"compute".

→ Maximum utilization of fuzzy search in the system that accept is (OCR) optical character Read.

OCR :- In OCR a hardcopy item is scanned in to a binary image. (or)

Process of scanning the hardcopy into binary image is called OCR.

## Term Masking :-

- Term Masking is the ability to expand a query term by masking a portion of the term and that maps to the unmasked portion of the term.
- The value of term masking is much higher in systems that do not perform "stemming".
- There are two types of search term masking:

ie



Sometimes the fixed-length and variable-length

- Sometimes both "fixed and variable length" are called "don't care" functions.

### \* fixed length Masking :- (\$)

- fixed length Masking is a "single position Mask".
- It Mask out any symbol in a "particular position".
- fixed length term Masking is not frequently used.
- It is typically not critical to the system.
- It is indicated by "\$".

for example :-

### SEARCH STATEMENT

\* Multinational

### SYSTEM OPERATION

Matches "multi-national",  
"multinational", "multinational"  
But it does not match the  
"multi national" since it is two  
Processing tokens.



\* Variable length Masking :- (\*) -

→ Variable length Don't-care functions allows masking ~~at~~ of any number of characters within a processing token.

→ The Variable length Masking may be in the front, at the end, at both ends.

i.e. \* Subfix Search :-  
The front end means, if the masking is performed at the "Starting position" of the processing token is called "Subfix Search".

i.e. "\* COMPUTER" → Subfix Search.

\* Prefix Search :- If the masking is performed at the "ending position" of the processing token is called Prefix Search.

i.e. "COMPUTER\*" → Prefix Search.

\* Imbedded String Search :- If the masking is performed at "both ends" of the processing token is called "Imbedded String Search".

i.e. "\* COMPUTER\*" → Imbedded String Search

→ The Variable length Masking is indicated by "\*".

For Example :-

| SEARCH STATEMENT | SYSTEM OPERATION                                                            |
|------------------|-----------------------------------------------------------------------------|
| → *computes      | * It Matches, "microcomputer",<br>"minicomputer" cor) "computes"            |
| → Comput*        | * It Matches,<br>→ "Computers",<br>→ "computing",<br>→ "Computes".          |
| → *Comput*       | * It Matches,<br>→ "micro computers",<br>→ "minicomputing",<br>→ "Compute". |

Fig. Term marking

VI) Numeric and Date Ranges :-

- Term Marking is useful when applied to words, but does not work for finding ranges of numbers (or) Numeric dates.
- To find the numbers larger than "125" using a term "125\*" will not find any number.
- As a part of system Normalization process, characterizes words as "numbers (or) dates".

For Example :-

\* for Numbers and Dates :

if "125-425" means finding the numbers <sup>range</sup> in between two numbers.



ix | 4/2/1993 — 5/2/1995 means finding date range <sup>(41)</sup>  
in between given dates.

\* for Infinite Ranges :-

for representing the infinite ranges between the numeric dates it uses special type of operators

i.e.  $[ > (or) \geq (or) < (or) \leq ]$

Example : " $> 125$ " means Larger Numbers above the 125.  
" $\leq 233$ " means Smaller (or) equal to 233.

Vii) Concept & thesaurus Expansion :-

\* Concept :-

→ A concept <sup>class</sup> is a "tree structure" that expands each meaning of a word into potential concept.

→ Concept classes are sometimes implemented as a "Network structure" that links the word stems.

→ Concept class representation:

The Concept class representation, assist a user who has minimal knowledge of "concept domain" so that user can easily expand related concepts.

for Example:-

(48)

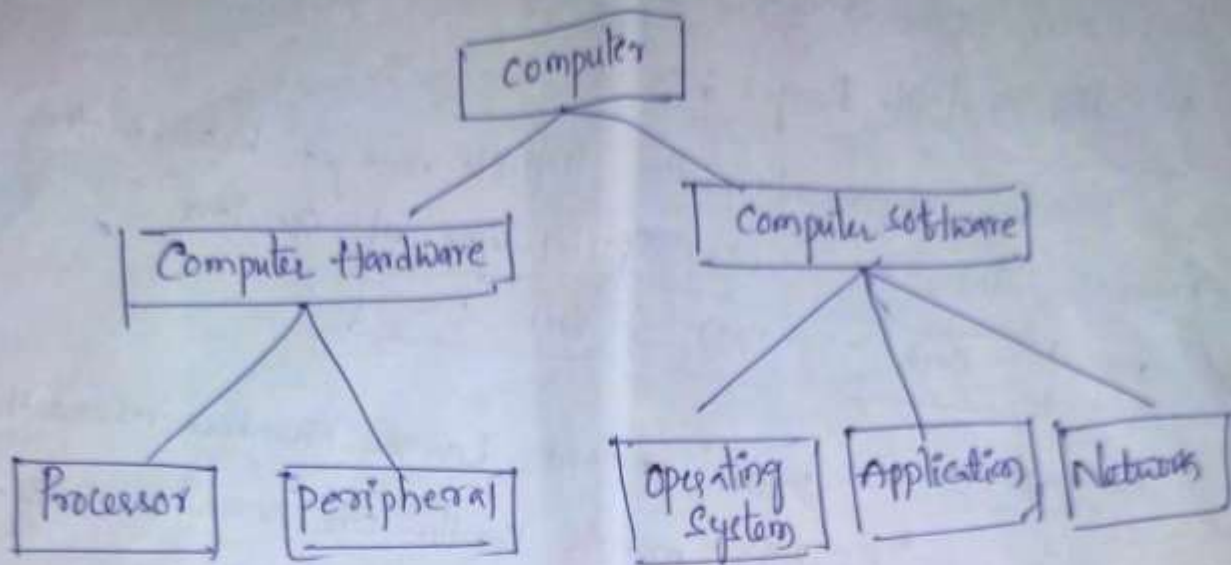


Fig:- Hierarchical Concept class structure for "Computer".

Thesaurus Expansion:-

A Thesaurus is typically a one-level (or) two-level expression of a term to other terms that are similar meaning.

for Example:-

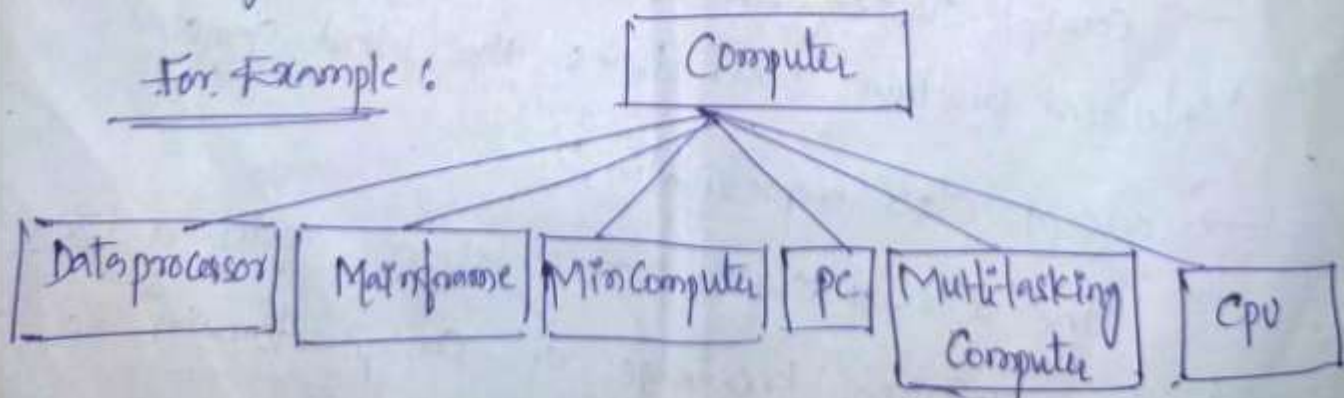


Fig: Thesaurus for term "Computer".



### Viii) Natural Language Queries:-

→ A Natural language query consists only of "normal terms" in the user's language, without any special syntax (or) format.

Natural Language Queries allow a user to enter a free statement that describes the information, <sup>ie</sup> that the user wants to find.

The longer phrase returned the more accurate results.

The most difficult logic case associated with Natural Language Queries is the ability to specify "negation" in the search statement.

Example :- i.e. Example of Natural Language Query:

find for me all the items that discuss "oil reserves" and current attempts to find new oil reserves. Include any items that discuss the international financial aspects of the oil production process. Do not include items about the oil industry in the United States.

→ The minimized input for above example is

oil reserves and attempts to find new oil reserves, international financial aspects of oil production, not united states oil industry.

(50)

The Boolean query for same search statement is:

$$\left( \text{"locate"} \text{ AND "new" and "oil reserves"} \right) \text{ OR } \left( \text{"international"} \text{ AND "finance"} \text{ AND "oil production"} \right) \text{ NOT } \left( \text{"oil industry"} \text{ AND "United States"} \right).$$

### ix) Multimedia Queries :-

The User Interface becomes far more complex, with the introduction of the availability of multimedia items.

- The image <sup>could be</sup> used to search the images that are part of an item.
- They also could be used to locate specific scene in a video product.



## 2. Browse Capabilities :-

Browse capabilities provide the user, capability to determine which items are of interest and select those items to be displayed.

There are 2 ways of "displaying" summary of items that are associated with Query.

i.e i) Item line status

ii) Data Visualization.

### Item line status :-

from the Summary,

The user can select the "specific items" and "zones within the items" for display.

### Data Visualization :-

The system allows easy transitioning between the "Summary display" and "review of specific items".

Browse Capabilities are of 3 types.

i) Ranking

ii) Zoning

iii) Highlighting.

## i) Ranking :-

(Foot)

⇒ The Introduction of Ranking based upon

- \* predicted relevance values,
  - \* the status summary displays, &
  - \* the relevance score associated with items, along with
- brief description.

### Relevance score:

→ The Relevance score estimate the item that satisfies the search statement.

→ The Relevance Scores are Normalized to a value between  $[0.0 \text{ and } 1.0]$ .

→ Here

"The Highest Value of  $1.0$ " means the system is sure that the item is relevant to search statement.

The Lowest Value of  $0.0$  means "not relevant".

⇒ In addition to ranking based upon the "characteristics of item" and "the database".

⇒ Ranking uses "collaborative filtering" for selecting and Ordering outputs.

is a technique that can filter out the items that user might like.

In this case, when user reviewing the items it provide feedback to the system on the relative value.

The "collaborative filtering" has been very successful in sites such as AMAZON.com, Movie-Finder.com are deciding what products to display to users based on their queries.



## ii) Zoning :-

→ The process of dividing the standard input into logical sub-divisions is called zoning / zones.

Zone : The "title" is frequently called as zone.

It provide the additional information to the user with "relevance weight" to avoid selecting the non-relevant items for review.

→ The Zoning is used for in minimizing what an end user needs to review from a hit item in Passage based Search and Retrieval.

## iii) Highlighting :-

→ Highlighting indicates how strongly "the highlighted word" participated in the selection of the item.  
So that user quickly focus on potentially relevant

parts of the text.

→ Highlighting has been useful in Boolean systems to indicate the cause of retrieval.

i.e. it provide the direct mapping between the "terms in the search" and "terms in the item".

→ The highlighting may vary by introducing colors and intensities to indicate the importance of a particular word.



### 3. Miscellaneous Capabilities :-

The additional functions that facilitate the user's ability to Input Queries, and reduce the time taken to generate the queries.

The Miscellaneous Capabilities are 4 types:

- i) Vocabulary Browse
- ii) Iterative Search and Search History Log
- iii) Canned Query
- iv) Multimedia.

#### i) Vocabulary Browse :-

Vocabulary Browse Provides the capability to display in alphabetical Sorted Order Words from the document database.

All "unique words / processing tokens" in the database are kept in Sorted Order along with a number count of uniqueness in which the word is found.

The User can enter "a word" (or) "word fragment" and the system will begin to display the dictionary around the entered text.

i.e. "The" system indicates what word fragment the user entered, and then it alphabetically displays other words found in the database.

The User can continue scrolling in either direction of reviewing additional terms in the database.

Exp: It helps the user to determine the impact of using a fixed (or) Variable length mask on a search term, and potential mis-spellings.



ie computer in effect is searching for  
 "Computation" (or) "compulsive" (or) "compulsory".  
 for potentially <sup>some one entered the</sup> misspelled word "computer" <sup>that</sup> really meant  
 "computer".

For Example :-

Perm: Comput

| TERM           | OCCURRENCES |
|----------------|-------------|
| ie<br>* Comput | 265         |
| * Computation  | 1245        |
| * Compute      | 1           |
| * Computer     | 101800      |
| * Computers    | 18          |
| * Computerize  | 29          |
| * Computes     |             |

Fig: Vocabulary browse list with entered term "comput".

### iii Iterative Search and Search History Log :-

"Iterative searching" and "Search history log" summarize  
 previous search activities so that user access the ~~logs~~  
 previous results from the current user session.

### Iterative Search :-

→ The Results of the previous search can be used to create a new query, rather than typing a complete new query. This has same effect as taking the Original Query. i.e. "result is same as Original Query" by adding the additional "AND" condition to the search statement.

This process is called "Iterative Search".

### Search History Log :-

→ The Search History Log is the capability to display "all the previous searches" that were executed during the current session.

→ It provide the facility to locating the previous searches as starting points for new searches.

### iii Canned Queries :-

→ The capability to "Name" a Query, and "Store" it to be retrieved, & "executed" during a later user session

is called Canned Query/ Canned Queries.

→ The Canned Queries is also called as "Stored Queries".

→ Queries that start with a Canned Query are significantly larger than ad hoc Queries.



#### iv. Multimedia :-

(58)

Multimedia introduces new challenges, for potential users, that satisfy the discovered query, and multimedia techniques are used for display them.

The 'transcribed text' used as navigation technique through the audio.

i.e. they appropriately label this new paradigm

What You See Is (Abstract) What You Hear

(WYSIAWYH).

→ The transcribed text could be used as an index into ~~future~~ future retrieval of the audio source.

i.e. The total information provided by the original speech.

## Unit-II Cataloging and Indexing

58

### Cataloging :-

- Cataloging is a process that creates 'metadata' and represents 'information sources' such as books, sound recordings, moving images etc.
- The Cataloging provides the information such as
  - \* the name of the creator,
  - \* the title and the subject term that describes the source, through creation of a bibliography record.

### Indexing :-

The indexing is one of the most important process in the IR system.

The transformation from the received item to searchable data structure is called indexing.

Indexing helps to retrieve the information "accurately" and "efficiently".

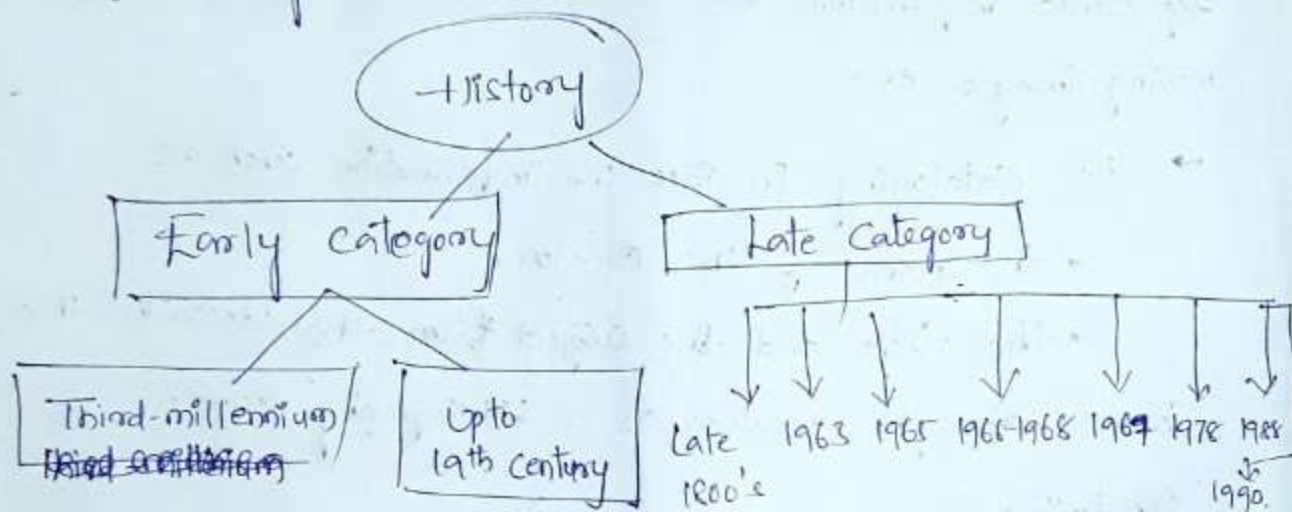
⇒ Indexing Originally called as "cataloging" and it is the oldest technique for identifying the contents of items.

→ Objective: The Objective of the cataloging is give the access points to a collection, that are expected and most useful to the users information.



## \* History and Objectives of Indexing:

The Basic information required for an item is "What is the item" and "What it is about", has ~~been~~ not change over the centuries.



### Early Category :-

As early as,

In third-Millennium, in "Babylon"

The Libraries of "Cuneiform tablets" were arranged by the Subject (By Hyman - 89).

### Up to 19th century :-

Up to 19th century there was little advancement in "Cataloging".

i.e. The advancement is only change in the methods used to represent the Basic Information.

## Late category :-

\* In the late 1800 :-

The subject indexing became "hierarchical"  
Eg: Dewey Decimal System (DDS)  
(COI)  
Dewey Decimal Classification (DDC).

## DDS/DDC :-

The Dewey Decimal System is the World's most widely used way to organize Library collections.

i.e The Dewey Decimal System used by the libraries to arrange "books" via "Subject".

→ DDC allows new books to be added to a library in their appropriate location based on subject.

It was first published by "Melvil Dewey" in 1876, in United States.

\* In 1963 :-

"The Library of Congress initiated".

i.e A study on the computerization of bibliographic surrogates.



## Library of Congress (LC) :-

(62)

(61)

The LC is "research library" that officially serves the "United States Congress".

The Library of Congress "collections are universal" not limited by subject (or) format (or) national boundary.

It includes "research materials" from all parts of "the world" and in more than 430 languages.

The Library of Congress

→ "Established" in April 24, 1800 and

→ "Located" in Washington.

\* In 1965 :- "Catalog system - DIALOG"

The earliest commercial cataloging system

"DIALOG", which was developed by "Lockheed Corporation".

for NASA.

\* From 1966 - 1968 :-

The Library of Congress ran the "MARC I pilot project".  
MARC → Machine Readable Cataloging

It standardizes the structure,

\* contents and

\* coding of bibliographic records.

\* In 1969 :-

The system (Congress Library System) became Operational in 1969.

\* In 1978 :-

The DIALOG became commercial, with three government files of indexes to technical publications.

\* In 1988 :-

The "DIALOG" was sold to "Knight-Ridder" at that time "DIALOG" contained "320" Index databases used by over "91,000" subscribers in 86 countries.

\* In 1990 :-

The significant reduction occur in cost of Processing power and memory in modern computer along with full text of item in electronic form, from publishing stages.

Finally,

Until recent, Indexing was accomplished by creating a bibliographic citation in a structured file that refers to the Original text.

→ The Indexing process is typically performed by "Professional Indexers" associated with Library Organization.



(63)

from the Ancient papyrus, clay tablets, cones and brick fragments "inscribed" using an ancient writing system known as cuneiform.

## Cuneiform tablets

from the Library of Congress' collections.

→ The word "cuneiform" is derived from "Latin".  
cuneus for "wedge" and forma meaning "shape".

→ Cuneiform is a "Logo-symbolic" script that was used to write ~~several~~ several languages of the Ancient Middle East.

→ Cuneiform was developed by the Sumerians, who thrived during the 3rd-millennium B.C.

→ Sumerians influenced <sup>their</sup> culture and development beyond their original borders in Mesopotamia (Present-day southern Iraq) & earliest civilization.

→ Originally Cuneiform signs were "pictograms", later it also became symbolic.

The materials used in Cuneiform - "clay and seeds".

→ The "clay and seeds" with are <sup>readily</sup> available.

→ "Reeds" were used as "writing implements".

↑  
A tall slender leaves of plants

→ The tip of a reed style was impressed into a wet clay surface to draw strokes of the sign, thus it acquiring a "wedge-shape".

i.e. The Sumerians invented this writing system, which involves the use of wedge-shaped reed styles to make impression in clay.

⇒ The Library of Congress acquired its collection of cuneiform materials in 1929 from Kirkor (an art dealer).

## DIALOG :-

DIALOG Online Search system - (1966 became available).

→ DIALOG was the first interactive, Online search system addressing 'large databases' and ~~and~~ ~~allowing~~ allowing 'Iterative refinement of results'.

→ DIALOG was developed at "Lockheed Palo Alto Research Laboratory", and it extended through contracts with NASA.

Why DIALOG became Popular:

Its Speed, ease of use and wide range of data content attracted the professional users worldwide including

- \* Scientists
- \* attorneys
- \* educators &
- \* librarians.

→ DIALOG preceded major "Internet search tools" by more than two decades.



By 1972,

The commercial introduction ~~extended~~ extended to all the professions by allowing "the online access" to large collections of digitized materials by a way of "command language" that allowed the ~~researcher~~ searcher to interactively refine results.

DIALOG's ability to allow Interactive refinement of results:-

\* DIALOG's "major technical innovation is reflected in its name"

i.e. It enables a conversation between the "searcher" and the "computer".

\* "Search at its best is a conversation"

i.e. "An iterative, Interactive Process where we find we learn".

User Characteristics of DIALOG Language:

There are 5 Important Characteristics.

1. The search question is constructed at search time.
2. Dialog is designed for nonspecialists.  
i.e. It avoids communication barrier.
3. Command language is independent of the particular data it searches.

4. As "an Online-System".

i.e. It allows continual redefinition of the search - question based on the examination of intermediate results.

5. Control of process lies with the user.

i.e. Computer merely serves as a data-processing extension of user.

### DIALOG Commands:-

The DIALOG System provides the number of Commands with which the "Searcher" interacts with the "Computer".

A Search consists of

- functions
- i) Identifying
  - ii) Selecting terms and phrases that reflect the user's interest.
  - iii) Combining descriptors into search ~~searches~~ expressions
  - iv) Examining the retrieved citations and modifying search expressions.

Each of these functions accomplished with - a particular command.



The 4 Principal Commands are:

- i) EXPAND
- ii) SELECT
- iii) COMBINE
- iv) DISPLAY

### EXPAND :

EXPAND provides list of "synonyms", related terms and similar combination of words defining his/her user interest.

### SELECT :

SELECT is used to selecting the list of terms (ie term a, term b, term c --), ~~(or)~~ a

"a range of terms" i.e. (term a — term d) ~~(or)~~  
(or)

a list of ranges and terms.

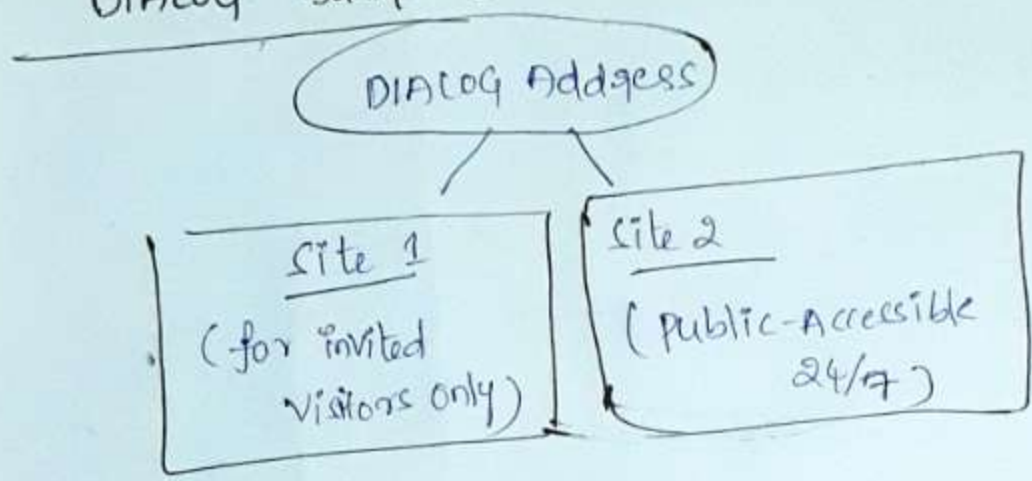
### COMBINE :

The List of Boolean Expressions are used to combine the documents.

### DISPLAY :

It allows the searcher to successively display the documents contained in the resultant set.

# Dialog Address:-



Site 1 :- (for invited - visitors only)  
Lockheed Martin Advanced Technology Center  
formerly Lockheed Palo Alto Research Laboratory  
Bldg. 201 (front lobby site)

Site 2 : (Publicly-Accessible 24/7): mountain view,  
"Computer History Museum", (on inside face of  
front patio brick wall.)

## Security / protection of site 1 & site 2:-

Site 1 : It requiring a specific US security clearance  
to access site. (with in a secure facility).

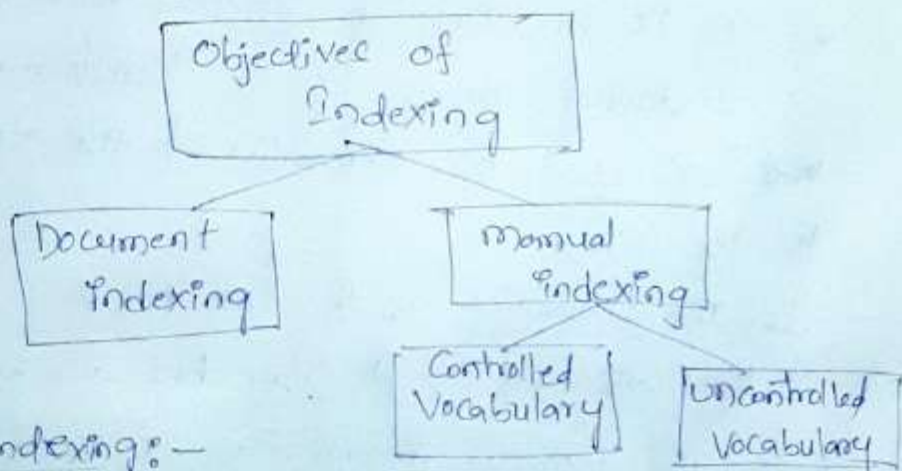
Site 2 : freely accessible to the public 24/7.



## Objectives :-

The Objectives of "Indexing" change with evolution of Information Retrieval system.

The Indexing Define the "Source" and "major concepts" of an item that provide the standard Mechanism for "Index terms".



## Document Indexing :-

The Document file provides the new class of Indexing called "Document Indexing".

## Manual Indexing :-

In Manual Indexing Environment,

### \* Controlled Vocabulary :-

The Use of controlled Vocabulary makes the Indexing process "slower" but potentially Simplifies the search process.

### \* Uncontrolled Vocabulary :-

The Use of uncontrolled Vocabulary makes the Indexing process "faster" but Search process much more difficult.

(70)

The availability of "Items" in "electronic form" changes the objectives of manual indexing.

i.e. The source information can be automatically extracted.

### (\*) Indexing Process

⇒ It is an important process in IR.

⇒ Indexing process is a transformation of an item that "extracts" the semantics of the topics discussed in the item.

⇒ The semantics of the item not only refers to the subject discussed in the item but also weighted systems.

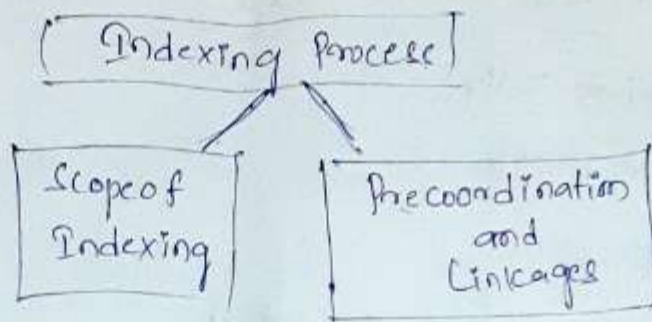
⇒ The extracted information is used to create "the processing tokens" and "the searchable data structures".

⇒ The index can be based on the

- full text of the item,
- automatic or manual generation of the item
- natural language representations of the item.

⇒ The result of this process stored in one of the data structures called "inverted data structure".





### Scope of Indexing :-

When scope of indexing performed manually, the following problems are occur with the interaction of two sources:

- i) The Author
- ii) The Indexer.

### The Author :-

The Vocabulary domain of the author may be different than that of the Indexer.

### The Indexer :-

→ The Indexer is not an expert on all the areas the item is being presented and the result is in different quality levels of indexing.

→ The Indexer must determine when to stop the indexing process.

## Factor effecting the scope of Indexing:-

(42)

i) Exhaustivity

ii) Specificity

### Exhaustivity:-

The extent to which  
extension

→ The different concepts in the item are indexed.

~~Exe~~

### Specificity:-

The preciseness (or) The exact index terms used  
- in "indexing" called Specificity.

## Linkages & precoordination:-

### Linkages:

The Linkages are used to correlate the  
related attributes associated with concepts discussed  
- in the item.

### Pre co-ordination:-

The process of creating the term "Linkages" at  
index creation time is called precoordination.

### Post Co-ordination:-

When the index terms are not coordinated at  
index time, the co-ordination occurs at search time  
is called Post co-ordination.

The post Co-ordination is implemented by "AND" <sup>indexing</sup> <sub>terms</sub>.



### Example :-

The different types of Linkages as shown below.

(43)

i.e. "The <sup>1</sup>drilling of <sup>2</sup>oil wells in <sup>3</sup>Mexico by <sup>4</sup>CITGO and  
The <sup>1</sup>introduction of <sup>2</sup>oil refineries in <sup>3</sup>peru by the <sup>4</sup>US."

### INDEX TERMS

### Methodology

→ oil, wells, Mexico, CITGO, refineries, Peru, BP, drilling. } No linking of terms.

→ (Oil wells, Mexico, drilling, CITGO)  
\* (U.S., oil refineries, peru, introduction) } linked (pre-coordination).

→ (CITGO, drill, Oilwells, Mexico)  
\* (U.S., introduction, Oil refineries, peru) } linked (precoordination) with position indicating role.

→ \* (SUBJECT: CITGO;  
ACTION: drilling; OBJECT: Oil wells;  
MODIFIER: in Mexico)  
\* (SUBJECT: U.S.; ACTION: introduces; OBJECT: Oil refineries;  
MODIFIER: in peru). } linked (pre-coordination) with modifiers & indicating role.

Fig: Linkage of Index Terms

## (\*) Automatic Indexing :-

Automatic indexing is "automatically" determine the "index terms" assigned to an item.

- Automatic indexing is deals with "Simplest case" and "Complex case":

### \* Simplest case :-

The simplest case deals with the total document of - Indexing.

### \* Complex case :-

The complex case deals with the "Human Indexer".  
The Human Indexer determines the "limited no. of "index terms" for major concepts in the item.

## Result of Automatic Indexing :-

The automatic indexing results fall into two classes:

- i) Weighted indexing system
- ii) unweighted indexing system

### i) Weighted Indexing :-

→ Weighted Indexing represents the "term values" placed inside the "index terms".

⇒ The Basic function of the weighted indexing associated with "frequency of occurrence" of an item.

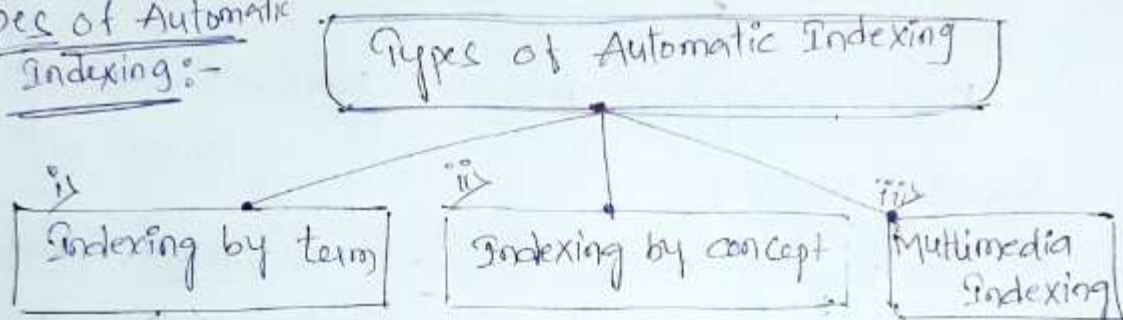


## ii) Unweighted Indexing :-

⇒ The unweighted indexing represents the existence of an ~~item~~ index term in a document. (45)

⇒ In unweighted indexing some times, "word locations" are kept as a part of searchable data structure.

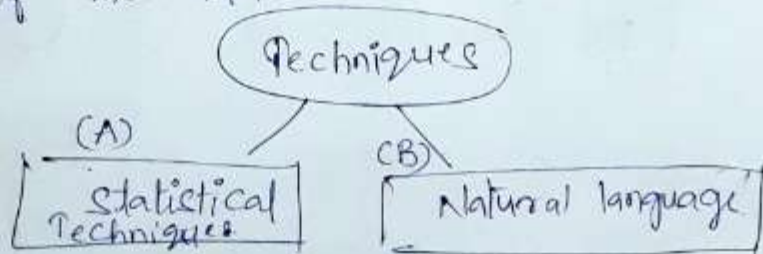
## Types of Automatic Indexing :-



## i] Indexing by term :-

⇒ The "terms" of original item are used as basis of "index process".

There are 2 major techniques used for creation of the index are:



## (A) Statistical Technique :-

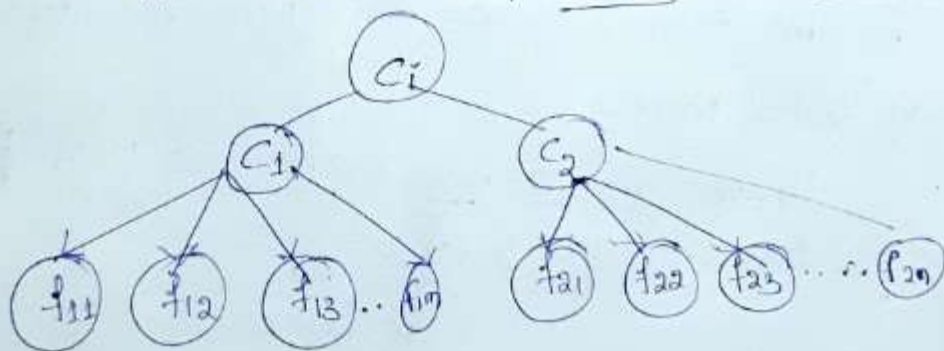
The statistical techniques can be based on the "vector models" and "probabilistic models" with a special case being "Bayesian models".

- All the models (i.e. Vector model, Probabilistic model and Bayesian model) classified as "statistical" (16) because their calculation of weights use the statistical information such as "frequency of occurrence of words" and their distributions in the searchable database.

Bayesian Model :-

A Bayesian model is a "directed acyclic graph" in which the node indicate the random variable.

It is the combination of "nodes" and "ArCs".



where  $C \rightarrow$  concept,  $f \rightarrow$  features.

To calculate the probability of " $C_i$ " from given " $f_{ij}$ ".

$C_i$  from given  $f_{ij}$

To perform the calculation 2 sets of probabilities are needed:

1. The Prior Probability  $\cdot P(C_i)$ .

means the item is relevant to concept  $C_i$ .



## 2. The "Conditional probability"

The conditional probability is indicated by  $P(f_{ij}/c_i)$ .

where,  $P \rightarrow$  Probability

$f_{ij} \rightarrow$  features

$c_i \rightarrow$  Concept

(11)

$\rightarrow$  The features  $f_{ij}$  where  $j=1, m$  are present in an item means the item contains topic  $c_i$ .

The automatic indexing task is used to calculate the "Posterior probability":

$$\text{i.e. } P(c_i / f_{i1}, \dots, f_{im}) \quad \text{where } j=1, m$$

The probability of the item that contains concept  $(c_i)$  given the presence of features  $(f_{ij})$ .

The Bayes inference / Bayesian formula used for Automatic indexing is:

$$P(c_i / f_{i1}, \dots, f_{im}) = \frac{P(c_i) P(f_{i1}, \dots, f_{im} / c_i)}{P(f_{i1}, \dots, f_{im})}$$

(or)

$$P(c_i / f_{i1}, \dots, f_{im}) = P(c_i) P(f_{i1}, \dots, f_{im} / c_i) / P(f_{i1}, \dots, f_{im})$$

⇒ The Result of search by 'posteriors', ~~using~~  
 the Bayes rule can be simplified to linear decision rules.

$$P(C_i | F_{f1}, \dots, F_{fm}) = \sum_k I(F_{fk}) W(F_{fk}, C_i)$$

(18)

where,

'g, w' are Interpreting the coefficients

i.e.  $w \rightarrow$  weights corresponding to each feature / concept  
 $f_k$ : Index term. Pair.

$g \rightarrow$  'g' as the sum of the weights of the features.

$I(F_{fk}) \rightarrow$  an indicator Variable. that equals 1  
 if  $f_k$  is present in the item (otherwise it is equals zero).

$P(C_i | F_{f1}, \dots, F_{fm})$   $\rightarrow$  Posterior probability.

(B). Natural Language processing :-

$\rightarrow$  The another approach to defining 'indexes to items'  
 is called Natural Language processing.

$\rightarrow$  The DR-LINK (Document Retrieval through Linguistic  
Knowledge) System processes items at the morphological,  
 lexical, semantic, syntactic and discourse levels.

Each level uses information from the previous level to  
 - Perform its additional analysis.



## ii) Indexing by concept:-

(19)

→ The basis for 'Concept Indexing' is that, there are "many ways to express" the same idea.

→ Indexing by term treat each of these occurrences as a "different index" and then use the other query expansion technique to expand a query to find the different ways the same thing has been represented.

### Concept Indexing:

Concept indexing determines a "canonical set of concepts" based upon a test set of "terms".

The concept indexing is also called as "Latent Semantic Indexing" because it indexing the latent semantic ~~indexing~~ information in items.

In concept indexing, it does not associated with each 'concept' (i.e. words / <sup>Set of</sup> words that can be used to describe it) but it is a mathematical representation (Ex. a vector / Context vector).

### Example of Concept Indexing:-

The example of Concept Indexing is the "Matchplus system" developed by HNC Inc.

→ The word stems, items & queries are represented by high dimensional vectors called "Context vectors".  
in (high dimension at least 300 dimensions).  
- Store.

The Matchplus system uses "neural networks" to facilitate.

facilitate machine learning concept of relationships" and  
Sensitivity to similarity of use. (80)

The two neural networks are used in context indexing, age

1. One Neural network : (Stem context vectors)

One Neural network learning algorithm generates "Stem Context Vectors" that are sensitive to similarity of use.

2. Second Neural network :

It performs query modification based upon user feedback.

Context vector / Vectors

Context Vector :

Words stems, items and queries are represented by "high dimensional" Vectors called "Context Vectors".

To define context vectors, a set of 'n' features are selected on an adhoc basis.

i.e. for any word stem k, its context vector  $V^k$  is  
is an n-dimensional Vector with each component 'j'  
interpreted as follows:

$V^k$  "Positive" if k is strongly associated with feature j.

$V^k \approx 0$  if word 'k' is not associated with feature j.

$V^k$  "negative" if word 'k' contradicts feature j.



Once the context Vectors for items are determined, they are used to create the index for an item. (21)

### iii) Multimedia Indexing :-

→ Indexing associated with multimedia differs from the previous discussions of indexing.

→ The automated indexing takes place in 'multiple passes' of the information 'versus' direct conversation to the indexing structure.

first pass : The first pass in most cases is a 'Conversation' from the analog input mode into a 'digital structure'.

The digital structure algorithms are 'extract' the different modalities that will be used to represent the item.

#### Next pass :-

Indexing video (or) images can be accomplished at the 'raw data level'.

-ex:- The aggregation of raw pixels.

i.e. The raw pixels, (distinguish) contains the primitive attributes such as 'color' and 'luminance' and at semantic level meaning full objects are recognized.

-ex: an airplane in image/video frame.

## Video processing:-

→ The system will periodically collect a frame of video input for processing. (2)

→ It might compare that frame to the last frame captured to determine differences between the frames.

→ If the difference is below the threshold it will discard the frame.

Ex: Virage.

## Image processing:-

In Image processing,

→ Semantic level indexing requires "pattern recognition of Objects" with in the images.

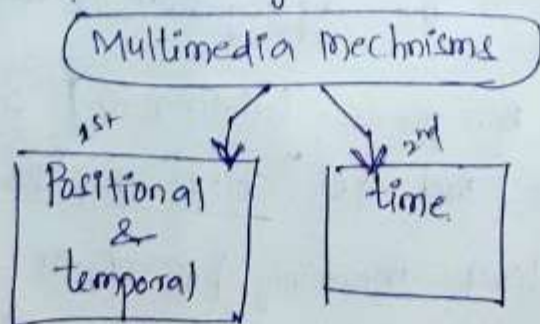
Ex: MIT's PhotoBook,

IBM's QBIC &

Multimedia Database from Informix/Virage.

Finally: → (multimedia mechanisms)  
to store,

The Multimedia items need some mechanism to correlate the different modalities during the search.





first mechanism:

→ Positional & temporal mechanism :-

This mechanism is used when the modalities are interepressed / interepressed in a linear sequential composition. (25)

Example :-

A document that has "images" (or) "audio" inserted can be considered a linear structure.

→ Second mechanism: i.e. "time" :-

The second mechanism is based upon time because the modalities are executing concurrently.

i.e. - Ex. - Television, SMIL

The typical Video Source of television is inherently a multimedia source.

It contains video, audio and potentially closed captioning.

SMIL → synchronized multimedia Integration Language.

It is a mark-up language designed to support multimedia presentations that integrate text with audio, images and video.

## (\*) Information Extraction :-

Extraction means → removal,  
- taking out, drawing out, pulling out etc.

→ Information Extraction is a task of extracting the (84) <sup>nt</sup> structured information from unstructured and/or semi-structured machine-readable documents.

→ There are two processes associated with information Extraction:

i) Determination of facts

ii) Extraction of text

### i) Determination of facts :-

Determination of facts go into "structured fields" in a database.

In this case, only a subset of the important facts in an item may be identified and extracted.

### ii) Extraction of text :-

Extraction of text can be used to "summarize an item".

In summarization, all the major concepts in the item should be represented in the summary.

→ The process of extracting facts to go into indexes is called "Automatic file Build".

It's "goal" is to process "Incoming Items" and extract index terms" that will go into a structured database.



→ An Information Retrieval System's "goal" is to provide an in-depth representation of the contents of an item. (85)

→ An Information Extraction System: Only analyzes those portions of a document that potentially contain information relevant to the extraction criteria.

### Metrics to compare Information Extraction:-

The ~~prev~~ previously defined measures of precision and recall are applied with slight modifications to their meaning.

Precision:- Precision refers to how much information was extracted accurately Versus the total information extracted.

Recall:- Recall refers to how much information was extracted from an item Versus how much should have been extracted from the item.

<sup>It shows,</sup>  
i.e. The amount of correct and relevant data extracted Versus the correct and relevant data in the item.

### Additional Metrics:-

↓  
Overgeneration    ↓  
                                fallout

The Additional Metrics used for Information Extraction for

- are: "Overgeneration" & "fallout".

## Overgeneration :-

Overgeneration measures the amount of irrelevant/non-relevant information that is extracted.

(86)

## Fallout :-

Fallout measures how much a system assigns incorrect slot fillers, as the number of potential incorrect slot fillers increases.

These measures are applicable to both "human" and "automated" extraction process.

⇒ The best source of analysis of data extraction is from the "Message Understanding conference proceedings".

⇒ The conference were held in 1991, 1992, 1993 & 1995.

⇒ The conference are sponsored by the Advanced Research Project Agency :

## Objective of data extraction :- "slot"

Objective of data extraction is update structured database with additional facts.

⇒ The term "slot" is used to define a particular category of information to be extracted.

⇒ Slots are organized into templates / semantic frames.

⇒ Information extraction requires multiple levels of analysis of text of the item.

ie It must understand the words and their context.   
 ⇒ This focusing is very similar to Natural Language processing.



# Data Structure

(89)

## ① Introduction to Data Structure :-

A Data structure is a "storage" that is used to store and organize data.

It is a way of arranging data on a computer so that it can be accessed and updated efficiently.

There are usually "two major data structures" in any information system.

### first/One major Datastructure :-

→ One major Data structure "stores and manages" the received items in their "Normalized form".

→ The process of supporting this structure is called the "document Manager".

### Other Major Datastructure :-

Other major Data structure contains the "processing tokens", and "associated data" to support search.

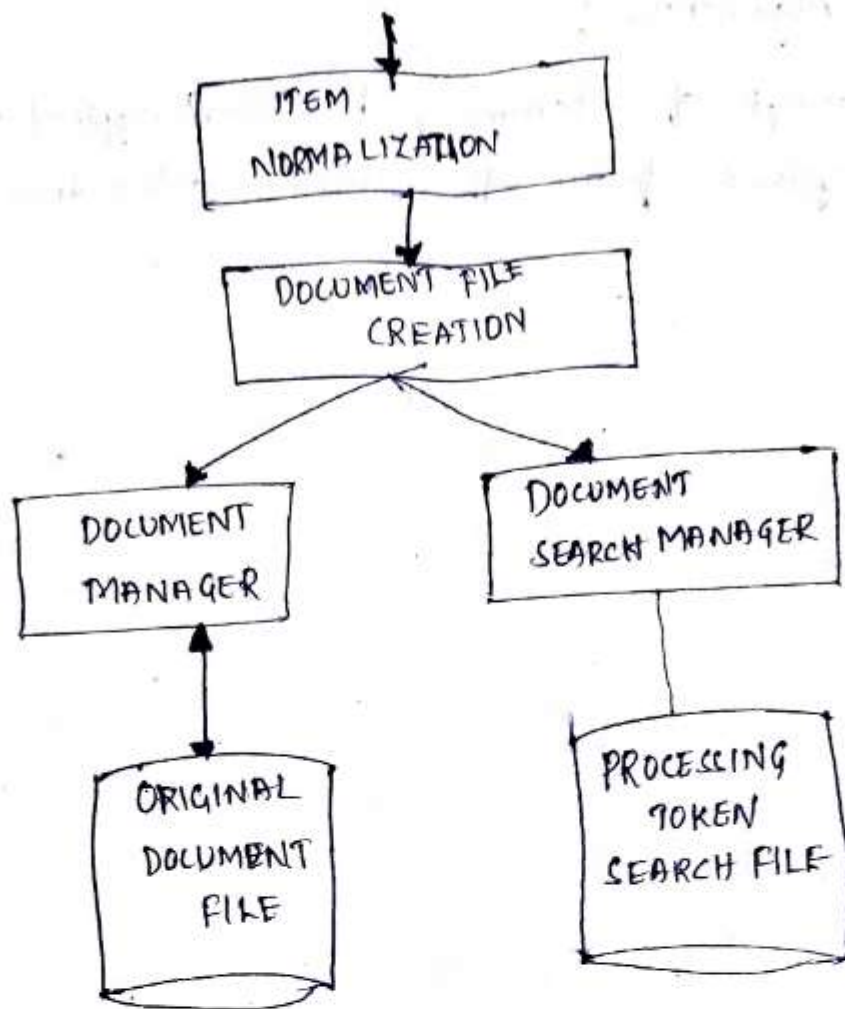


Fig. Major Datastructures

→ The above figure expands the document file creation function (we discussed in 1st chapter).

→ The Results of a search are references to the items that satisfies the search statement, which are passed to the document manager for retrieval.

→ To understand this datastructures background, the reader / user should pursue a text on finite automata and language (regular expressions).

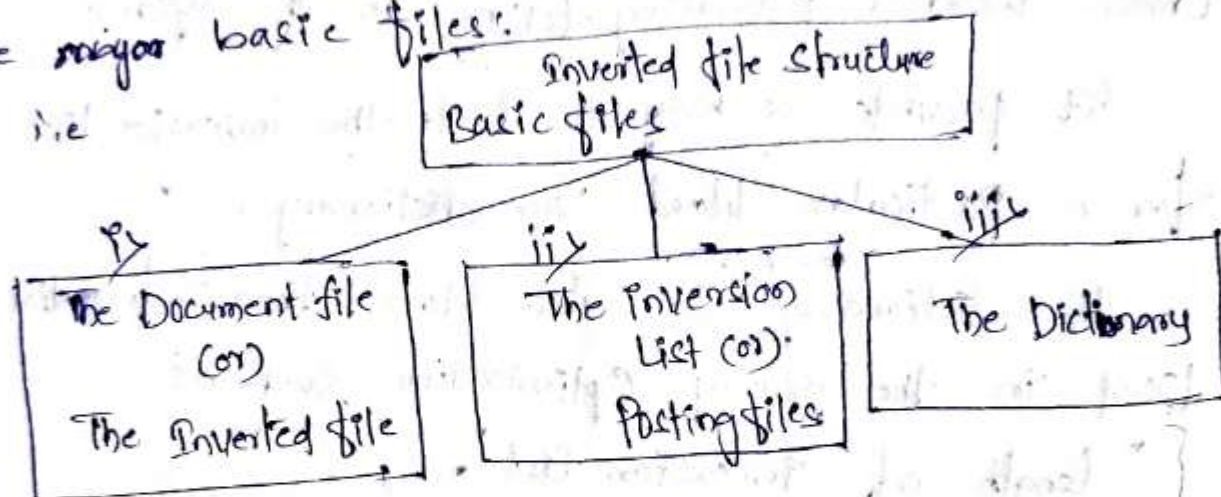


## ⊗ Inverted file Structure:-

The most common data structure used in both Database management and Information Retrieval Systems is the Inverted file structure.

It is used to minimize the secondary storage access when multiple search terms are applied across the total database.

The Inverted file structures are composed of three major basic files:  
i.e.



## ⊗ Inverted file / Document files:-

The name ~~inversion~~ "Inverted file" comes from its underlying methodology of storing an inversion of the documents.

from the inversion of document perspective, for each word a list of documents in which the word is found.

ii) The Inversion list :-

(90)

Inversion list sometimes called as "posting list".

In this each document in a system give a

"Unique numerical Identifier".

The Identifier is stored in the Inversion list.

iii) The Dictionary :-

The Dictionary is typically a stored list of all Unique words (processing tokens) in the system.

It provide a way to locate the Inversion list for a particular word via dictionary.

The dictionary can also store other information used in the query optimization such as "length of inversion list".

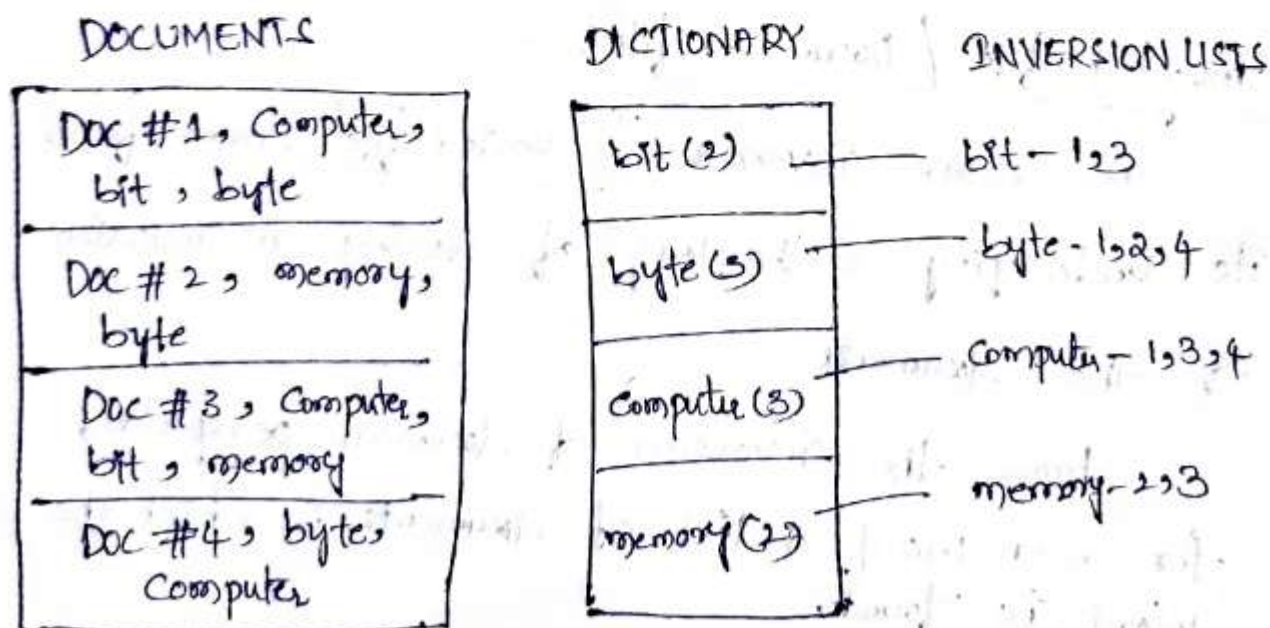


Fig : Inverted file structure



from the figure diagram :-

(91)

- The Documents contains information of words (such as bit, computer, bytes) i.e. how many words are there in a particular documents.
- The Dictionary contains how many number of times the particular word is appear in the document, it written inside parenthesis ( ).
- The Inversion List contains in which document the particular word is appearing.

finally,

The Inverted file structure uses some silent features i.e.

1. It increases precision and Recall.
2. It also uses browse capabilities such as Ranking and Zoning.
3. It uses the NLP (Natural Language processing)
4. It is mainly used to store the concepts and Relationships.
5. Rather than using a dictionary to point the Inversion list B-Trees can be used.

In B-Tree Inversion list may be appear at  
"leaf level / leaf node".

(92)

\* B.T. A B-tree of Order "m" is defined as:

- \* A "root node" with between "2 and 2m" keys.
- \* All other Internal nodes have between "m" and "2m" keys.
- \* All keys are kept in Order from smaller to larger.
- \* All leaves are at the same level (or) differ by at most one level.

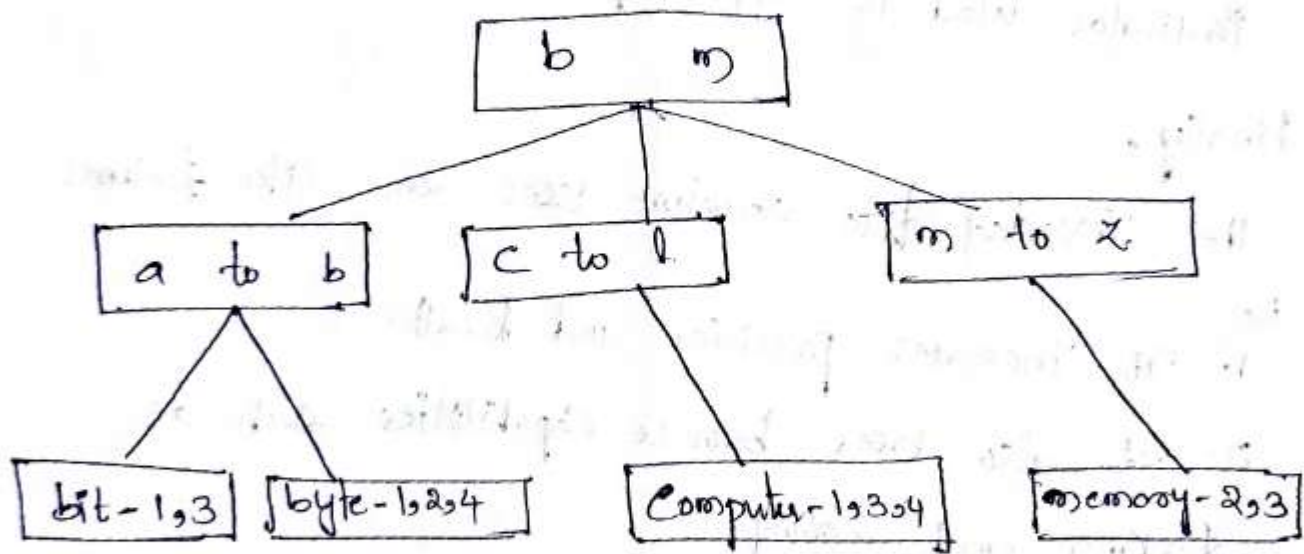


Fig: B-Tree Inversion list.



## (\*) N-Gram Data Structures :-

(93)

→ A Variant of the 'Searchable data structure' is N-Gram Data Structure, that breaks processing tokens into smaller string units and uses the token fragments for search.

→ N-grams can be viewed as a special technique for conflation (stemming). It is used as a unique data structure in Information systems.

→ N-grams are 'fixed-length' consecutive series of "n" characters.

Unlike stemming,

ie The stemming determines the stem of a word that represents the semantic meaning of the word <sup>working → work</sup> but n-grams do not care about semantics.

Instead of that, they are algorithmically based upon a fixed number of characters.

→ The Searchable data structure is transformed into Overlapping n-grams, which are then used to create the Searchable data base.

Examples : Bigramme, Trigramme and Pentagramme for the word phrase " sea colony".

(94)

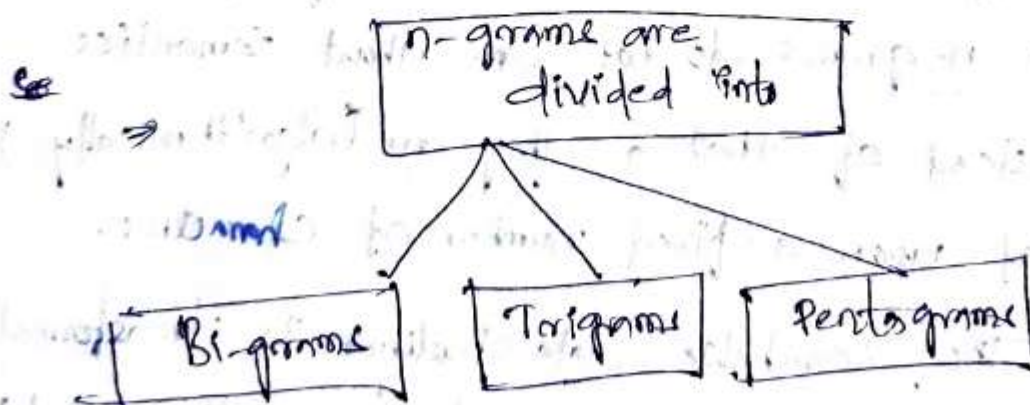
⇒ for n-gramme : with  $n > 2$  : Some systems allow interword symbols to be part of the n-gram.

The symbol "#" is used to represent the interword symbol.

i.e. # ⇒ interword symbol

⇒ The # may be represented any set of symbols such as blank, period, semicolon, colon etc.

⇒ It is possible that same n-gram can be created multiple times from a single word.





For Example :-

(95)

Bigrams, trigrams and pentagrams for the word phrase "Sea colony".

\* Bigrams for word "Sea colony"

Word                      Bigram  
Sea → se ea

colony → co ol lo on ny

(Bigrams  
no interword symbols)

∴ [se ea co ol lo on ny]

\* Trigrams for word "Sea colony"

Word

Trigram

Sea →

Sea

colony →

col olo lon ony

(Trigrams  
no interword symbols)

∴ [Sea col olo lon ony]

\* Pentagrams for word "Sea colony"

Word

Pentagram

Sea →

Sea (no change)

Colony →

colo olon lony

(Pentagram  
no interword  
symbols)

∴ [Sea colo olon lony]

\* Trigramme with Interword Symbol '#' for word 'Sea colony' 96

Place the Trigramme inside bigrams and place # interword symbol at both ends.

i.e. # Se sea ea #

bigram # Se sea ea # interword symbol

interword symbol co cal old don ony ny

# Co cal old don ony ny #

∴ [ # Se sea ea # # Co cal old don ony ny # ]

\* Pentagramme with interword symbol '#' for word 'Sea colony'

Place the Pentagrams inside trigrams and place # interword symbol at both ends.

Sea

i.e. # Sea #

colony ⇒ colo olon lony → 4-grams

colony ⇒ colon olony → 5-grams

# colo colon olony lony #

4-grams 5-grams 5-grams 4-grams

∴ [ # Sea # # colo colon olony lony # ]

4-grams 5-grams 5-grams 4-grams



## Example 2

Reveal Details find Bigrams, trigrams &

Pentagrams.

(97)

### ② Uses of N-Grams :-

① The first use of n-grams are at the time of "World War-II", it was used by "cryptographers".

The cryptographers namely "Fletcher Pratt" states / defines the "bigram and trigram tables" so that any cryptographer can decipher Simple Substitution Cipher.  
takeably partition/divide up

Another cryptographer namely "Adams" describes the "Use of Bigrams" as a method for "conflating terms" (or) stems.

i.e. The conflating terms do not follow normal definition of stemming because the "stemming" has "semantic meaning" whereas "n-grams" have "no semantic meaning".

② Another Major Use of N-gram (particularly trigrams) is in "Spelling error detection and correction".

i.e.

ie There are 4 categories of errors.

### Error Category

- \* Single character Insertion
- \* Single character Deletion
- \* Single character Substitution
- \* Transposition of two adjacent characters

Example → Computer

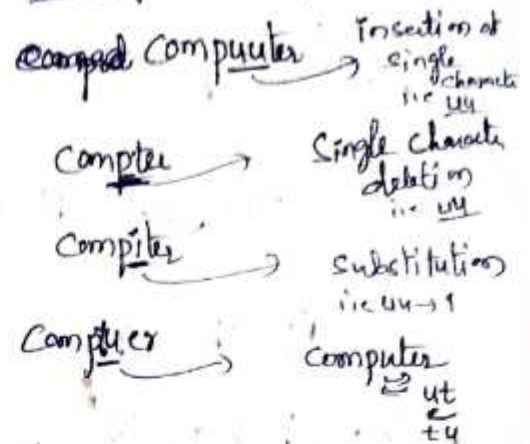


Fig: Categories of Spelling Errors.

③ N-Gram patterns also used for "Identifying the language of the item".

④ Trigrams have been used for "text compression to manipulate the length of index terms".

⑤ Advantages & Disadvantages of N-grams:-

Advantages:-

- i) longer n-grams, ignoring the word boundaries.
- ii) <sup>There is a</sup> A finite limit on the no. of searchable tokens.

$$\boxed{\text{Max seg}_n = (\lambda)^n}$$

Where,  $\text{Max seg}_n \Rightarrow$  Maximum no. of Unique "N-grams" that can be generated.



ie Max Seq  $\rightarrow$  calculated as a function of 'n'. (99)  
n  $\rightarrow$  Indicates the length of N-gram.  
 $\lambda$   $\rightarrow$  Indicates the no. of processible symbols from the alphabet.

### Disadvantages :-

- i. The disadvantage of N-gram is "Increased size of Inversion list" that store the linkage datastructure.
- ii. N-grams are poor representation in "Concepts and their relationships".

### (\*) PAT Data structure

The name PAT is the short form for <sup>PATRICIA</sup> Patricia trees.

PATRICIA stands for

P  $\rightarrow$  Practical

A  $\rightarrow$  Algorithm

T  $\rightarrow$  To

R  $\rightarrow$  Retrieve

I  $\rightarrow$  Information

C  $\rightarrow$  Coded

E  $\rightarrow$  In

A  $\rightarrow$  Alphanumeric.

ie Practical Algorithm To Retrieve Information Coded In  
Alphanumerics.

⇒ The PAT Data Structures were described by

(100)

Knuth - 92

Gooden - 83

Knuth - 73 &

Knuth - 68 has "Patricia trees"

⇒ Concept:

The Original Concepts of PAT tree data structure were described as "Patricia trees", and they are used for searching text and images.

These applications are in "Genetic database".

The PAT datastructure represent the information in 2 ways.

i.e. \* PAT trees &

\* PAT Arrays.

PAT trees and PAT arrays are addressing the different view of continuous "text" input.

The Input stream transformed into a searchable data structure consisting of "substrings".

Creation of PAT trees:

In creation of PAT trees each position in the input string is the "Anchor point" for a sub-string.

The substring start at "Anchor point" and include all new text upto end of the input.



→ All substrings are 'Unique'.

(10)

Sub String :-

\* A Substring can start at 'any point in text' and uniquely indexed by its 'starting location' and length.

\* If all the strings are to the end of the input, Only the starting location is needed.

Because the length is different from the location and the total length of the item.

\* It is possible to have a substring go beyond the length of the input stream by adding the additional null characters.

These substrings are called sstring.  
sstring means (Semi-Infinite String).

Example : Some possible sstrings for an 'input text' given below.

TEXT :

sstring 1

sstring 2

sstring 5

sstring 10

sstring 20

sstring 30

1 2 3 4 5 6 7 8 9 10 11 12 13 14 15 16 17 18 19 20 21 22 23 24 25 26 27 28 29 30 31  
Economics for Warsaw is complex.

1 2 3 4 5 6 7 8 9 10 11 12 13 14 15 16 17 18 19 20 21 22 23 24 25 26 27 28 29 30 31  
Economics for Warsaw is complex.

economics for Warsaw is complex.

omics for Warsaw is complex.

for Warsaw is complex. (without space count)

Warsaw is complex. (with space count)

ex. (with space count)

Fig: Examples of sstring.

In PAT trees,

102

\* The key values are stored at the "leaf nodes" in the PAT tree.

for a text input of size "n" there are "n" leaf nodes and "n-1" at most higher level nodes.

Example of the strings used in generating a PAT.

Input: 100110001101 (∴ Binary representation of HOME)

String 1 1001....  
String 2 001100...  
String 3 01100...  
String 4 11...  
String 5 1000...  
String 6 000...  
String 7 001101  
String 8 01101

H = 100

O = 110

M = 001

E = 101

i.e. 100110001101  
H O M E

Fig: strings for input "100110001101"

→ the word home produces the input

100110001101



⇒ No. of string = 8 (nodes)

i.e  $n=8$

⇒ We get  $n-1$  = levels

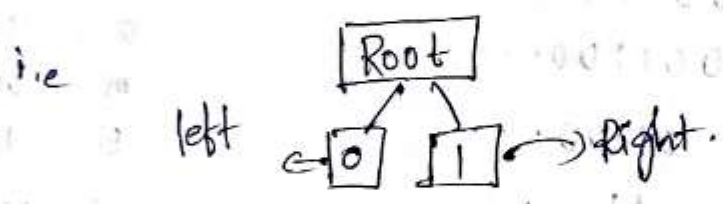
$8-1=7$

i.e for 'n' no. of node we get the 'n-1' levels

Rules for PAT Tree :-

The individual bits of string decides the branching

- The pattern with Zero (0-bit) branching 'left'
- The pattern with one (1-bit) branching 'Right'



Note

Take any query / string to reach external node.

query = 001100  
5 bits  
1st bit

- We first inspect bit-1 (if it is '0', branching will be done at left)
- then bit-2 (if it is '0' branching will be done at left)
- bit-3 (if it is '1' branching will be done at Right)
- bit-4 (if it is '1' branching will be done at right)
- bit-5 (if it is '0' branching will be done at left)

Once we reach desired node we have to make one final comparison with one of string.

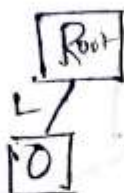
i.e. Example: INPUT : 100110001101  
 nodes:  $n=8$  ,  $n-1=7$  (levels)

(104)

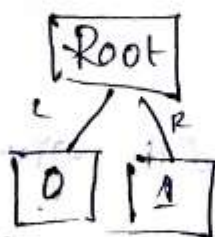
i) Initial Tree contains 'Root - R..



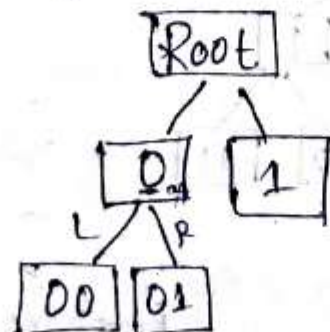
ii) add "0-bit" as left branching &



iii) add "1-bit" as Right branching:

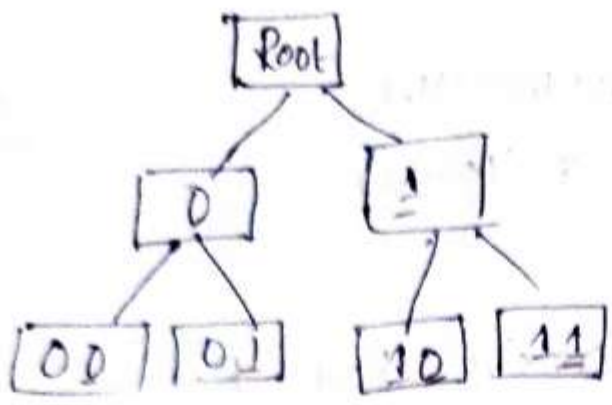


iv) for desired position '0' (Zero) compare with two bits i.e. again Zero & one.

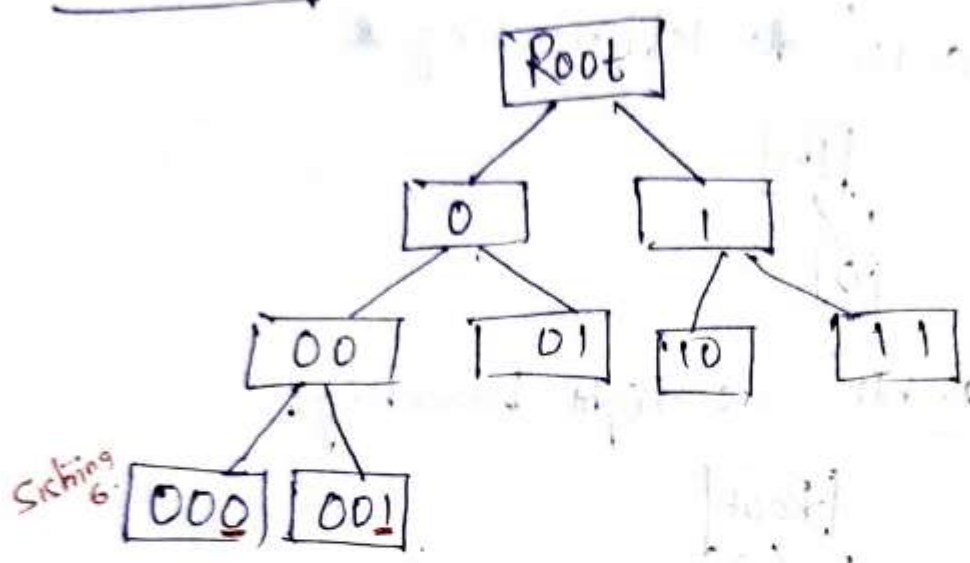


v) Similarly for 1 (Right most bit.) Compare with two bits i.e. again Zero & one.

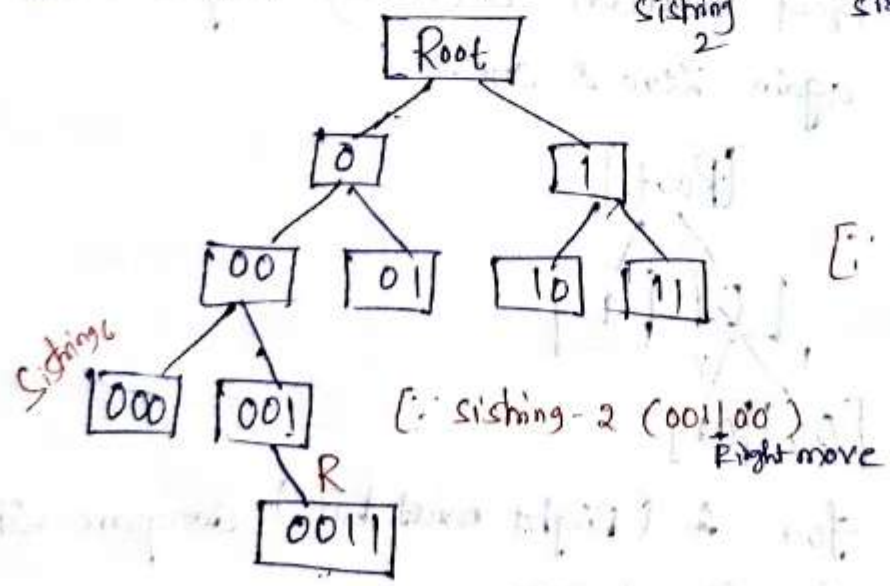




from <sup>Given</sup> strings we are going to construct PAT tree



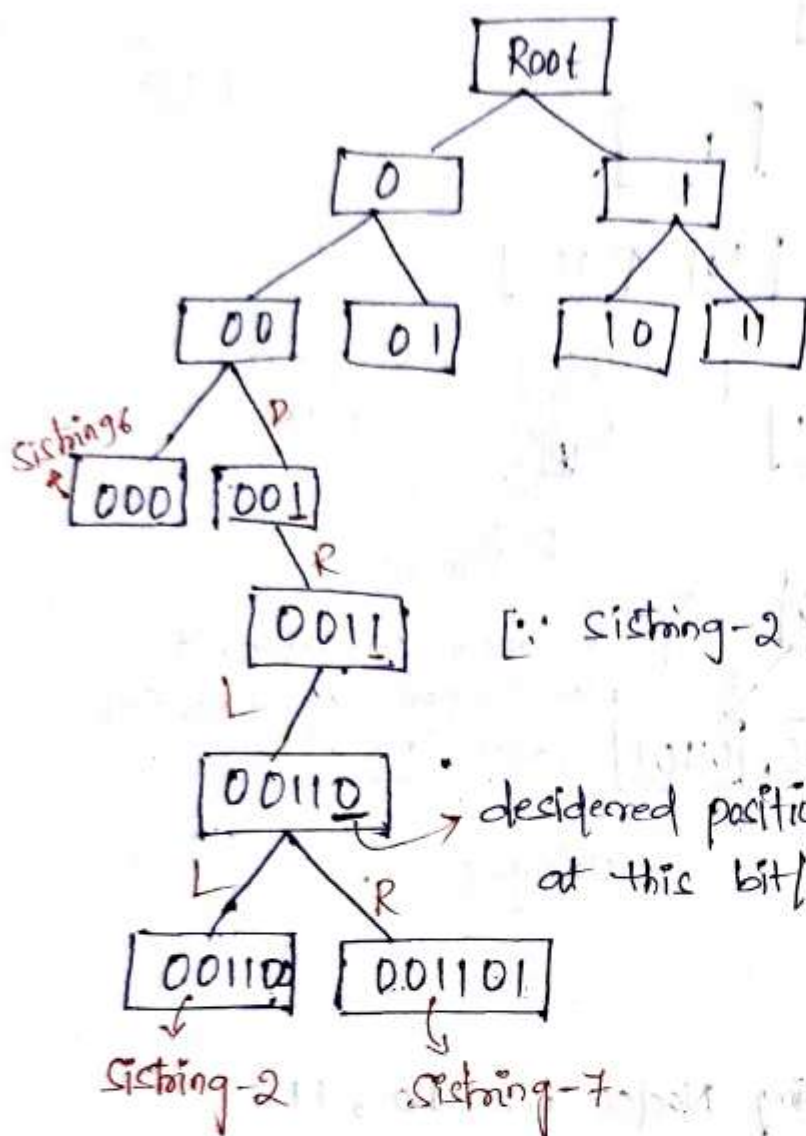
→ Now compare 001 node with given strings.  
In given input strings: 001100, 001101  
ie String 2 String 4



[∴ String-2  
001100  
move Right

[∴ string-2 (001100)  
Right move

Next compare remaining bit in string 2. ie 001100  
Left (zero) move



∴ Sistring-2 i.e. 001100

↓ Left move

at this bit/place compare both bits 0 & 1

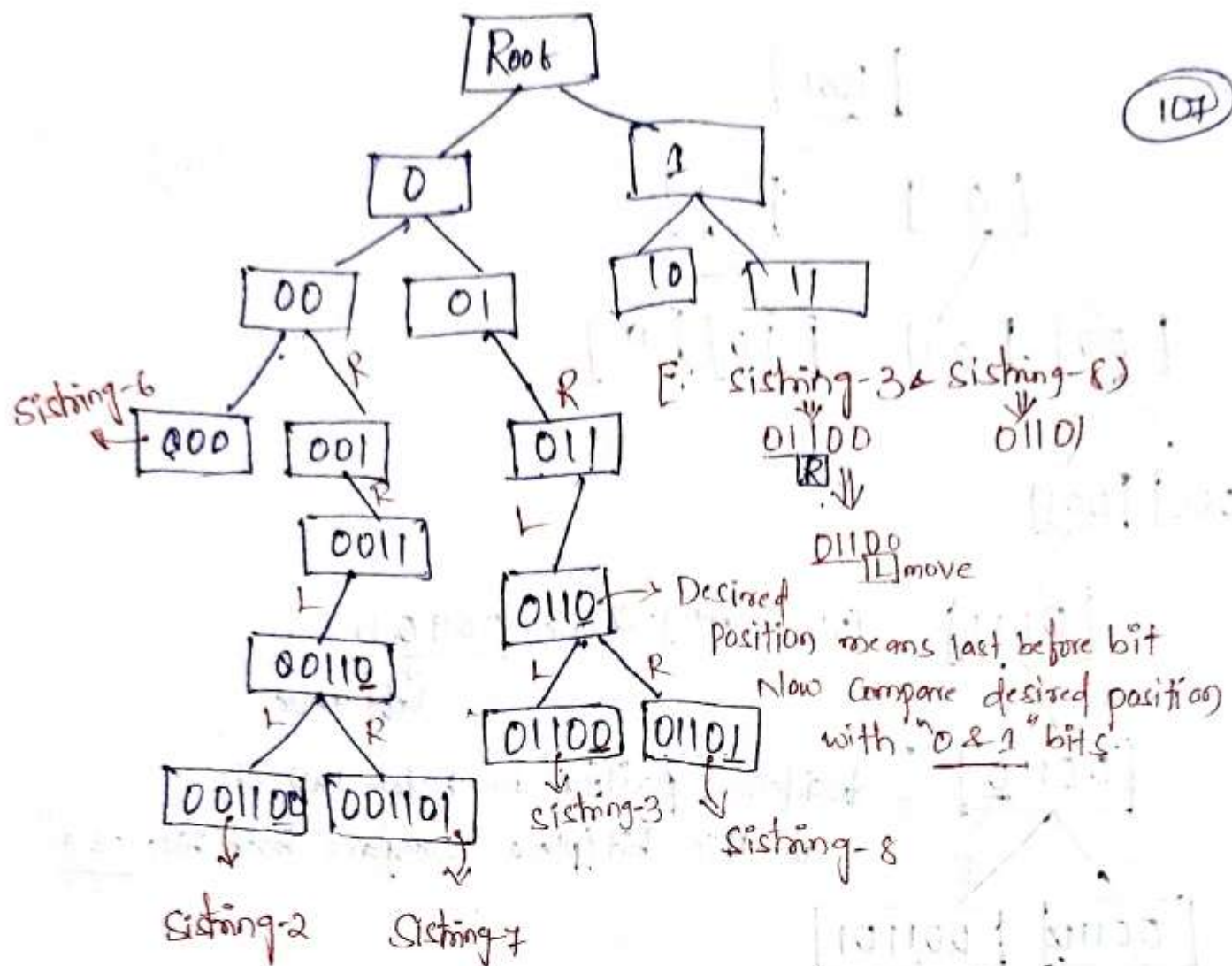
⇒ Similarly Compare Remaining Nodes. ~~Now comp.~~  
i.e. (01, 10, 11)

Now Compare 01 node (i.e. 01)

→ check is there any sistring start with 01.  
at sistring-3 (i.e. 01100) & sistring-8 (i.e. 01101)  
~~start~~ start with 01.

Now construct <sup>remaining</sup> PAT Tree.





⇒ Similarly Compare Remaining Nodes i.e 10, 11.

Now compare "10" Node.

- \* Check if there any sistring start with 10.
- \* At Sistring-1 (1001) and Sistring-5 (1000) are start with "10".

Now construct Remaining PAT tree.

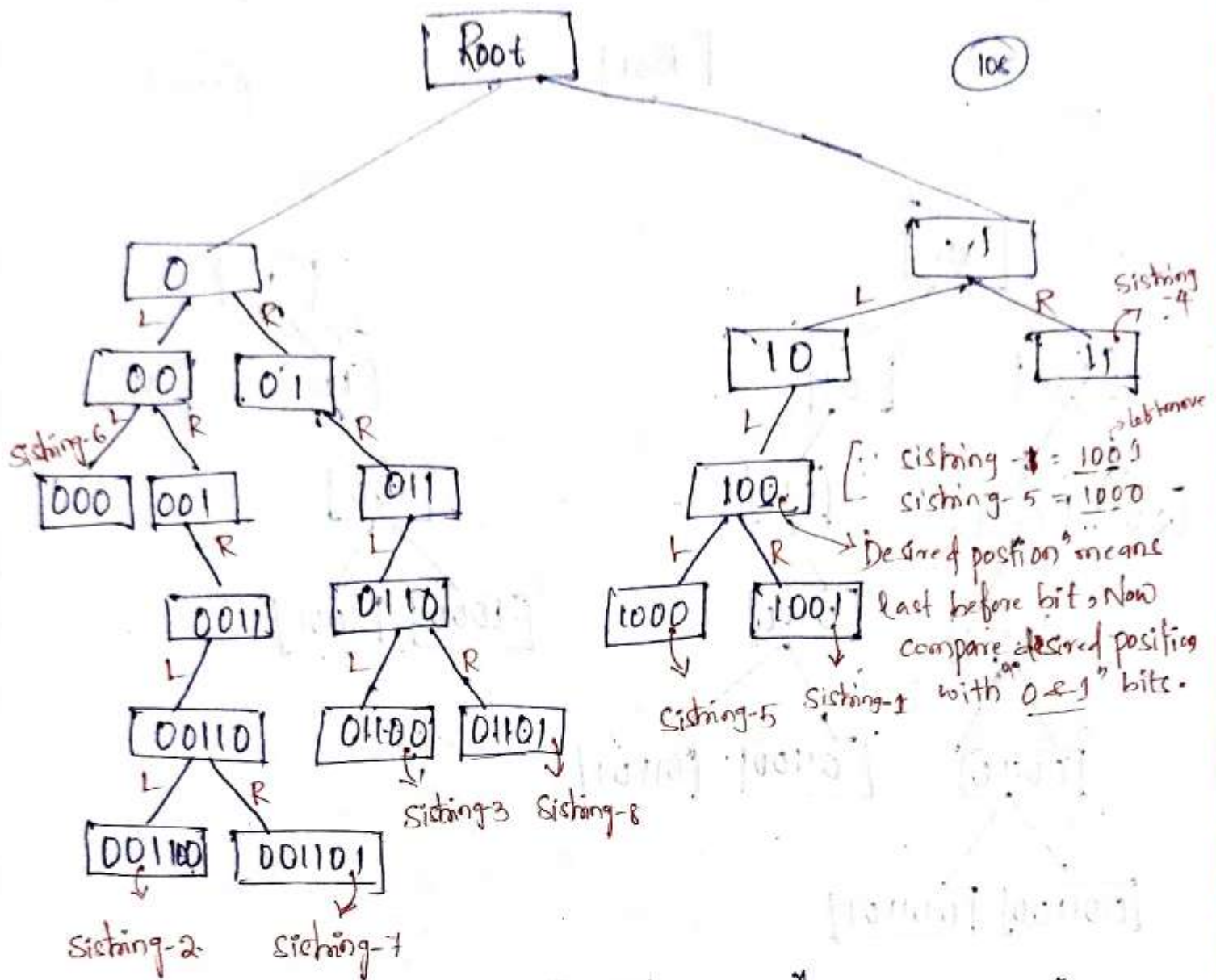
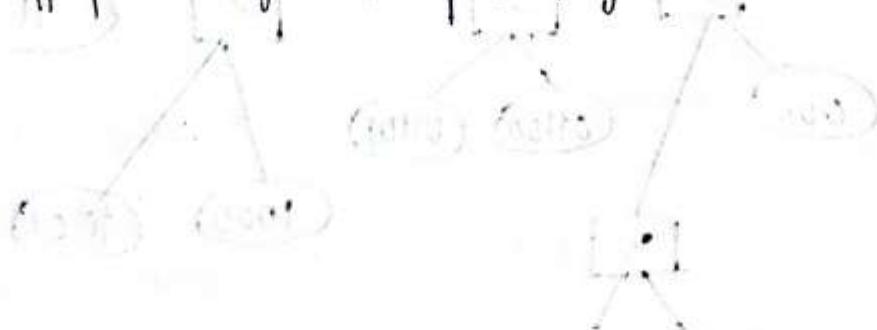


Fig: PAT Binary Tree for input "100110001101"

⇒ for the above PAT tree (skip-bits) some bits are skipped through its path length.





109

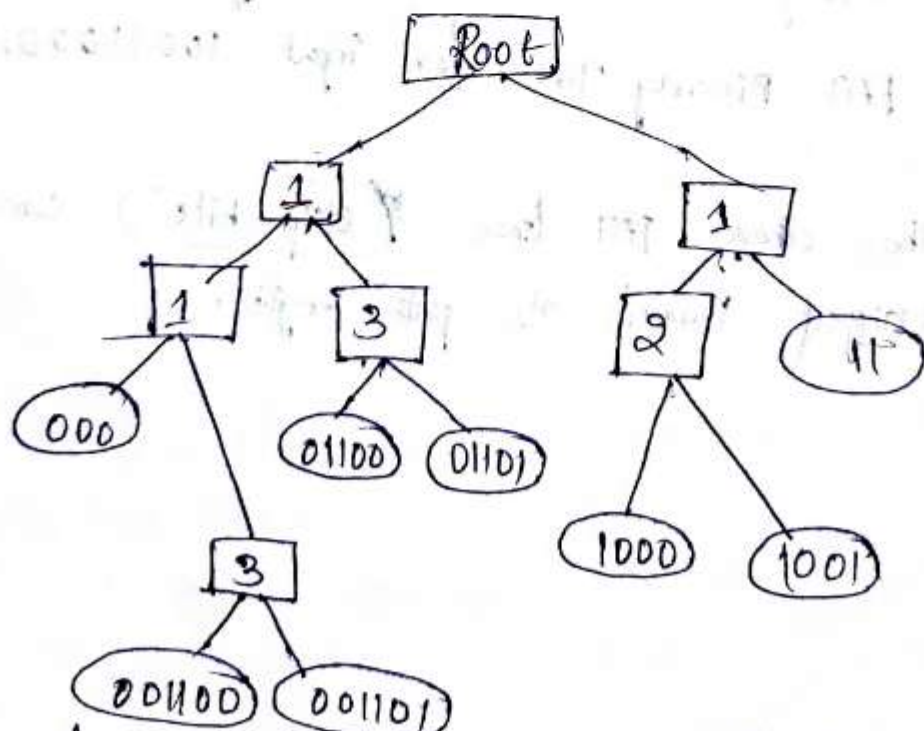
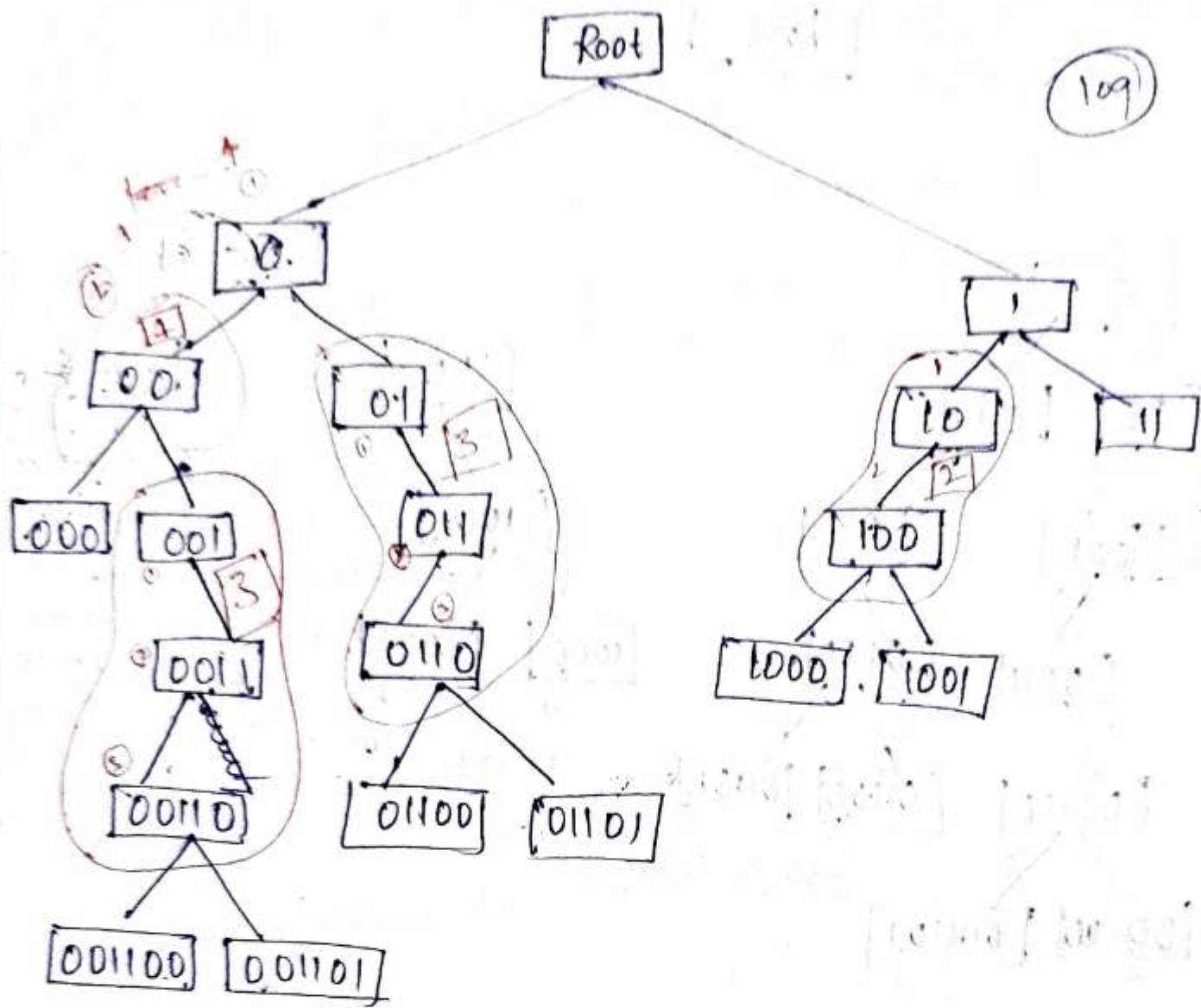


Fig: Huffman tree skipping bits for "100110001101"

## (\*) Signature file structure :-

A signature is a compressed version of a database.

All the signatures that represent the documents are kept in a file called "signature file".

### Goal :

The goal of a signature file structure is to provide a 'fast test' to eliminate the majority of items that are 'not related' to a query.

The items that satisfies the test can either be evaluated by another search algorithm to eliminate the additional false delivered to the user to review.

The text of the items is represented in a 'highly compressed form' that facilitate the fast test.

Because, file structure is 'highly compressed' and 'unordered', so it requires significantly less space than inverted file structure, and

New items can be concatenated to the end of the structure. 'Versus' significant inversion list update.

In this, the items are deleted from information database marked as deleted.



Signature file Search is a Linear Scan of compressed version of items that producing the response time linear with respect to file size.

### Creation of Signature file:-

The signature file typically uses the "Superimposed Coding" to create the signature of the document.

The coding is based up on the Words in the item. The words are mapped into "Word Signatures".

### Word Signatures-

A Word Signature is a fixed-length code with a fixed number of bits set to "1".

The bit positions that are set to "1" are determined via a "hash function" of the word.

The longer <sup>to long</sup> signatures with "1" partitioned into "blocks" of size.

For Example:- The Block size is set at five words, the code length is 16-bits and number of bits that are allowed to be "1" for each word is five.

- i.e. here \*
- \* Block size is indicated by (D) → Distinct-non common words
  - \* Code-length/fixed length indicated by (F)
  - \* Bits per word is <sup>set to one</sup> indicated by (m)

method-1

(1/2)

TEXT : Computer science graduate students study (assume block size is five words).

WORD

signature

Computer

0001

0110

0000

0110

Science

1001

0000

1110

0000

Graduate

1000

0101

0100

0010

Students

0000

0111

1000

0100

Study

0000

0110

0110

0100

Block signature :

1001

0111

1110

0110

sig: superimposed coding.

Here (1) Block size = "5" five words

five words are: (1) Computer, (2) science, (3) graduate, (4) students, (5) Study.

(F) code-length = "16-bits"

$$16 = 2^4$$

so bit length is "4"

(m) no. of bits allowed for "1" is "5" (no. of bits per word)

\* The search of the signature file requires  $O(N)$  search time.



for example: 2

11 - method

(113)

TEXT : DataBase Management system (assume block size is 4)

code length / fixed-length of bits = 20-bits

WORD

signature

(19)

|                 |      |      |      |      |      |
|-----------------|------|------|------|------|------|
| Data            | 0000 | 0000 | 0000 | 0010 | 0000 |
| Base            | 0000 | 0001 | 0000 | 0000 | 0000 |
| management      | 0000 | 1000 | 0000 | 0000 | 0000 |
| system          | 0000 | 0000 | 0000 | 0000 | 1000 |
| Block Signature | 0000 | 1001 | 0000 | 0010 | 1000 |

Here

D  $\rightarrow$  Distinct non-common Words  $D=4$

They are: <sup>①</sup>Data, <sup>②</sup>Base, <sup>③</sup>management, <sup>④</sup>system

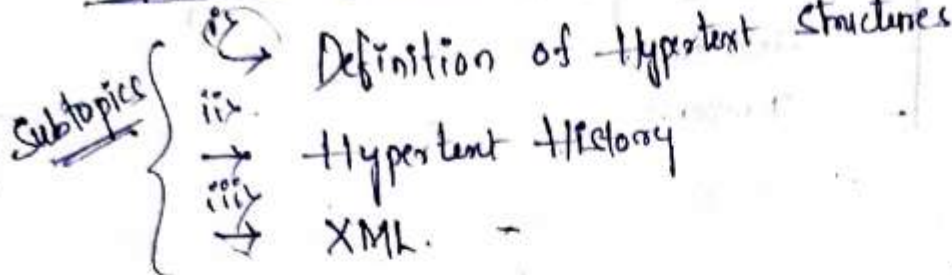
F  $\rightarrow$  Signature size / fixed-length bits

$F=20\text{-bits}$

m  $\rightarrow$  (How many bits set to 1) no. of bits per Word.

- $m=4$
- $\Rightarrow$  The signatures files have been used in following environments:
- i) PC-based media size databases.
  - ii) WORM devices Write Once Read many (Optical disks).
  - iii) Parallel machines.
  - iv) distributed environments.

## ① Hypertext and XML Data Structures



- \* Hypertext document is a "Non-linear" document with many links.
- \* Father of Hypertext is "Vannevar Bush" - 1985.
- \* He wrote an article in 1945 on Hypertext.

### Introduction / Basic points :

The advent / arrival of the Internet and its exponential growth and "wide Acceptance" has introduced new mechanisms for representing information.

This structure is called "Hypertext" and it is different from traditional information storage data structure in format and use.

The Hypertext is stored in Hypertext Markup Language (HTML) and extensible Markup language (XML).

### 1. Definition of Hypertext :-

The Hypertext data structure is used extensively in the Internet environment and it requires an electronic media storage for the files. (or)

→ The HTML defines the internal structure for information exchange across the World Wide Web on the Internet.

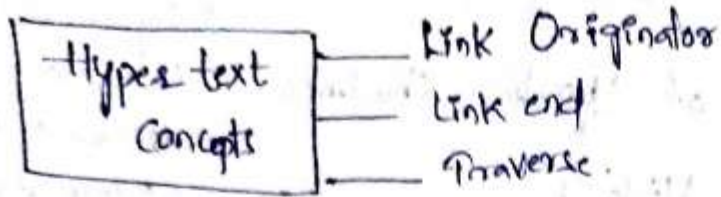
→ The HTML document is composed of text of the files along with HTML tags that describe how to display the document.

→ Tags are expressed b/w "<" and ">" symbols.  
i.e. <body> </body> <h1> </h1>



## Hyper-text concepts :

115



### Link Originator :-

- \* It is the starting point to the link.
- \* It is surrounded by the "symbols" to indicate that it is a hypertext link.
- \* The link-Originator is an anchors.

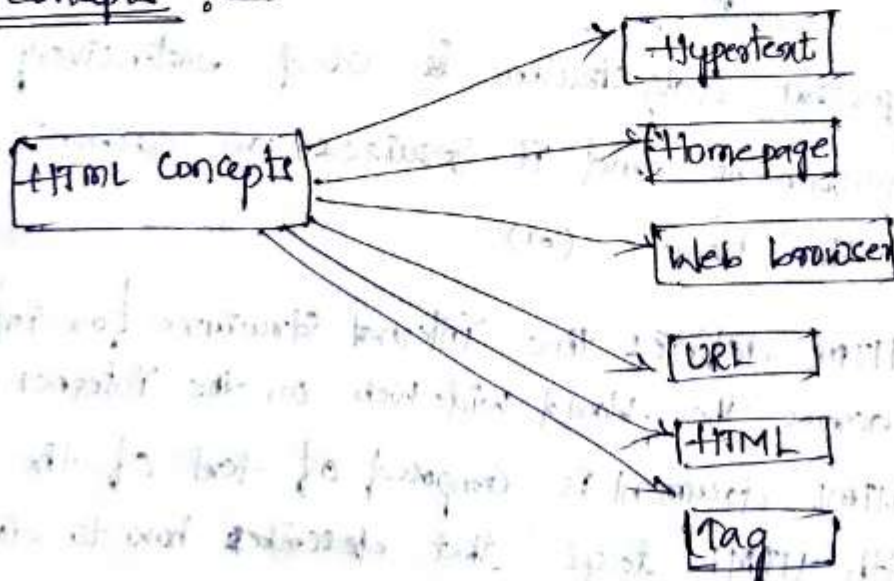
### Link end :-

The other side of the link Originator is link end.

### Traverse :-

The act of travelling from a link Originator to its associated link end.

## HTML concepts :-



Hyper text :-

Text <sup>with</sup> links the other text.

Homepage :-

The Root of the Hypertext structure.

Web browser :-

A tool to retrieve a document using "URL" from the Web server.

URL :- Uniform Resource Locator.

The address of Web page is called "URL".

HTML :- Hypertext Markup Language.

<sup>the</sup> HTML language is used to design the webpage.

Tags :-

HTML is made with the "tags" and they are enclosed in "angular brackets" i.e. < >

iii) XML :-

The extensible markup language is starting to become a standard data structure on the web.

Its first recommendation (1.0) was issued on February 10<sup>th</sup> 1998.

Hyperlinks for XML are being defined in the "XLink" (XML linking language) and XPointer (XML pointer language) Specifications.



XML Objective is extending HTML with Semantic Information.

The logical data structure within XML defined by a (DTD) Data Type Description

XML was designed to store and transport the data.  
→ It is designed to be both human and machine readable.

The ~~late~~ W3C (World Wide Web Consortium) is developing HTML as a suite of XML tags.

The following is a simple example of XML tagging:

<company> Widgets </company>

<city> Hyderabad </city>

<state> TS </state>

<product> Widgets </product>

The W3C is also developing a Resource Description Format (RDF) for representing Properties.

## ④ Hidden Markov Models :- (HMM)

118

In Hidden Markov Model the data is hidden to you / unknown to you.

Q: Why it is called Hidden Markov model?

A: The reason is to constructing an interface model based on the assumptions of Markov process.

The Markov process assumption is simply that "future is independent of the past given the present".

\* Andrei Andreyevich Markov (1856-1922)

Andrei Andreyevich Markov was a Russian Mathematician. He is best known for his work on stochastic processes.

A Primary Subject of his research later became known as "Markov chains" and "Markov process".   
 → random probability distribution

→ Hidden Markov models (HMM) have been applied for the last 20 years to solving problems in "Speech recognition", and lesser extent in the areas

\* locating named entities (Bikel-97),

\* Optical character recognition (Bazraf-98) &

\* topic identification (Kubala-97).



⇒ More recently HMM's have been applied to information retrieval search with good results. (119)

⇒ One of the first comprehensive and practical descriptions of Hidden Markov Models was written by Dr. Lawrence Rabiner (Rabiner-89).

The easiest way to understand HMM is by an example.

Example:

Let's take the example of three state Markov Model of the "stock Market".

The states will be one of the following that is observed at the closing of the Market.

i.e. State 1 ( $S_1$ ): Market Decreased.

State 2 ( $S_2$ ): Market did not change.

State 3 ( $S_3$ ): Market Increased in value.

The movement between "states" can be defined by a "state transition matrix" with "state transitions" (this assumes you can go from any state to any other state):

$$A = \{a_{ij}\} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}$$

$$A = \{a_{ij}\} = \begin{bmatrix} 0.5 & 0.3 & 0.4 \\ 0.1 & 0.6 & 0.3 \\ 0.6 & 0.1 & 0.5 \end{bmatrix}$$

(120)

Given that the Market fall (fell) on one day (state-1), the matrix suggests that, the probability of the market not changing the next day is 0.1.

then, it is allow questions such as

The probability that the Market will increase for the next 4 days then fall.

This would be equivalent to the sequence of

$$SEQ = \{S_3, S_3, S_3, S_3, S_1\}$$

In Order to Simply this model,

Let's assume <sup>that</sup> instead of the current state being dependent upon all the previous states, let's assume it is only dependent upon last state. (This order is called discrete, first order Markov chain).



⇒ It is calculated by the formula:

$$P(SFQ) = P[S_3, S_3, S_3, S_3, S_1]$$

$$\Rightarrow P[S_3] \times P[S_3/S_3] \times P[S_3/S_3] \times P[S_3/S_3] \times P[S_1/S_3]$$

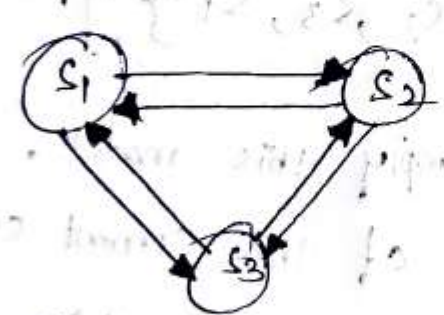
$$\Rightarrow S_3(\text{init}) \times a_{3,3} \times a_{3,3} \times a_{3,3} \times a_{1,3}$$

$$\Rightarrow (1.0) \times (0.5) \times (0.5) \times (0.5) \times (0.4)$$

$$\Rightarrow 0.05$$

⇒ In this equation we also assume the probability of the initial state  $S_3$  is  $S_3(\text{init}) = 1.0$ .

The following "graph" depicts the Model.



From Graph:

- \* The Directed lines indicates the state transition Probabilities  $a_{ij}$ .
- \* There also an implicit loop from every state back to itself.
- \* In example every state corresponded to Observable event. (i.e. change in the Market).

## Formal Definition of HMM:-

(122)

A more formal definition of a discrete Hidden Markov Model is summarized by Mittendorf and Schauble (Mittendorf-94): as consisting of the following.

1.  $S = \{s_0, \dots, s_{n-1}\}$  as a finite set of states  
Where " $s_0$ " is always indicates the "initial state".

Typically these states are "inter connected" such that any state can be reached from any other state.

2.  $V = \{v_0, \dots, v_{m-1}\}$  is a finite set of output symbols.

3.  $A = S \times S$  is a transition probability Matrix

Where  $a_{ij}$  represents the probability of transitioning from state "i" to state "j"

Such that  $\sum_{j=0}^{n-1} a_{ij} = 1$  for all  $(\forall) i = 0, \dots, n-1$ .

\* Every value in the Matrix is a positive value between 0 and 1.

\* For the case;

Where every state can be reached from every other state, every value in the matrix will be non-zero.



4.  $B = O \times V$  is an Output Probability Matrix, (123)

Where element  $b_{j,k}$  is a function determining the probability  
and  $\sum_{k=0}^{m-1} b_{j,k} = 1$  for  $V$  (for all)  $j = 0, \dots, n-1$ .

5. The Initial state distribution.

From the given HMM Definition;  
it can be used for both sequence of output and their  
Probabilities.

The complete specification of HMM requires:

- \* Specification of states,
- \* Output Symbols &
- \* Three Probability Measures for
  - ①  $\rightarrow$  The state transitions,
  - ②  $\rightarrow$  Output probability functions
  - ③  $\rightarrow$  The initial states.

The distributions are frequently called  $A$ ,  $B$  and  $\pi$ .

The following notation is used to define the model:

$$\lambda = (A, B, \pi)$$

## ⑧ Stemming Algorithm :-

124

→ Stemming means cutting (or) trimming the unnecessary data.

→ Stemming is a technique used for base/form abstract the base form of affixes.

Affixes indicate the ending terms.

→ Stemming Programs are commonly referred as Stemming Algorithms and Stemmers.

The "main goal" is to improve the performance of the system.

"Stemming is just like cutting the branches of trees".

### Examples :-

1. eating  
eaten  
eate

Eat

2. Calculations

Calculate

Calculating

Calculate

3. Computerise

Computing

computer

computerizing

compute



Sub topics of stemming algorithms are:

125

- i) Porter stemming Algorithm
- ii) Dictionary look-up stemmers
- iii) Successor stemmers.
- iv) conclusions.

i) Porter stemming Algorithm :-

\* It is introduced in 1980 to perform "actions" based on "Conditions".

\* It is a rule based "Algorithm", based on conditions.

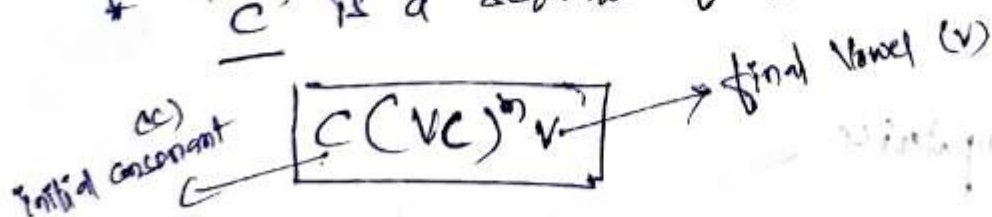
\* The <sup>stemming/Porter</sup> algorithm <sup>based on</sup> contains the <sup>root</sup> stem, suffix, Prefix and associate conditions.

Conditions :-

① The measure "m" of a stem is a function of sequence of Vowels (a, e, i, o, u) followed by a consonant.

\* If "V" is a sequence of "Vowels" and

\* "C" is a sequence of "consonants", then "m" is:



Where

initial "C" and final "V" are Optional and

"m" is the number of "VC" repeated.

# Examples :-

126

① FREE, Why

find the given word no. of consonants and no. of vowels:

Consonants: CC VV  
FREE

Why  
↓↓↓  
CCC

m=0

Hence  $m=0$  because there is no "VC" combination.

② FREES, whose

FREES  
↓↓↓ ↓↓  
CCVVC

whose  
↓↓↓↓↓  
CCVVCV

Hence  $m=1$  because there is 1 "VC" combination.

③ Prologue, compute, Private

for above words find no. of vowels & consonants and later find the no. of (VC) combination.

Prologue  
↓↓↓↓↓ ↓↓  
CCVVCVV

compute  
↓↓↓↓↓ ↓↓  
CCVCCV

Private  
↓↓↓↓↓ ↓↓  
CCVVCVV



② \* <X> :

This statement is indicating "stem ending with letter X."

Example :-

Index, tax, Xerox, prefix, suffix etc.

③ \*V\* :

Stem contains a Vowel.

Example :

pen → vowel.  
↓ ↓ ↓  
CVC, Pin ... etc

④ \*d :-

Stem ending with double consonant.

Example :

sky, dry, strong  
↓ ↓ ↓   ↓ ↓ ↓   ↓ ↓ ↓ ↓ ↓  
C C C   C C C   C C C V C C  
          double  
          consonants

⑤ \*O :-

It indicates the stem ends with Consonant + Vowel -  
Consonant Sequence (C-V-C) where final consonant

is not W, X, (or) Y.

Example :

~~Pen~~  
~~↓ ↓ ↓~~  
~~CVC~~

(\*)

Pen  
↓ ↓ ↓  
CVC ⇒ C-V-C sequence  
last consonant is  
is not w, x, y.

## Rules :

128

The Rules are divided into steps to define the order of applying the rules.

The following are some examples of rules:

### PORTER STEMMER STEPS

Step 1  $\Rightarrow$  (A), (B), (C)

(A)

i)  $SES \rightarrow S$

Eg: Process  $\rightarrow$  process, Stresses  $\rightarrow$  stress

ii)  $RES \rightarrow R$

Eg: ties  $\rightarrow$  ti

iii)  $SS \rightarrow S$

Eg: class  $\rightarrow$  class

iv)  $S \rightarrow \epsilon$

Eg: Cats  $\rightarrow$  Cat  $\Rightarrow$  i.e. cat

(B)

i)  $(m > 1) EED \rightarrow EE$

If the stem is ending with "EED" is replaced by "EE".

Example :-

\* Agreed

$\rightarrow$  For the above word find the no. of (VC) combination

So that we get "m" value. If the "m" value is ( $>$ ) greater than 1 then it satisfies the above condition.



ie Agreed  
 ↓ ↓ ↓ ↓ ↓ ↓  
 V C C V V C  
 ① ②

$m=2$  for above word 'm' value is '2'.

so that it satisfies the condition.

$m=2$

~~m=2~~ (mx) EED  $\rightarrow$  EE

$\rightarrow (2 \times 1) \text{ Agreed} \rightarrow \text{Agree}$

ii)  $*V*ED \rightarrow \epsilon$

If the stem containing vowel, followed by 'ED' is replaced by ' $\epsilon$ ' (empty).

Example: Died  $\rightarrow$  Di  
 $\rightarrow$  Di  $\epsilon \Rightarrow$  Di

iii)  $(*V*)ING \rightarrow \epsilon$

If the stem ending with 'ING' is replaced by ' $\epsilon$ '.

Example: Dancing  $\rightarrow$  Dance  $\epsilon \Rightarrow$  Dance

\* Dancing  $\rightarrow$  Dance  $\epsilon \Rightarrow$  Dance

Making  $\rightarrow$  Mak  $\epsilon \Rightarrow$  Mak

\* Making  $\rightarrow$  Mak  $\epsilon \Rightarrow$  Mak

$\Rightarrow$  walking  $\rightarrow$  walk

Searching  $\rightarrow$  search

(C)

136

i)  $AT \rightarrow ATE$

Ex:  $\star$  Conflated

conflated  $\rightarrow$  conflates  $\left[ \because \text{from } 1b \text{ condition} \right]$   
 $\star V \star ED \rightarrow \epsilon$

i.e. conflated  $\rightarrow$  conflat  $\rightarrow$  conflate  $\left[ \because \text{from } 1c \text{ condition} \right]$   
 $AT \rightarrow ATE$

ii)  $BL \rightarrow BLE$

Example : 1 : Troubling  $\rightarrow$  another condition i.e. 1c

for ~~to~~ Word Troubling  $\rightarrow$  1b condition from.

i.e. Troubling  $\rightarrow$  Troubl  $\left[ \because \text{1b condition} \right]$   
 $\star V \star \text{ing} \rightarrow \epsilon$

Remaining Words: Troubl

Troubl  $\rightarrow$  Trouble  $\left[ \because BL \rightarrow BLE \text{ 1c condition} \right]$

Example : 2 "Bubbled"

$\rightarrow$  Bubbled find "BL" in word i.e. Bubbled

Bubbled  $\rightarrow$  Bubble  $\left[ \because \star V \star ED \rightarrow \epsilon \right]$

Bubbled  $\rightarrow$  Bubbl  $\rightarrow$  Bubble  $\left[ \because BL \rightarrow BLE \right]$



iii) "(\*d ! (\*L (or) \*s (or) \*z)) → single character." (131)

The stem is ending with "double consonant" and that consonant is "not L (or) s (or) z?"

Example :-

① Hopping

↓↓↓↓↓↓↓  
CVCCVCC

Hopping → Hopp<sup>→ null</sup>ε [∵ \*V\*ing → ε]

<sup>ε V CC</sup>  
↑↑↑  
Hopp → Hop<sup>→ single consonant</sup>

→ single character

∴ Hopping → Hopp → Hop

② Panned

Panned → Pann<sup>→ ε</sup> [∵ \*V\*Ed → ε  
ed → ε]

Remaining stem: Pann

<sup>ε V CC</sup>  
↑↑↑  
Pann → pan

Ending with double consonant is replaced by single character.

Pann → Pan

∴ Panned → Pann → Pan

③ swimming

→ swimming → swim  $[\because *V*ing \rightarrow \epsilon]$

Remaining stem 'swimm':

swimm → swim  
 $\downarrow \downarrow \downarrow \downarrow$   
 CC VCC  
 ↳ double consonants replaced by single consonant

∴ swimming → swim → swim

④ falling

falling → fall  $[\because *V*ing \rightarrow \epsilon]$

Remaining stem : fall :

Fall → Fal (Condition not Verified)

∴ Falling → Fall → Fal\*

IV)  $(m=1 \& *0) \rightarrow F$

→ stem ending with ~~double consonant~~ C-V-C sequence and final consonant is not W, X, Y

The measure "m" (VC combination) i.e 1 and the stem is ending with CV (CVC) combination is replaced by F



Example :  $\text{Filing} \xrightarrow{\text{VC}} \text{Fil} \xrightarrow{\text{VC}} \text{File}$  (initial)

(123)

①  $\text{Filing} \rightarrow \text{Fil}$  [  $\because$  \*vating  $\rightarrow$  e ]

$\text{Fil} \rightarrow \text{File}$  [ cvc sequence followed by E ]

(VC)  $m=1$

$\therefore$  VC combination  $m=1$ , cvc followed by E)

Filing  $\rightarrow$  Fil  $\rightarrow$  File Condition Verified.

② fail  $\rightarrow$  fail (not Verified)

$\text{fail} \xrightarrow{\text{VC}} \text{fail}$  (VC combinations  $m=1$ )

$m=1$  but cvc sequence is not there

so condition failed.

$\nabla ( * V * ) Y \rightarrow I$

Our aim is to take out "Y" from the word and the remaining stem containing at least "one vowel".

If the condition is satisfied "Y" is replaced with "I".

Example :  $\text{Happy} \rightarrow \text{Happi}$

Step 2 :- Derivation Morphology - I

i) (m > 0) ATIONAL  $\rightarrow$  ATE

Example :

Relational  $\rightarrow$  Relate

Relational  
↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓  
CVCVCVCVCVC  
① ② ③ ④  
 $\therefore m=4$

$m > 0 \text{ i.e. } 4 > 0$

ii) (m > 0) IZATION  $\rightarrow$  IZE

Example :

Generalization  $\rightarrow$  Generalize

Generalization  
↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓  
CVCVCVCVCVCVCVC  
① ② ③ ④ ⑤ ⑥

$\therefore m=6 \Rightarrow 6 > 0 \text{ \& } IZATION \rightarrow IZE$

iii) (m > 0) BILITY  $\rightarrow$  BLE

Example :

Sensibility  $\rightarrow$  Sensible  
↓ ↓ ↓ ↓ ↓ ↓ ↓ ↓  
CVCVCVCVCVCVC  
1 2 3 4

Give combinations = 4  
 $m=4$

$\therefore m=4 \text{ i.e. } 4 > 0 \text{ \& } BILITY \rightarrow BLE$





# Step 4 Derivation Morphology - III

136

i) (m>0) ANCE  $\rightarrow \epsilon$

Example: Allowance  $\xrightarrow{\epsilon}$  Allow

Allowance  
 $\downarrow \downarrow \downarrow \downarrow \downarrow \downarrow \downarrow \downarrow$   
 VC CV VC CV  
 ① ② ③

$m=3$

$\therefore (3>0) \text{ Allowance} \xrightarrow{\epsilon} \text{Allow.}$

ii) (m>0) ENT  $\rightarrow \epsilon$

Example: Dependent  $\rightarrow$  Depend

$\downarrow \downarrow \downarrow \downarrow \downarrow \downarrow \downarrow \downarrow$   
 CV CV CV CC  
 ① ② ③

$m=3$

$\therefore (3>0) \text{ Dependent} \xrightarrow{\epsilon} \text{Depend}$

iii) (m>0) IVF  $\rightarrow \epsilon$

Example: Elective  $\rightarrow$  elect

Elective  
 $\downarrow \downarrow \downarrow \downarrow \downarrow \downarrow \downarrow \downarrow$   
 VC VCCVCV  
 ① ② ③

$m=3$

$\therefore (3>0) \text{ Elective} \xrightarrow{\epsilon} \text{Elect}$



Step (5) A & B :-

(A)

i)  $(m > 1) E \rightarrow \epsilon$

Bye  $\rightarrow$  By

Example: ~~Pro~~ Probate  $\rightarrow$  Probat

ii)  $(m > 1 \& ! * 0) NESS \rightarrow \epsilon$

condition  $m \geq 1$  and stem not ending with CVC combination.

Example: Goodness  $\rightarrow$  Good

↓↓↓↓↓↓↓↓

CVC CVC

$M=2$  Goodness  $\rightarrow$  Good

↓  
 $\epsilon$

not in CVC combination.

(B)  $((m > 1) \& * d \& * l) \rightarrow$  Single letter.

Example:

controll  $\rightarrow$  Control

↓↓↓↓↓↓↓↓

CVC CCVC

① ②

$\rightarrow$  Replaced by single letter 'l'

check the condition  $m$  (ie  $m=2$ ) and the stem is ending with double consonant and 'l' is replaced by single letter.

Example 2: Roll  $\rightarrow$  Roll

$m=1$ , Condition not verified.

The two other techniques in stemming are

• ii) Dictionary lookup stemmers

iii) successor stemmers

iv) conclusions

ii) Dictionary lookup stemmers:

An alternative way to determine a "stem" is a "Dictionary lookup mechanism".

In this approach, simple stemming rules are may be applied.

The Original term (or) stemmed version of the term is looked up in a dictionary and replaced by the stem that is the best representation of it.

This technique has been implemented in the "INQUERY and Retrieval ware system".

The INQUERY system uses a stemming technique called "k-stem".

k-stem: k-stem is a morphological analyzer that converts word variants to a root form.

Eg:- "Memorial" and "Memorize" reduced to "memory".

But memorial and "memorize" are not synonyms and they have different meanings.



K-Stem requires a word to be in the dictionary before it reduces one form to "Other". (139)

Some endings are removed, even if the root form is not found in the dictionary.

Ex: "ness", "ly"

\* K-stem uses the following "Six major" data files to control and limit the stemming process:

- ① Dictionary of Words (Lexicon).
- ② "Supplemental list" of words for the dictionary.
- ③ "Exception list" for those words that should retain an "e" at the end (Ex: "suite" → "suite" but "suited" → ~~suit~~ "suits")
- ④ "Direct - Conflation": → allow definition of direct conflation via word pairs that override the stemming algorithm.
- ⑤ "Country - Nationality":  
It indicates conflations between "nationalities and countries". ("British" <sup>maps</sup> → "Britain").
- ⑥ "Proper Nouns":  
A list of Proper Nouns that should not be stemmed.

Retrieval Index system :-

This system lies in its thesaurus/semantic network support the data structure that contains over 4,000,000 words.

### iii) Successor stemmer :-

(140)

Successor stemmers are based on the "length of prefixes" that optimally expand the stem. It is also known as "successor variety".

It is used to determine the "Word", word is divided into "Segments".

for Example :-

The Successor Variety for the first three letters (i.e. Word Segment) of a five letter word is no. of words that have the same first three letters but fourth letter is different.

Graphical Representation :-

→ A Graphical representation of Successor Variety is shown in ~~fig.~~ Symbol tree.

→ The "Symbol tree" for the same

bag, barn, bring, both, box and bottle.

→ The Successor Variety for any prefix of word is the no. of childrens that are associated with the node in the "Symbol tree". representing that prefix.

for example :- The Successor Variety for the first letter "b" is "three".

The Successor Variety for the prefix "ba" is "two".



|                                                     |   |
|-----------------------------------------------------|---|
| B                                                   | 3 |
| <u>Ba</u>                                           | 2 |
| → Baq                                               | 0 |
| Bar                                                 | 1 |
| → Barn                                              | 0 |
| <u>Bs</u>                                           | 1 |
| Bri                                                 | 1 |
| Brin                                                | 1 |
| → Bring                                             | 0 |
| <u>Bo</u> (∵ Bo followed by x and Bo followed by t) | 2 |
| → Box                                               | 0 |
| <u>Bot</u>                                          | 2 |
| Both                                                | 0 |
| Bott                                                | 1 |
| Botta                                               | 1 |
| → Bottle                                            | 0 |

Fig: Successor Vari

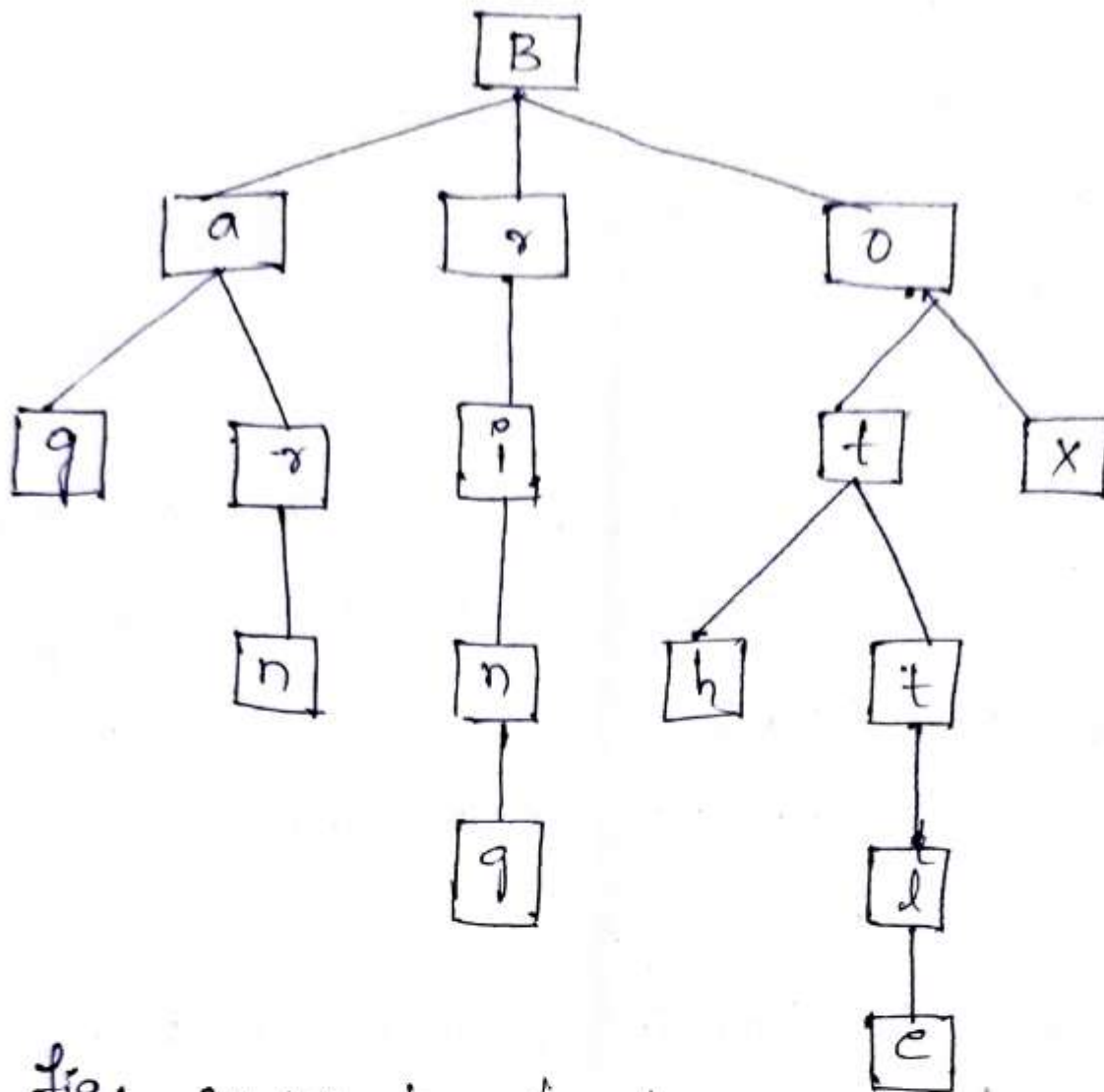


Fig: Symbol tree for terms bag, barn, bring, box, bottle and both.

#### Successor Variety methods :-

The Successor Varieties of a words are used to segment a word by applying ~~at~~ one of the following 4-methods:

1. Cutoff Method
2. Peak and plateau
3. Complete Word Method
4. Entropy Method.



### ① cutoff Method :-

cutoff value is selected to define "stem length" and the values are varies from word-to-word / each possible set of words.

### ② Peak and plateau :-

"a segment break" is made after a character whose successor variety exceeds that character and immediately preceding it.

### ③ Complete Word Method :-

Break on boundaries of complete words.

### ④ Entropy Method :-

Used to Distribution of Successor Variety letters.  
(Let  $|D_k|$  be no. of words begin with "k" length sequence of letters "a".)

### iv) Conclusions

Frakes summarized the study of various stemming.

Frakes came to following conclusions:-

ii) Stemming can effect retrieval (recall) and where effects were identified they ~~were~~ were positive.

- ii) Stemming is as effective as manual collation.
- iii) Stemming is dependent upon the nature of the Vocabulary.

Extra Information:

⇒ To quantify the impact of stemmers, "parce" has defined a stemming performance measure, called "ERROR RATE RELATIVE TO TRUNCATION (ERRT)".

ERRT :- ERRT can be used to compare stemming algorithms.

→ This approach depends upon "concept groups".

After applying a stemmer that is not perfect, "concept groups may still contain multiple stems rather than one."

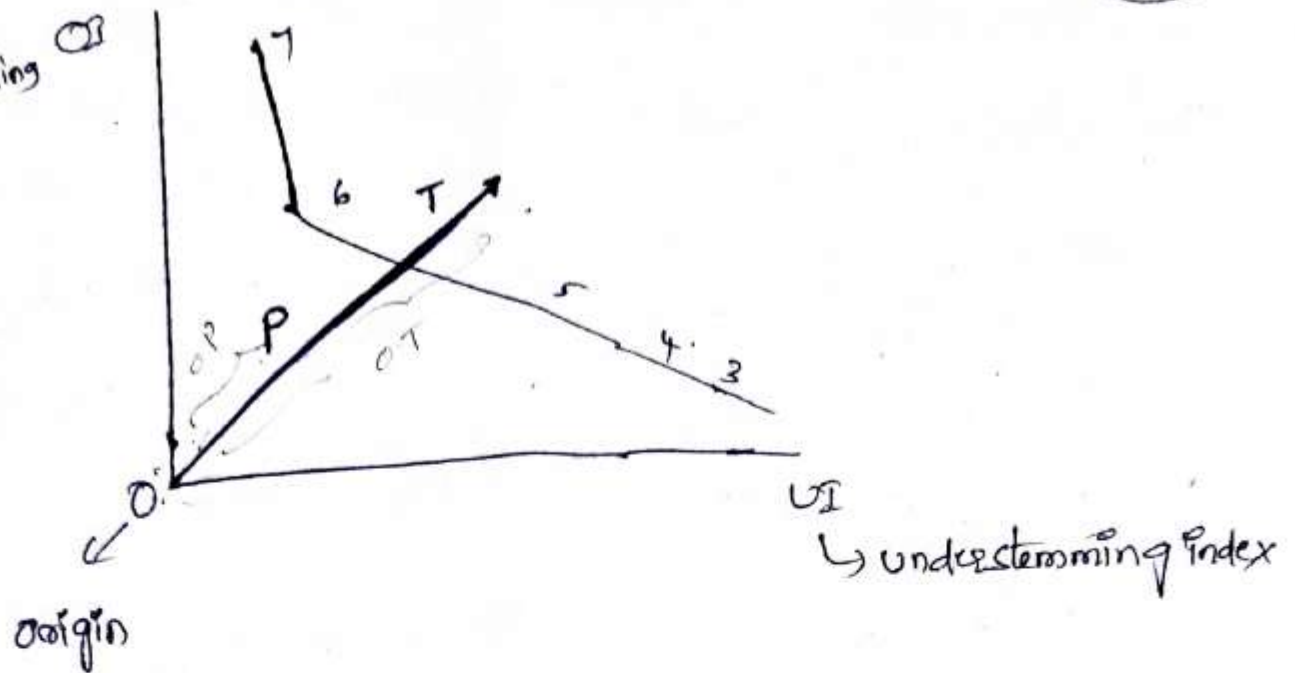
→ This introduces an error reflected in the understemming index (UI).

This is also possible that same stem is found in multiple groups, is <sup>error state</sup> reflected in the overstemming index (OI).

⇒ The UI & OI values can be reflected calculated based on truncated word lengths.

The Perfect case is where  $UI \leq OI = 0$  (zero).





$\Rightarrow$  FRRT is calculated as the distance from the Origin to the  $(U_i, O_i)$  coordinate of the steamer being evaluated (Op) versus distance from origin to the Pure termination (OT).