

10-12-19

Introduction

data:- collection of raw facts.



processed data

unprocessed data

- The data is in the form of textual, video, audio, number, pictorial
- collection of information or collection of records is known as a file.
- Database is systematic collection of data. Database support

Database - relations - table - rows / cols.

Application software - users uses these applications

system software - in built - like compilers, translators, interpreters.

warehouse:- repository of heterogeneous types of data
↓
storing

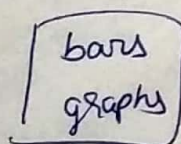
- Large amount of data we can store in it.

Warehouse - gb & Tb of data to be stored

In warehouse there is lot of data to be stored and retrieve the data also.

↓
mining - extracting

or we can tool



outlook data



patterns

different types of data are

db of data

file of data

spread sheet

data warehouse

↓
data about data

db

spread sheet of data

relational data

↓
warehouse

11-12

Unit 1 - Introduction to Data Mining

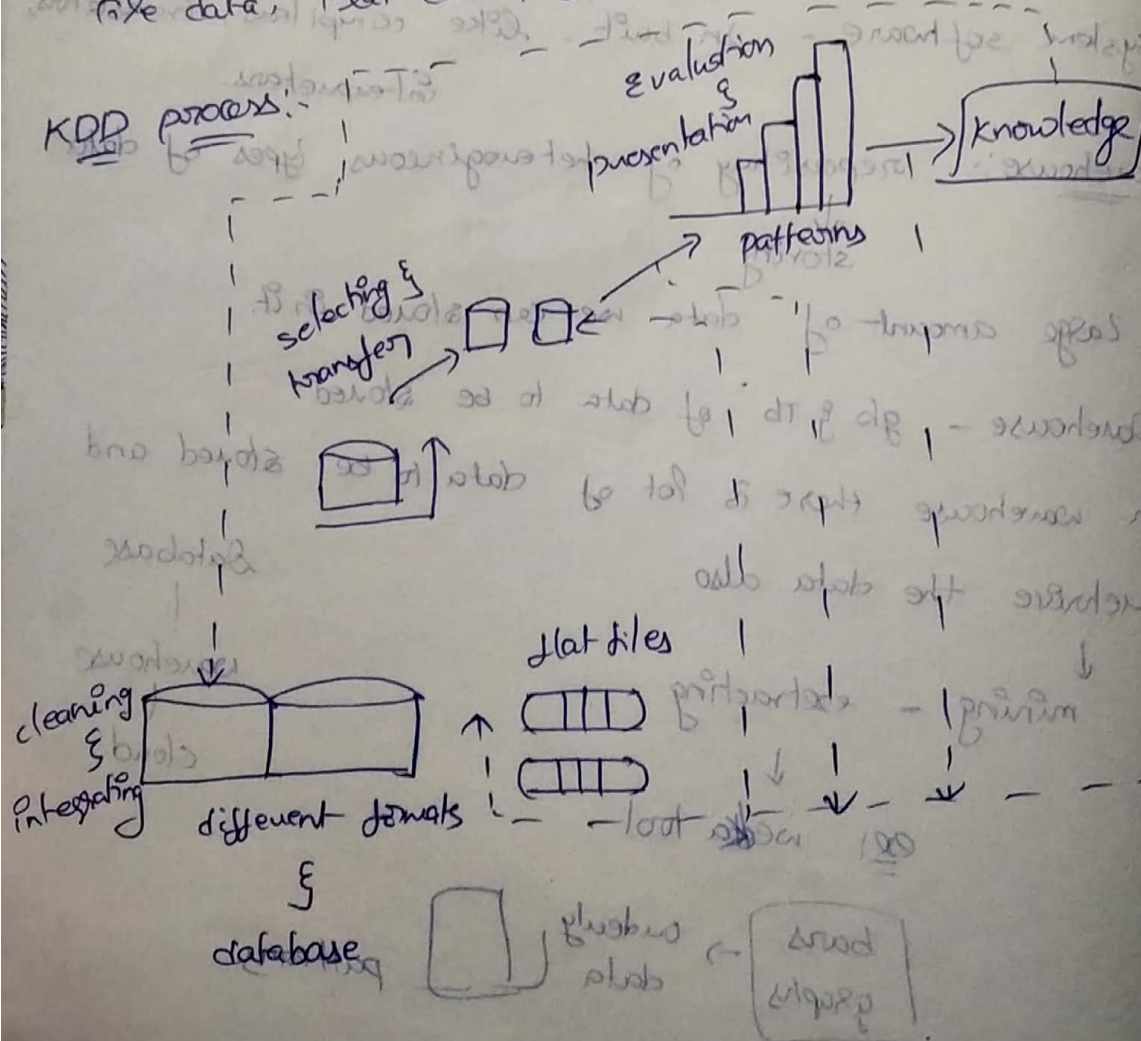
1st topic:- Knowledge discovery in database

→ extract the data is known as data mining

→ From the warehouse we can store data

→ warehouse - repository of heterogeneous types of data
file data, flat data, data warehouse, spread sheet type of data

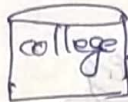
KDD process:-



1. Flat

1. Data integration - combining of data it can store different types of data like data base, flat file, spread sheet and data mart. In data base we can store relational form of a data. A flat is also a per line we can take only one record is known as flat file. This flat file doesn't contain any ^{type of} relation to another file, table.

→ Data mart is nothing, but sub-section of data warehouse

ex:-  Datawarehouse

2. Data cleaning - It remove the noisy, inconsistency, duplication. To remove these data using normalization, for the cleaning purpose.

S. No	name	course
1	A	c, c++
2	B	Java, python

1 A c

1 A c++

2 B Java

2 B python

Like above table we remove the noisy & duplicate

data with the use of 2nd normal form.

3. Data selecting & transformation -

→ 90% of the cleaned data should be selected for the transformation.

* For selecting the cleaned data we are using some techniques

they are called pre-processing techniques.

The following are come under the pre-processed techniques

they are 1. feature selection

2. Dimensionality reduction

3. Normalisation

4. Data subsetting

→ Transformation is nothing but convert the one type of format of data to another

ex:- dataset - collection of data
Data mining:- ^{post processing} ^{classification} ^{cluster} large - massive

Data mining is the process of mine the data i.e., post processing can be applied here i.e., on the transformed data we can apply the post processing techniques they are filtering patterns, visualization, pattern interpretation. Then finally from the visualization or from the pattern or from the evaluation the knowledge is an happen or the useful information is provided.

2nd topic:- motivating challenges in Data mining
→ strength

The challenges are

1. scalability
2. High dimensionality Reduction
3. Heterogeneous complex data
4. Data ownership Central & distributed
5. Key challenges

1) scalability - scalable is nothing but big in size we can extract huge amount of data at data warehouse we can't extract the massive data with help of previous algorithms.

scalability is nothing but one of the challenges faced by data mining.

1. Originally the data warehouse can accommodate the data in terms of GB & TB i.e., is in scalable manner (Big in size)
2. But the massive of data sets of data extraction

Is a challenge faced by existing data mining algorithms like classification algorithm, clustering algorithm, Aggregation algorithm

3. For these we need to develop scalable data mining algorithms to mine massive data sets

2. High dimensionality :-

Dimensionality - Data, Dataset, attributes
↳ dimension

The types of data - qualitative, quantitative
character number

→ nominal & ordinal

→ qualitative - categorized data - classification

nominal - categorized

ordinal - un "

numbers - discrete, continuous

discrete - measuring the finite values

continuous - measuring the infinite values.

a) Data set have the attributes

b) when the dataset contain the attribute it is called as dimension

c) At present we have the data sets with minimum number of attributes for data processing is possible but it is not support the data set with 100's or 1000's of attributes for this it is the challenge faced by data mining due to this need to develop data mining algorithms to handle high dimensionality.

d) It uses low dimensionality i.e., it supports the data is in the form of continuous & categorical

3. Heterogeneous & complex data:-

- * Traditional data analysis methods deal with data sets it contains same type of attributes.
- * But the complex data set contain diff types of data i.e. it supports image data, video data, text video.
- * Need to develop Mining methods to handle complex data sets.

4. Data ownership & distribution:-

- * Here sometimes the data needed for an analysis is not stored in one location (or) owned by one organisation instead the data is Geo-graphically distributed among resources belonging to multiple entities.

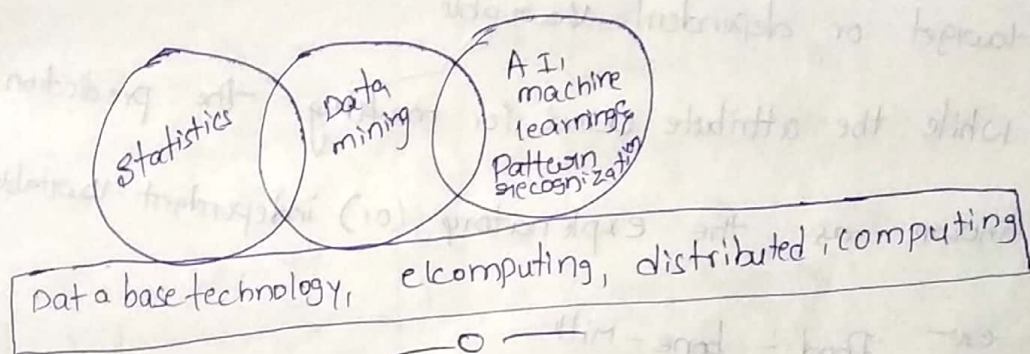
5. Key Challenges:-

- * Scalable is one of the key challenge
- * High dimensional is another "
- * Ownership & distribution is another "
- * Heterogeneous & complex data " " "

3rd topic The Origins of data mining

- * The Origin of data mining is possible through the diff domains.
- * From the statistics domain it can be happen from AI (Artificial Intelligence) it was happen from machine learning it was happen from pattern recognition it was happen.

- * Sampling, estimation & hypothesis testing form statistics from all of these the data mining coming to picture.
- * Search algorithms, modeling techniques, learning theory machine language from AI, (ML) & pattern recognition.
- * Data base Systems are used for storage effective, index & query processing effective, index & query processing effective storage.
- * Parallel computing is used for massive or some data sets address (Bulk amount)
- * According distributed Computing gather the info from various locations



with topic Data mining Tasks

Task meaning work

* They are two types of Data mining task

1. Predictive

2. Descriptive

Predictive:- we can estimate the data (variable data may

Processed / mining has two types.

1. Pre process

2. Post process

Preprocessing:- We can remove the large, noisy data using Normalization is called preprocessing.

Post processing:- We can select, transform the data we can appear some algorithms is called post processing.

* We can predict the value of particular attribute based on the value of other attributes.

* An attribute can have not a single value have the multiple values.

* The attribute to be predicted is commonly known as the target or dependent variable.

* While the attribute used for making the prediction or known as the explanatory (or) independent variable.

ex:- Food - bone - milk

Sound - bark - New

1.b

Descriptives:- It is nothing but designed the pattern.

* Pattern is nothing but some standard structure of the data.

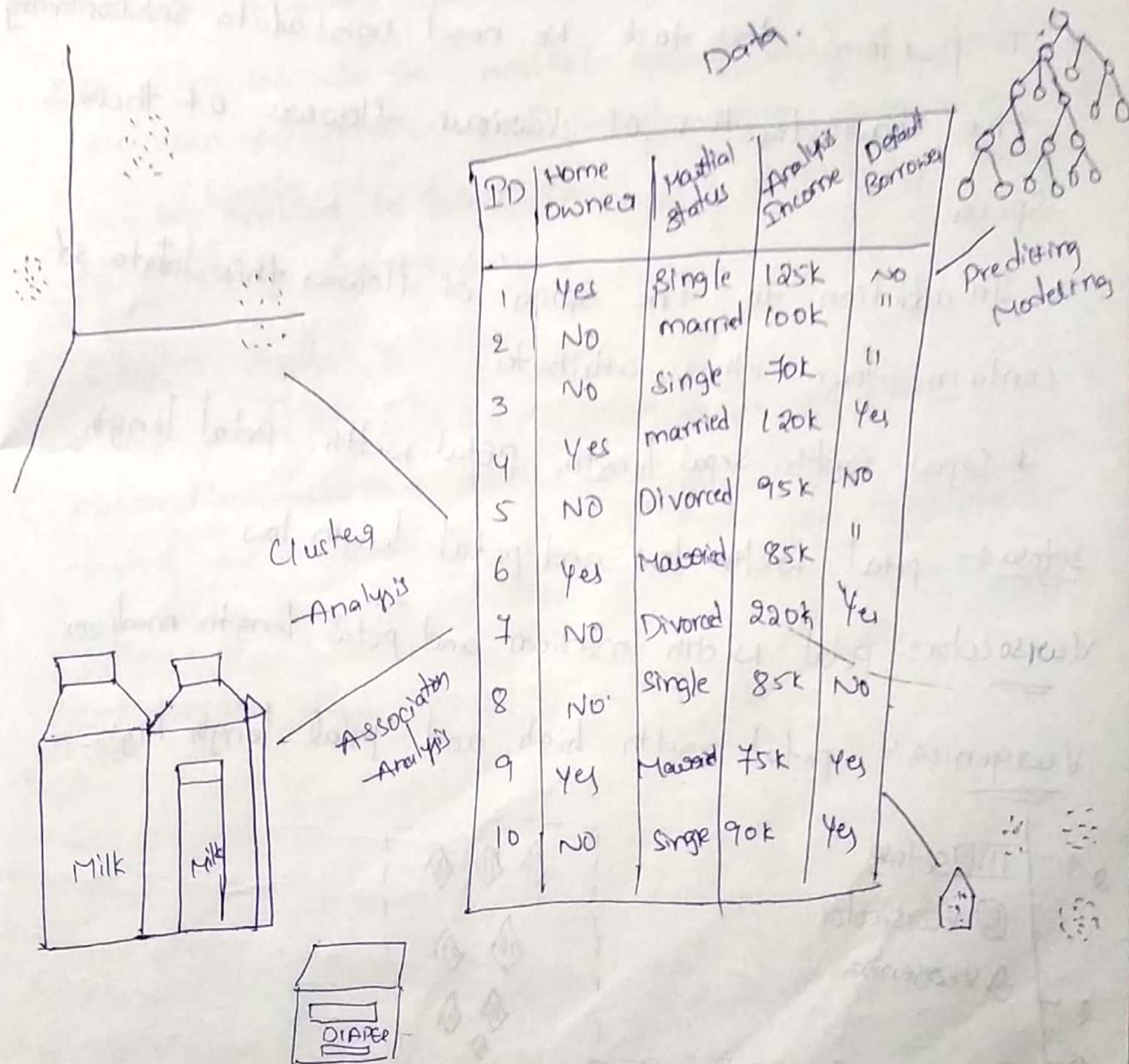
* It is used to summarize the underline relation in data.

* Whether the data is in the pattern of pre-processed data is in the predictive model. It can

answer the what type of data it can belongs

* we have -furtherly four core data mining tasks they are

1. Predictive modeling
2. Association Analysis
3. Clustering Analysis
4. Anomaly detection



Four of the core data mining Tasks

1. Predictive Modeling:-

* Consider an "Iris" flower.

* Classifying the Iris flower whether it belongs to one of the following three Iris species.

Setosa, Versicolor and Virginica

* To perform this task we need need a data set containing the characteristics of various flowers of these 3 species.

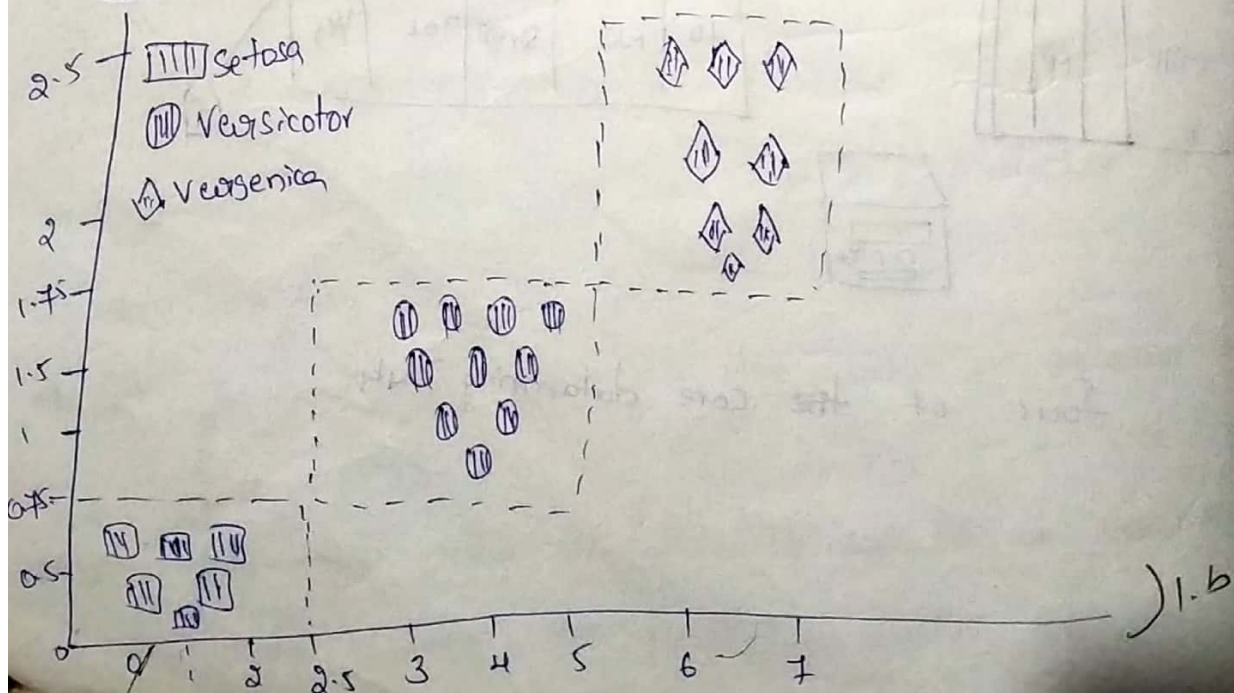
* In addition to the species of flower this data set contain four other attributes

* Sepal width, Sepal length, petal width, petal length.

Setosa:- petal width low and petal length low

Versicolor:- petal width medium and petal length medium

Virginica:- petal width high and petal length high.



2. Association Analysis :-

- Association Analysis is used to discover patterns (order) that describes strongly associated features in the data.
 - The discovered patterns are typically represented in the form of rules.
 - The goal of association analysis is to extract the most interesting patterns in an efficient manner.
- ex: we can use the market basket analysis. Here we can concentrate the grocery store association analysis can be applied to find items are frequently bought together by customers.

3. Cluster Analysis :-

- This analysis is used to find groups of closely related observations so that observation that belongs to the same cluster are more similar to each other than observations that belong to other clusters.

ex: we use the document clustering.

market basket data

Transaction id

items

- 1 {Bread, Butter, Diapers, milk}
- 2 {coffee, sugar, cookies, salmon}
- 3 {Bread, Butter, coffee, Diapers, milk, eggs}
- 4 {Bread, Butter, salmon, chicken}
- 5 {eggs, Bread, Butter}
- 6 {salmon, Diapers, milk}
- 7 {coffee, sugar, chicken, eggs}
- 8 {Bread, diapers, milk, salt}
- 9 {Tea, eggs, cookies, Diapers, milk}
- 10 {Bread, Tea, sugar, eggs}

collection of news articles

Articles

words

1. dollar:1, industry:4, country:2, loan:3, deal:2, government:2
2. machinery:2, Labour:3, market:4, industry:2, work:3, country:1
3. job:5, inflation:3, rise:2, jobs:2, market:3, country:2, index:3
4. domestic:3, forecast:2, gain:1, market:2, sales:3, price:2
5. patient:4, symptom:2, drug:3, health:2, clinic:2, decision:2
6. pharmaceutical:2, company:3, day:2, vaccine:1, flu:3

4. Anomaly detection:

- Anomaly means abnormal that is not in a proper way. we can credit the anomalies of the data.
- Anomaly detection is one of the core data mine tasks. The task of this detection is identifying observations whose characteristics are significantly different from the remaining data such observations are known as anomalies (a) outliers

ex:- credit card fraud Detection.

A credit card company records the all the transactions made by every credit card holder, along with personal information such as credit limits, age, annual income, and address.

- Anomaly detection technique used can be applied to build a profile of legitimate (true) transactions for the users then any new transactions arise it can be checked with the profile of true transactions.
- The characteristics of the enabled transaction are very different with the profile it can be treated as outlier.

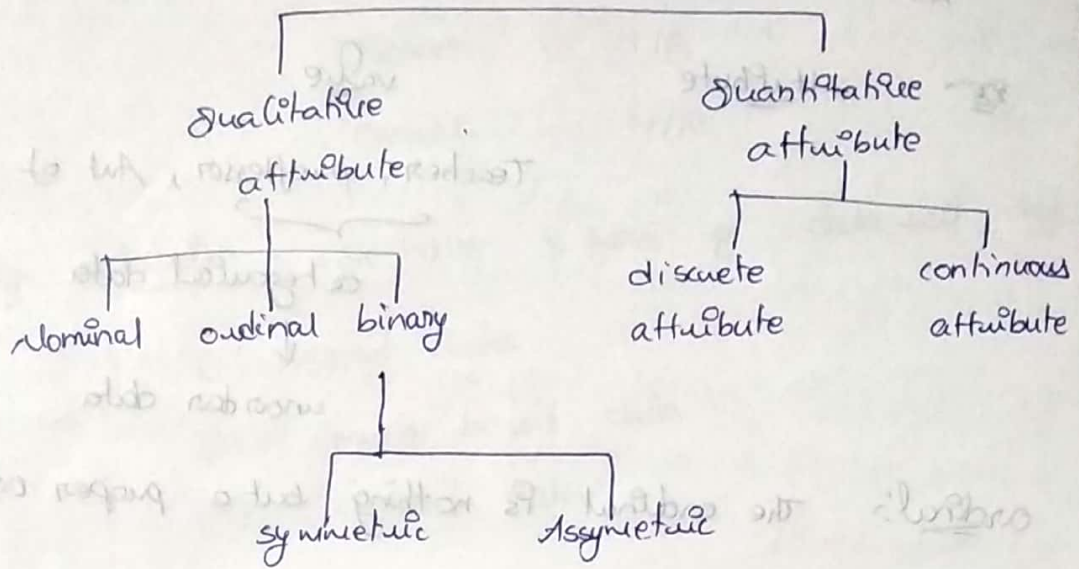
5th topic: Types of data

1. Data sets differ in a number of ways the attributes used to describe the data objects can be different types
2. we have two types of attributes they are quantitative and qualitative attributes according to the data
3. A data set is nothing but collection of data objects

Exam:

Different attribute types			
Attribute Type	Description	Examples	operations
Qualitative (Qualitative)	Nominal The values of a nominal attribute are just different names i.e., nominal values provide only enough information to distinguish one object from another ($=, \neq$)	Zip codes, Employee id numbers, eye color, gender	mode, entropy, contingency, correlation, χ^2 test
	Ordinal The values of an ordinal attribute provide enough information to order objects ($<, >$)	Hardness of minerals, {good, better, best} grades, street numbers	median, percentiles, rank & correlation, sign tests
	Interval For interval attribute, the difference b/w values are meaningful, i.e., unit of measurement exists ($+$, $-$)	calendars dates, temperature in Celsius or Fahrenheit	mean, standard deviation, Pearson's correlation t and F tests
Quantitative (Quantitative)	Ratio For ratio variables both differences and ratios are meaningful ($*$, $/$)	temperature in Kelvin, monetary, quantities, count, age, mass, length, electrical current	geometric mean, harmonic mean, percent, variation

→ Types of attributes :- The attributes are classified into two types



Qualitative attribute :- Qualitative measures something based up on the state.

Quantitative attribute :- The quantitative measures with the special characters of the data.

Discrete attribute :- The attribute contains the finite value is called discrete attribute.

Continuous attribute :- The attribute contains the infinite value is called continuous attribute.

Ratio :- It is nothing but the two values @ 1) two attributes b/w the ratio can be contained for example

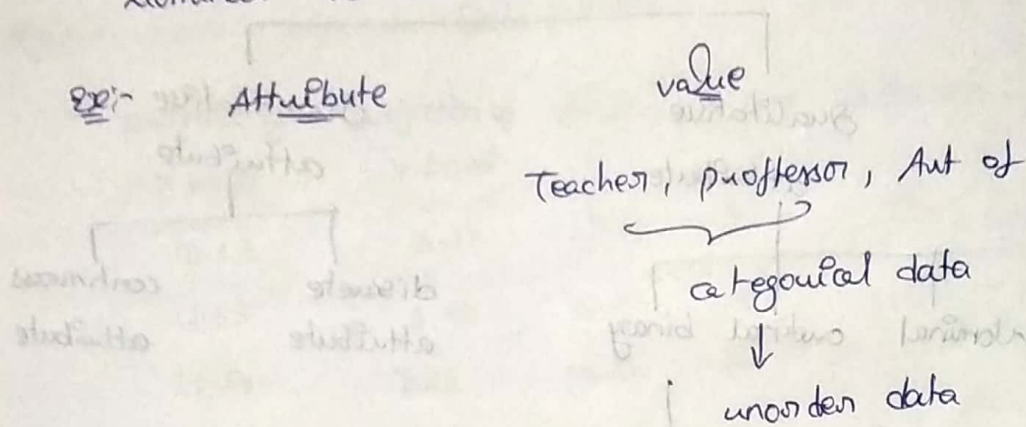
attribute values

height 4.5, 5.0

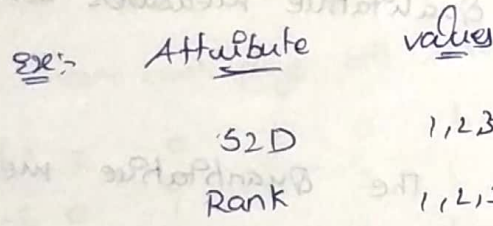
weight 50, 60

Ration b/w two values i.e., 4.5, 5.0 the mean, median and standard deviation are called b/w those ratios.

Nominal: The nominal is the unordered. Nominal attribute can take the values not in proper order. Nominal is also called categorical.



ordinal: The ordinal is nothing but a proper order (i) proper sequence.

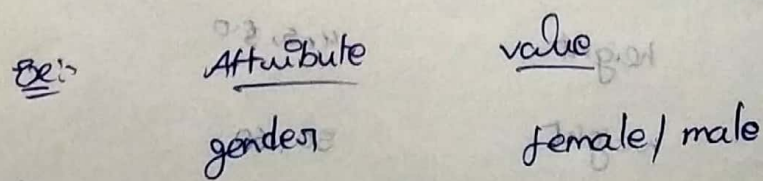


Binary: The binary attribute contains two values. These are true and false. The both values come under have same equal importance.

These two values are

- (i) symmetric
- (ii) asymmetric

symmetric: The symmetric values have both equal importance is called symmetric.



Asymmetric:- The both values not have equal importance & called Asymmetric.

ex:-

attribute	value
Report	T/F
cancel	Y/N

→ Types of dataset:- we have 3 types of data sets they are

- (i) Record data
- (ii) Graph based data
- (iii) ordered data

characteristics:- Generally to measured the data. The data belongs to which type otherwise the data belongs to which group whether it is record, grouped data (i) graph based grouped data (ii) ordered, grouped data. The dataset can have certain characteristics.

The characteristics are

- (i) Dimensionality
- (ii) sparsity
- (iii) Resolution

(i) Dimensionality:- The dimensionality of a data is the no. of attributes that the object is the dataset possesses.

(ii) sparsity:- Some data set have the asymmetric feature (asymmetric feature nothing but it can have the two values i.e. command the binary the attribute)

* most attributes of an object have values of non-zero because the non-zero values need to be stored and manipulated.

* The data mining algorithms work well only for sparse data.

(ii) Resolution: The data possibilities are obtained frequently at different level of resolution.

Ex:- (i) Earth seems very an event at resolution of few meters but is a relatively smooth at resolution of lots of kilometers.

(ii) The patterns in the data also depend on the level of resolution.

Record data (a) Record data

id	refund	marital status	tenable income	Defaulted Borrower
1	yes	single	125K	No
2	No	married	100K	No
3	No	single	70K	No
4	yes	married	120K	No
5	No	Divorced	95K	yes
6	No	married	20K	No
7	yes	Divorced	220K	No
8	No	single	85K	yes
9	No	married	75K	No

TID

items

1

Bread, soda, milk

2

Beer, Bread

3

Beer, soda, Diaper, milk

4

Beer, bread, Diaper, milk

5

soda, Diaper, milk

(b) Transaction data

projection of x load	projection of y load	Distance	load	thickness
10.23	5.27	15.22	27	1.2
12.65	6.25	16.22	22	1.1
13.54	7.23	17.34	23	1.2
14.27	8.43	18.45	25	0.9

(c) Data matrix

	Team	coach	play	ball	score	game	win	lost	time out	season
Doc-1	3	0	5	0	2	6	0	2	0	2
Doc-2	0	7	0	2	1	0	0	3	0	0
Doc-3	0	1	0	0	1	2	2	0	3	0

(d) document-term matrix

more of the data mining work assumes that the dataset is a collection of records (data objects). Each of which consists of fixed set of data fields (n) attribute.

The most basic form of record data shown in Fig (a). There is no explicit relationship among records (a) data fields. But every record has the same set of attributes.

Record data usually stored either in flat files (a) in relational DB's.

Transaction data (a) market - basket data:-

Transaction data is a special type of record data. where each record / transaction involves set of items. The set of products purchased by a customer with the store during one shopping trip constitute on transaction, this type of data is called market - basket data.

Transaction data is collection of set of items, that is viewed as set of records whose fields are asymmetric attributes.

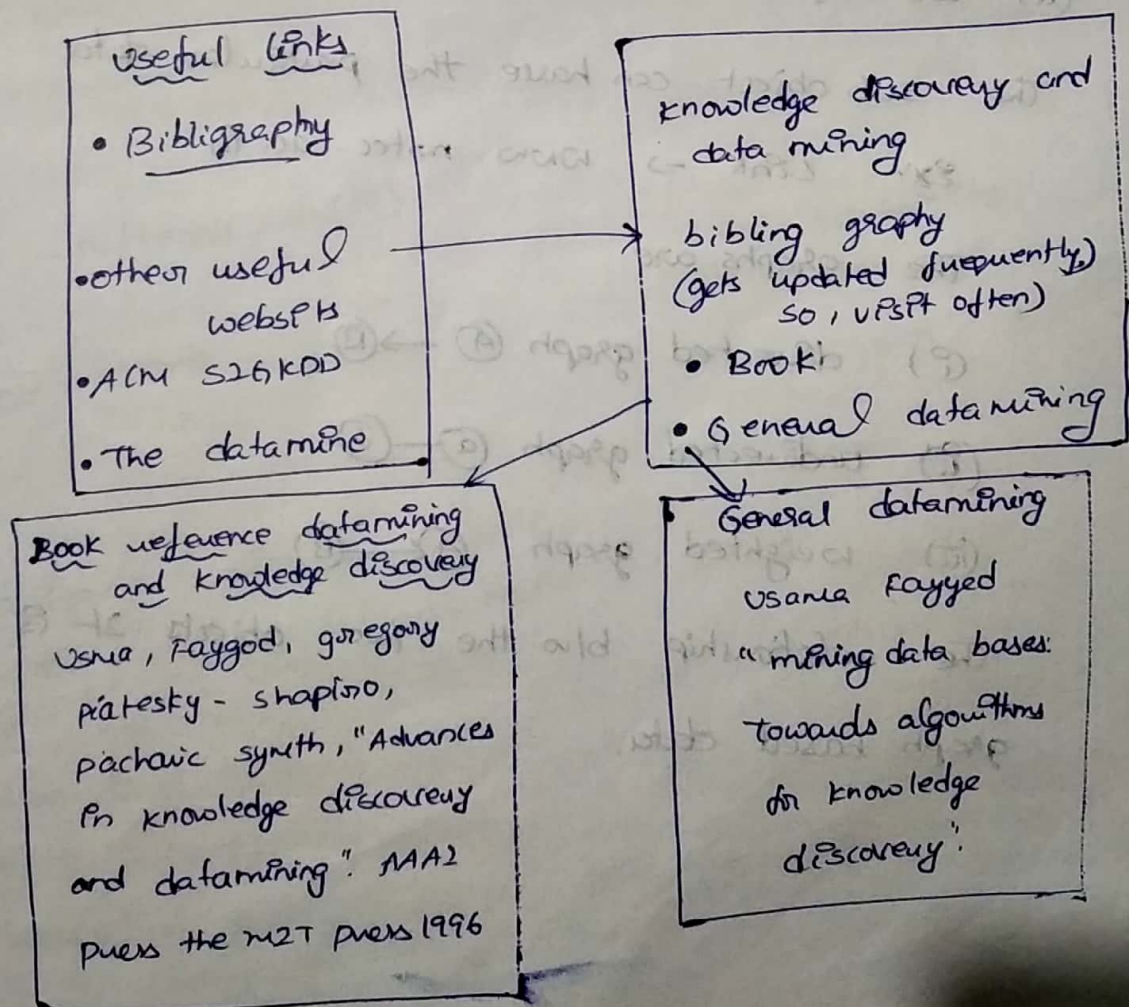
Data matrix:-

The data objects in a collection of data all have the same fixed set of numeric attribute the data objects can be thought of as points in multi-dimensional space.

A set of such data objects can be interpreted as $m \times n$ matrix is m rows and n columns. This matrix is called data matrix (a) pattern matrix.

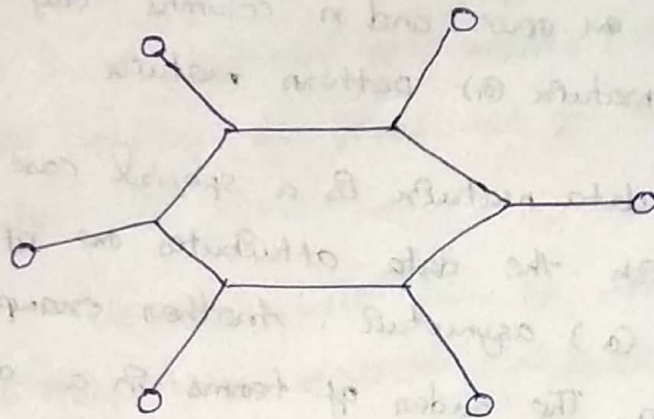
A sparse data matrix is a special case of data matrix. In which the data attributes one of the same type and (a) asymmetric. Another example for this is document data. The order of terms in a document is ignored, then the document represented as a term vector where each term is a component is the no. of times of the corresponding terms occurs in the document. This representation of a collection of a document is called document-term matrix.

Graph based data:-



Linking web pages

Benzene molecule:-



Graph based data:-

A graph is nothing but a collection of vertices & edges.
Here we have the 2 classifications in the graph based data.

- (i) The objects can have the relationships
- (ii) we object the sub object in one object

The sub object can have the particular data

Ex- Link $\rightarrow$ www.notecc.ac.in

The graphs are

- (i) directed graph $A \rightarrow B$
- (ii) undirected graph $A - B$
- (iii) weighted graph $A \xrightarrow{\quad} B$

The relationship b/w the two objects. It is called graph based data.

23-1-20

Unit - III

Datawarehouse & OLAP technologies

OLTP - online transaction processing

OLAP - online analytical processing

DW - subject oriented
heterogeneous

Time variant
separation

supply / customer
product
sales
view
model

OLAP / OLTD - users : Data content, Database, patterns
ER-model - OLTD - database designs

OLTP

ER-model
low level - users -

OLAP

star schema, snowflake
top level - users
knowledge workers,
managers,

SQL - patterns - simple short

tuple - records

view - model the
data of patterns

complex query - star net

tuple -

Data warehouse is also non-volatile
means the previous data is not
erased when new data is
entered in it. Data is used only
and periodically refreshed

throughput - at a time
more of works

transaction processing, reporting and concurrency control
mechanisms

Differences b/w operational / database systems & data ware houses.

→ comparision b/w OLTP & OLAP systems

<u>Feature</u>	<u>OLTP</u>	<u>OLAP</u>
characteristic	operational processing	informational processing
orientation	Transaction	Analysis
User	clerk, DBA, database professional	knowledge worker, ex managers, executive analyst
Function	day-to-day operations	long-term information requirements. decision support
DB design	ER based, application oriented	star / snowflake, subject-oriented
Data	current, guaranteed up-to-date	historical; accuracy maintenance overtime
summarization	primitive; highly detailed	summarized,
View	detailed, flat relational	Summarized, Multi-dimensional
Unit of work	short, simple transaction	Complex Query.
Access	read/write	mostly read
Focus	data in	information, out
Operations	Index / hash on primary key	lots of scan.
No. of records accessed	tens	millions.
No. of Users	thousands	Hundreds.
Metric	transaction throughput	query throughput, response time.
priority	High performance, high availability	High flexibility and user autonomy

multidimensional data model:-

Datware houses and OLAP tools are based on a multidimensional data model.

→ This multidimensional data model is used to view the data in the form of data cube.

→ Data cube is nothing but allow the data to be modeled and viewed in multiple dimensions. It is defined by dimensions and facts.

→ Dimension is nothing but an entity. Entity is any real world object. Here we take the entities with respect to an organization want to keep a record.

→ Each dimension may have a table associated with it called a dimension table. The dimension table further describes the dimension.

Ex: A table for item may contain the attributes item name, brand.

→ A multidimensional data model is typically organized around a central theme like sales.

→ This theme is represented by a fact table.

→ Facts are numerical measures i.e., the quantities by which we want to analyse relationships b/w dimensions.

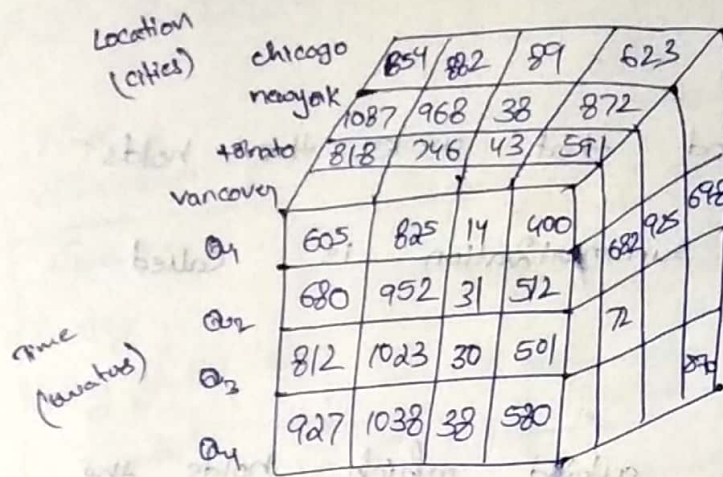
→ The fact table contains the names of the facts or measures as well as keeps to each of the related dimension table.

A 2D view of sales data for all electronics according to the dimensions time and item, where the sales are from branches located in the city of Vancouver, the measure displayed in dollars - sold (in thousands)

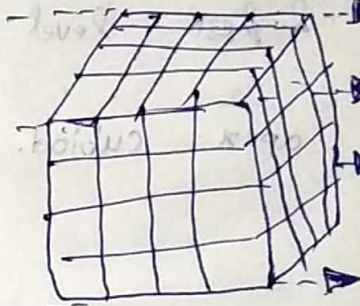
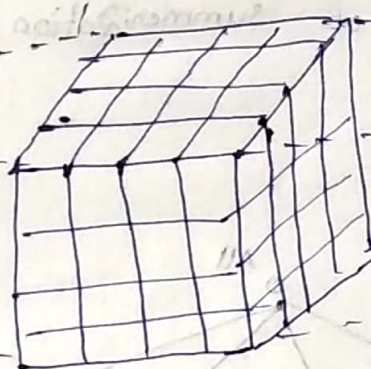
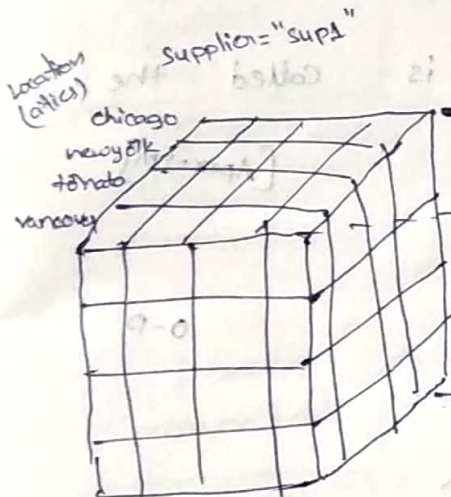
Location = "Vancouver"				
	item type			
time (quarter)	home entertainment	Computer	phone	security
Q ₁	605	825	14	400
Q ₂	680	952	31	512
Q ₃	812	1023	30	501
Q ₄	927	1038	38	580

A 3D view of sales data for all electronics, according to the dimension time, item and location. The measure displayed in dollars - sold (in thousands)

Location = "Chicago"	Location = "New York"	Location = "Toronto"	Location = "Vancouver"
item	item	item	item
home	home	home	home
Time out Comp phone sec	out comp phone sec	out Comp phone sec	out Comp phone sec
Q ₁ 854 882 89 623	1087 968 38 872	818 746 43 891	605 825 14 400
Q ₂ 943 890 64 698	1180 1024 41 925	894 769 52 682	680 952 31 512
Q ₃ 1034 934 59 789	1034 1048 45 1002	940 795 38 728	812 1023 30 501
Q ₄ 1109 992 63 870	1142 1091 54 982	978 864 59 784	927 1038 38 580



Computer/Security
home phone entertainment item(type)
supplier="sup1"
supplier="sup2"
supplier="sup3"



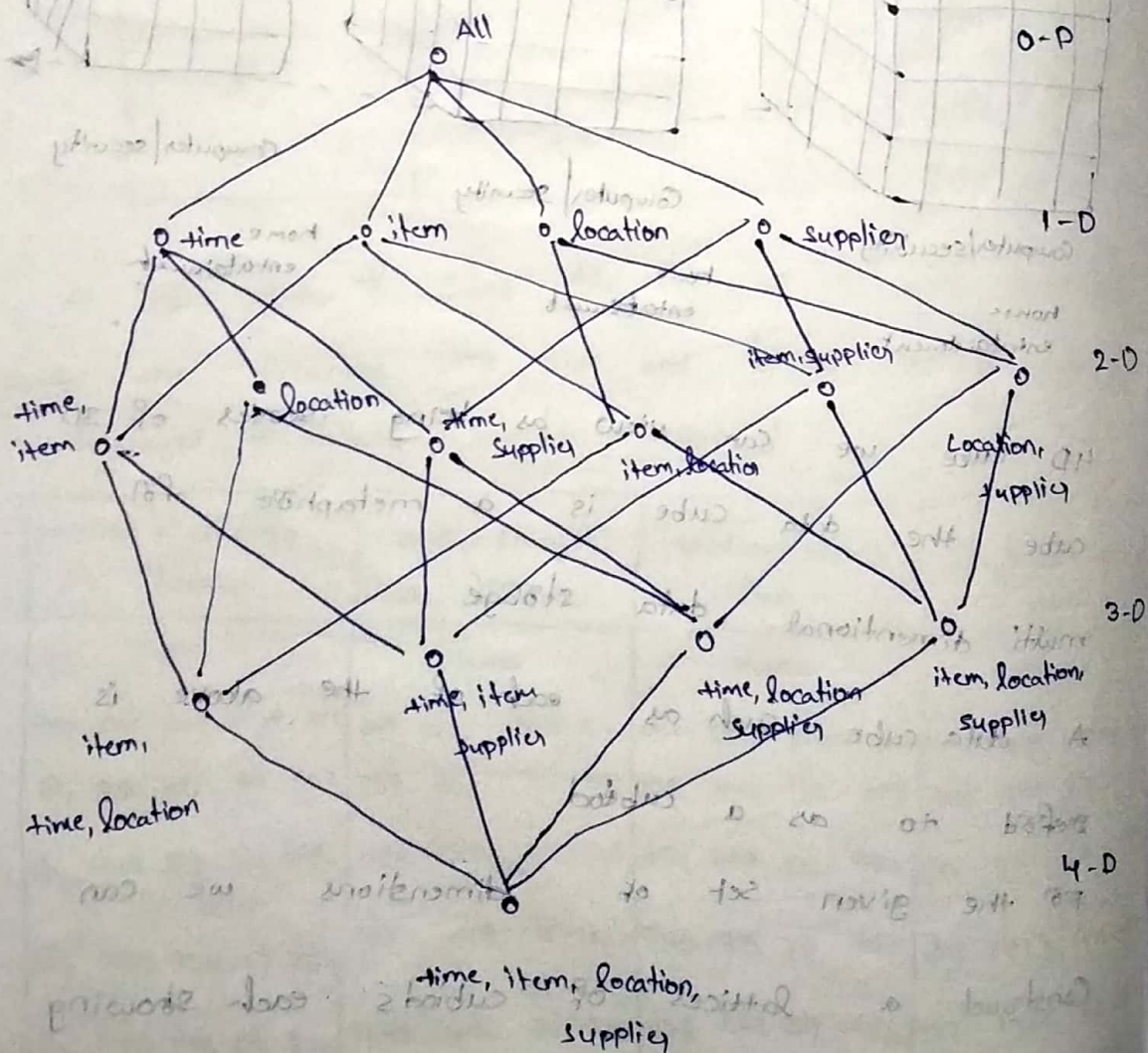
Computer/Security
home entertainment
4D cube we can view as being series of 3D cube
the data cube is a metaphor for multi dimensional data storage
A data cube such as each of the above is referred to as a cuboid
For the given set of dimensions we can

Construct a lattices of cuboids each showing

the data at a different level of group by

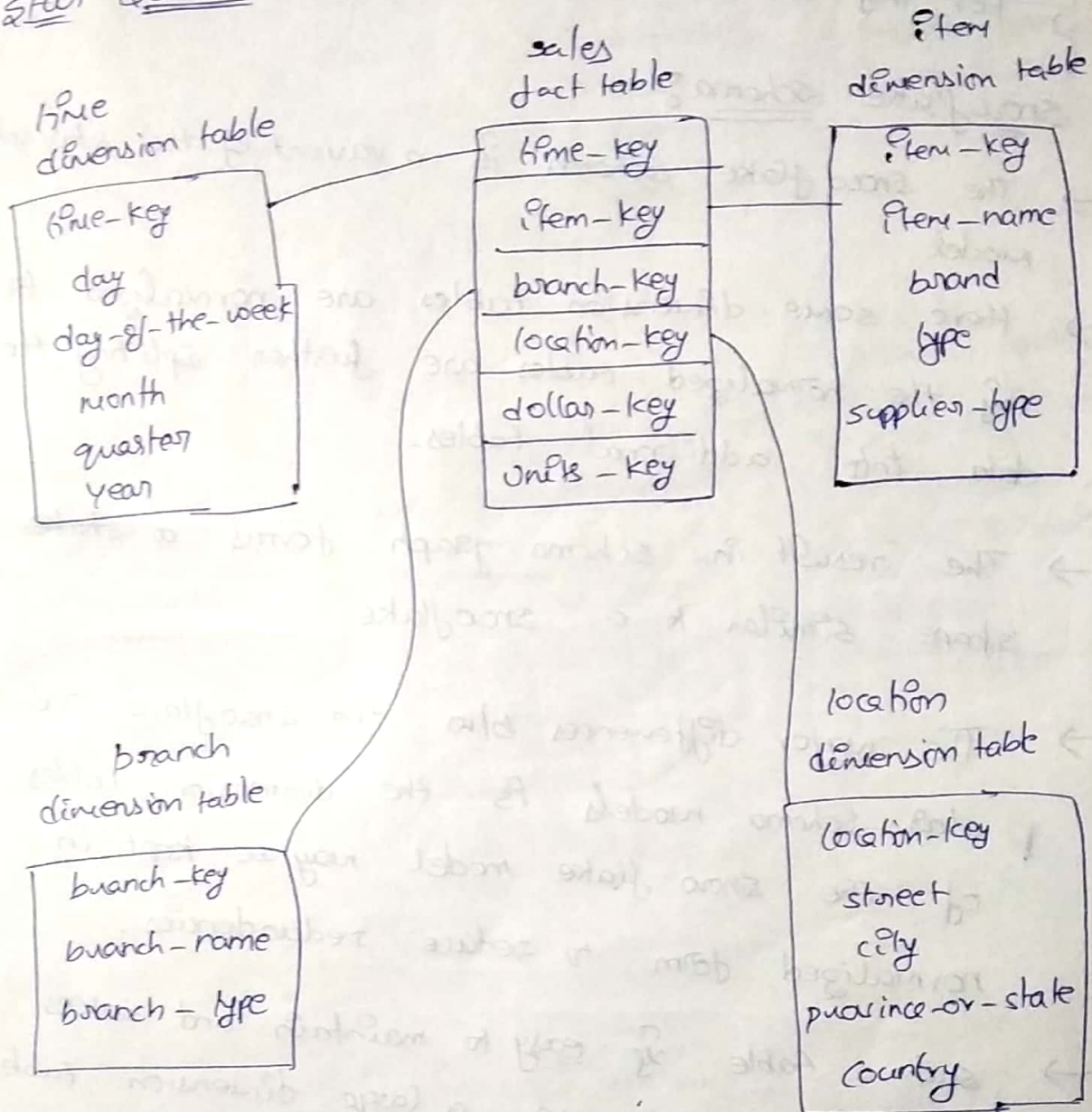
The cuboid that works the holds lowest level of summarization is called the base cuboid.

The O'D cuboid which holds the highest level of summarization is called the apex cuboid.



→ star, snowflake and fact constellation schemas for multidimensional databases

star schema:-



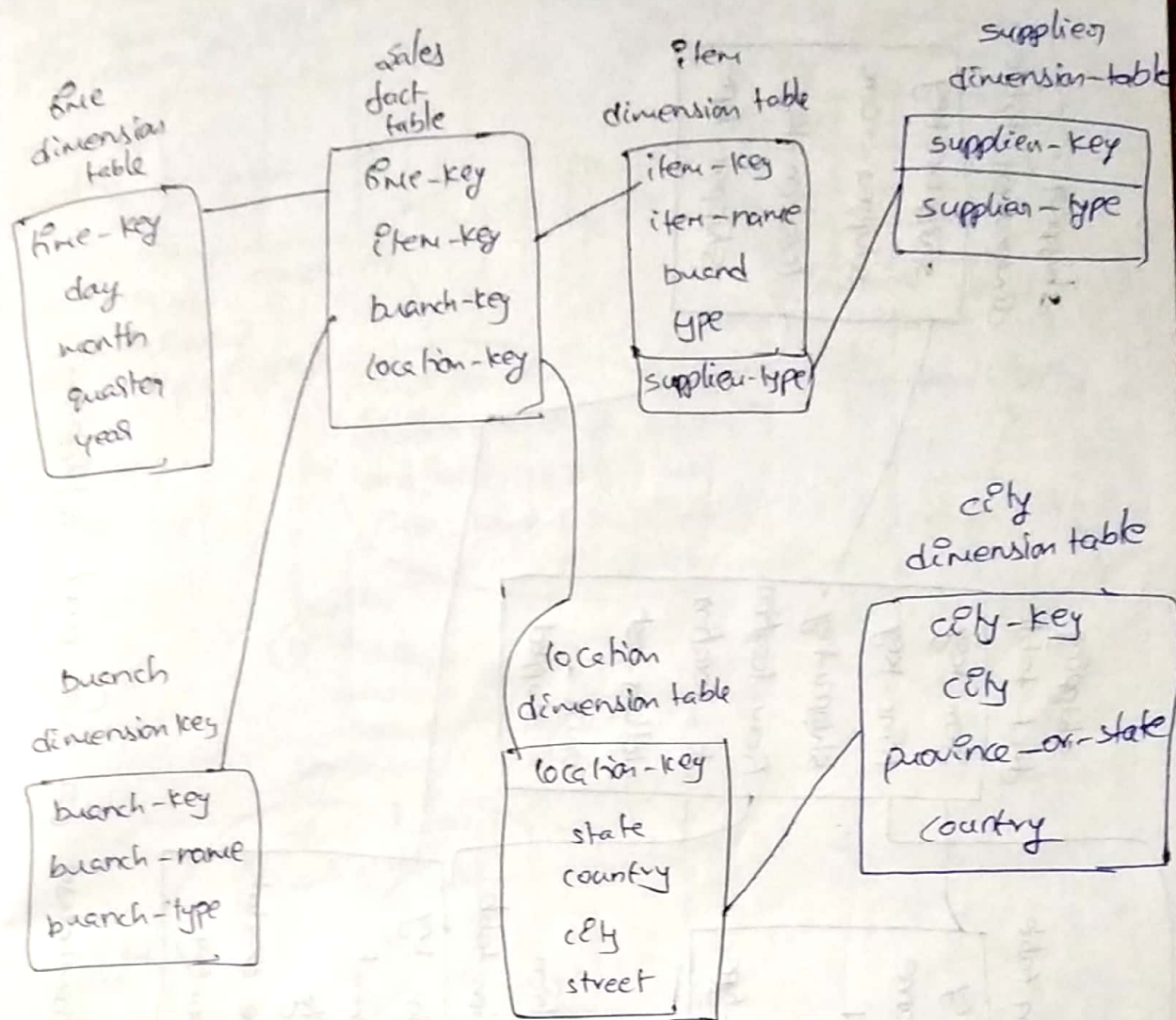
eg. star schema of a dataware house for sales.

→ the most common modeling paradigm is the star schema in which the dataware house contains a large central table (fact table) (sales) containing the bulk of the data with no redundancy and we have a set of smaller attendant tables (dimension tables), one for each dimension.

- The schema graph resembles (appears) a star burst with the dimension tables displayed in a radial pattern around the central fact table.
- For this reason this we call it as star schema.

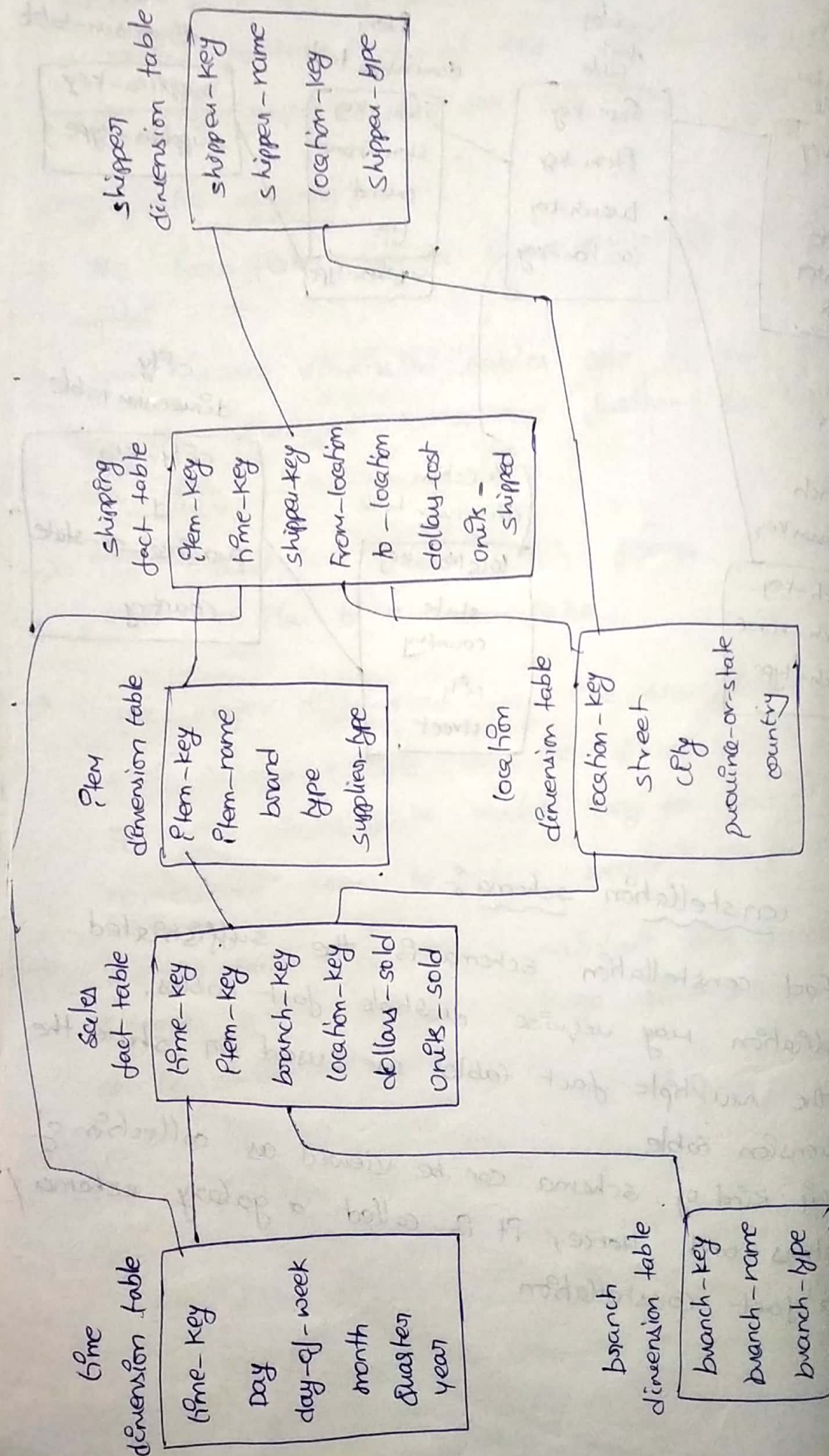
Snowflake schema:-

- The snowflake schema is a variant of the star schema model.
- Here some dimension tables are normalized for this the normalized tables are further splitting the data into additional tables.
- The result in schema graph forms a star shape similar to a snowflake.
- The major differences b/w the snowflake and star schema models is the dimension tables of the snowflake model may be kept in normalized form to reduce redundancies.
- such a table is easy to maintain and saves storage space because a large dimension table can be come anomalous.



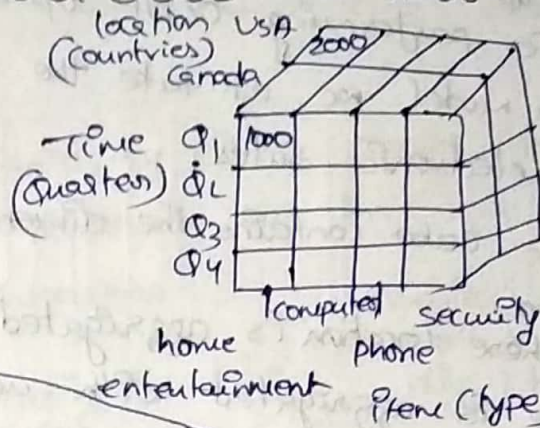
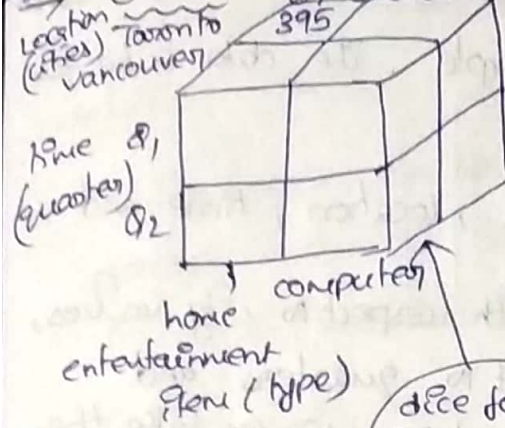
fact constellation schema :-

- Fact constellation schema is the sophisticated application may require multiple fact tables.
- the multiple fact tables are used for share the dimension table.
- This kind of schema can be viewed as collection of stars and hence, it is called a galaxy schema / a fact constellation.



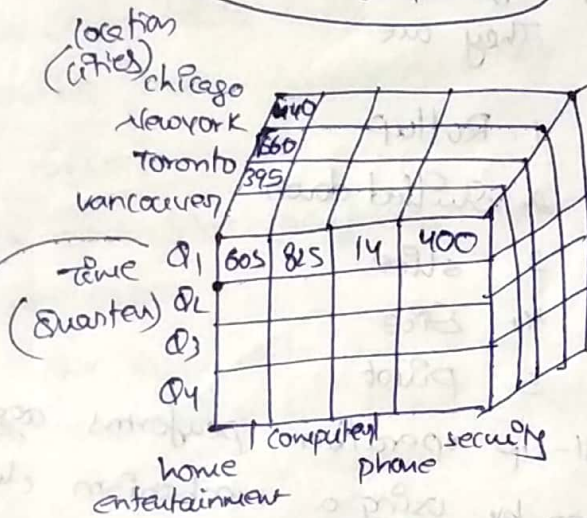
Fact constellation schema of a data warehouse for sales and shipping

OLAP operations in multidimensional data model:-

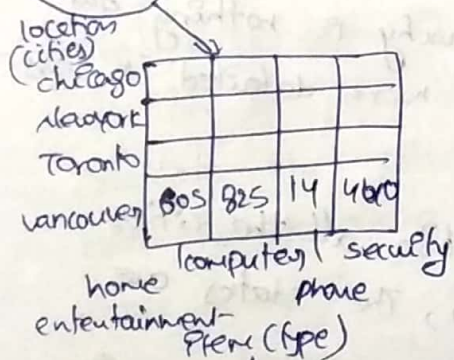


slice for
(location = Vancouver, Toronto)
and (time = Q1 and Q2) and
item = home entertainment/
computer

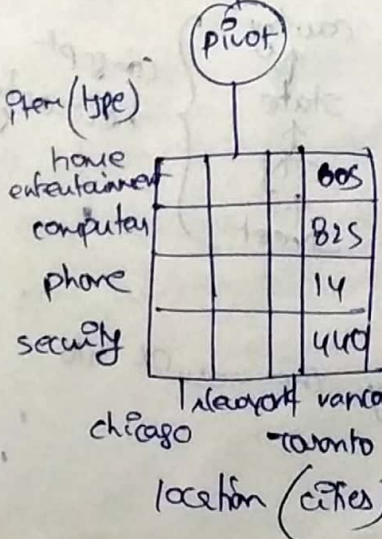
Rollup
on location
from cities
to countries



slice
for
time = Q1



rolled down
on time
from quarters
to month



Rollup:

→ For performing OLAP operations in the multidimensional data model, we can take the example. The data cube for all electronic sales.

→ The cube contains the dimensions, location, time and item.

→ where location is aggregated with respect to city values, time is aggregated with respect to quarters and item is aggregated w.r.t. to item types. We can take the all electronic sales data cube treated as the central cube.

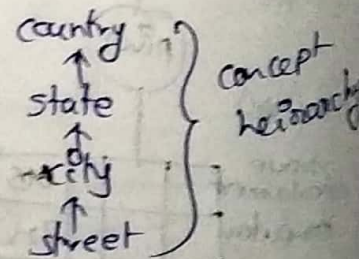
→ We have five types of OLAP operations performed on the central data cube. They are

1. Rollup
2. Drilled down
3. slice
4. Dice
5. pivot

1. Rollup:- The Roll-up operation performs aggregation on data cube either by using a mechanism climbing up a concept hierarchy (concept hierarchy is nothing but we can take the data in the form of more detailed to less detailed).

For example; with the dimension location all the cities combined together form a state then, the states are combined together form a country.

→ The Rollup operation shown aggregates the data by ascending the location hierarchy from the level of city to the level of country.



→ Rather than grouping the data by city, the result into cube groups the data by country.

2. Drilled down: Drilled down is the reverse of roll-up

→ It uses the mechanism (less detailed data - more detailed data)

→ Drilled down can be replaced by either stepping-down concept hierarchy for a dimension or here we introduce additional dimensions

→ Drilled down uses descending the time hierarchy from the level of quarter to the more detailed level of month.

→ The Resulting data cube details the total sales per month rather than summarized by quarter.

3. slice & dice:

→ The slice operation performs selection on one dimension of the given cube.

Example the sales data are selected from the central cube for one dimension "time" using the criterion time = "Q₁" is for the slice.

→ The dice operation defines sub-cube by performing selection on two or more dimensions

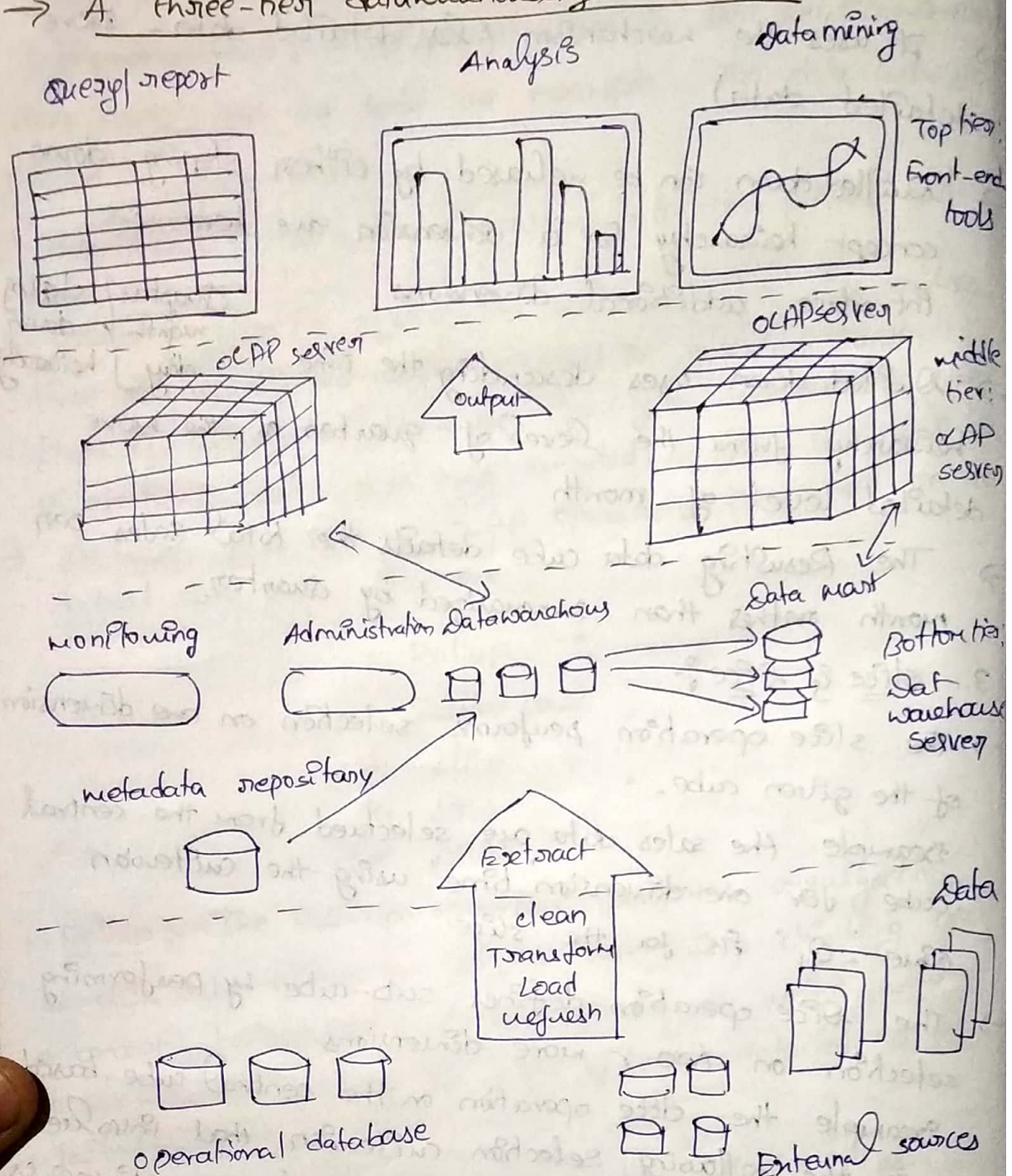
Example the dice operation on the central cube based upon the following selection criterion that involve three dimension Location = Toronto or Vancouver Product =

item = home entertainment / computer.

4. pivot (rotate):

→ Pivot is also called a rotate is a visualization operation that rotates the data access in view in order to provide an alternative representation of the data.

→ A. three-tier data warehousing architecture:-



→ steps for design and construction of data warehouse.

- a) To design an effective data warehouse 1st we need to understand business needs, and then construct the data warehouse belongs that need.
- b) we have 4 different views regarding the design of data warehouse.

1. Top-down view :-

- selection of the most relevant information necessary for the datawarehouse
- This information matches current and up-coming business need

2. Data-source view :-

- The information being captured stored and managed is also well documented at various levels of detailed and accuracy from individual data source tables to integrated data source table.

3. Data-warehouse view :-

- It includes fact and dimension tables the information is stored in the form these fact and dimension tables

4. Business-Query View :-

- the data in the data warehouse from the view point of the end user it can be obtained
- To perform all these views i.e., we design and construct data warehouse 2nd of all choose a business process to model (waterfall model) choose the grain of the business process. The grain is nothing but the data to be represented in the fact table for this process.
- choose the dimensions. The dimensions are applied to each fact table record.

Three-tier data warehousing architecture :-

Bottom-tier :-

- It is the warehouse database server the data from the data source can take prior it can put into warehouse it can be extracted or cleaned using application program

Interface known as gateway.

→ A gateway is supported by DBMS and allows client programs to generate sql code to be executed at a server.

Example. ODBC

Middle-tier:-

→ It is an OLAP server is typically implemented with a ROLAP (Relational OLAP model) is an extended relational dbms that maps the operations on multi-dimensional data to standard relational operations.

→ It's also implemented with MOLAP (multidimensional OLAP). This model is a special purpose server that directly implements multi-dimensional data and operations.

Top-tier:-

→ Top tier is nothing but a client here we use front-end tools i.e., which contains Query, reports and performs the analysis of data with the help of determining tools used.

→ Tg.

→ Types of OLAP servers :-

we have 4 types of OLAP servers they are

1. ROLAP
2. MOLAP
3. HOLAP
4. specialized SQL server

1. ROLAP :- (Relational online analytical processing)

→ It stands for ROLAP server. These servers are the intermediate servers.

→ why we call it as intermediate server? It stands in b/w relational backend server and client front-end tools.

→ They use a relational or extended relational dbms to store and manage the warehouse data.

→ ROLAP servers include optimization for each dbms back-end, implementation of aggregation and navigation logic and also provide additional tools and services.

→ The ROLAP technology have greater flexibility than MOLAP technology.

2. MOLAP :- (multidimensional OLAP)

→ It stands for multidimensional OLAP server.

→ These servers support multidimensional views of data through array based multidimensional storage in engine.

→ They map multidimensional views directly to data cube array's structure.

→ With multidimensional data stores the storage utilization may be low if the data set is sparse.

→ many MOLAP servers adopt a tool level storage representation to handle sparse and dense dataset.

3. HOLAP :- (Hybrid OLAP)

→ 21- stands hybrid OLAP servers.

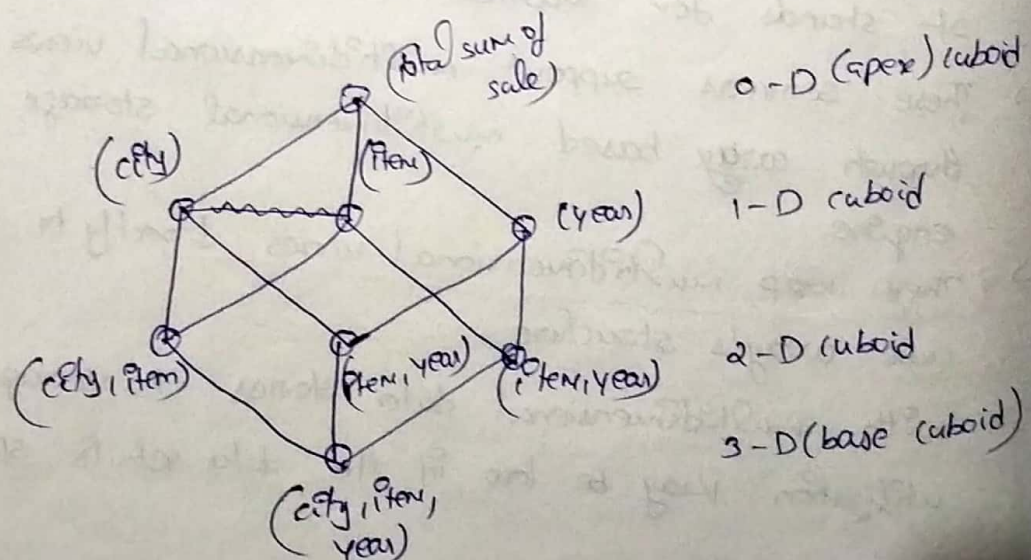
→ The hybrid OLAP approach combines the both the ROLAP & MOLAP technologies, benefiting the greater flexibility of ROLAP and the faster computations of MOLAP.

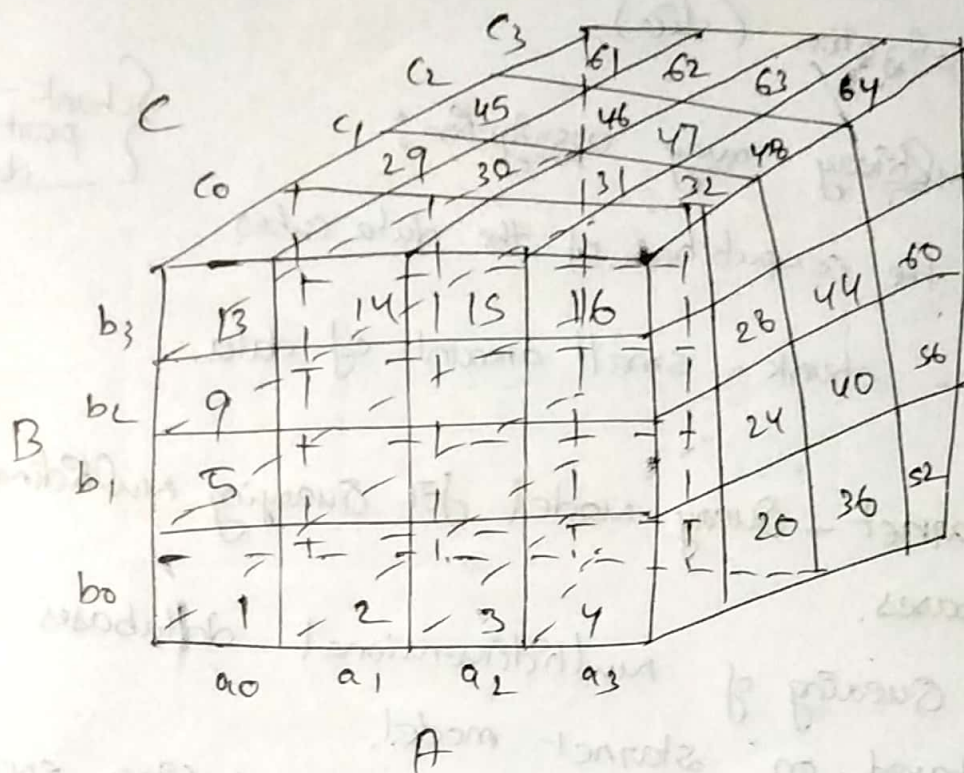
4. specialized SQL servers :-

→ For the growing demand of OLAP processing in relational-databases.

→ some relational and data warehousing dising (ex:- Red bull of informis) implement specialized SQL servers that provide advanced query language and query processing support for SQL service over star & snowflake schemas in a need only environment.

Data warehouse implementations :- (a) operations / view of data warehouse and OLAP technology for data mining





→ star net - Query model for querying multidimensional databases.

→ The data warehouse implementation can be happened with the help of efficient computations of data cubes

→ The example for this is the above lattice from this we compute the some of sales grouping by item & city - 2D, compute the some of sales grouping by only item 1-D, compute the some of sales grouping by city, item, year - 3-D.

→ partial materialization - From this selected computations of cuboids.

→ compute all of the cuboids is called as full materialization (dpc).

→ compute any of the non-base cuboid we call as no materialization (apex)

→ selectively compute proper subset of the whole set of possible cuboids is called as partial materialization (slice).

→ multistage away aggregation:-

{ chunk -
part of
data

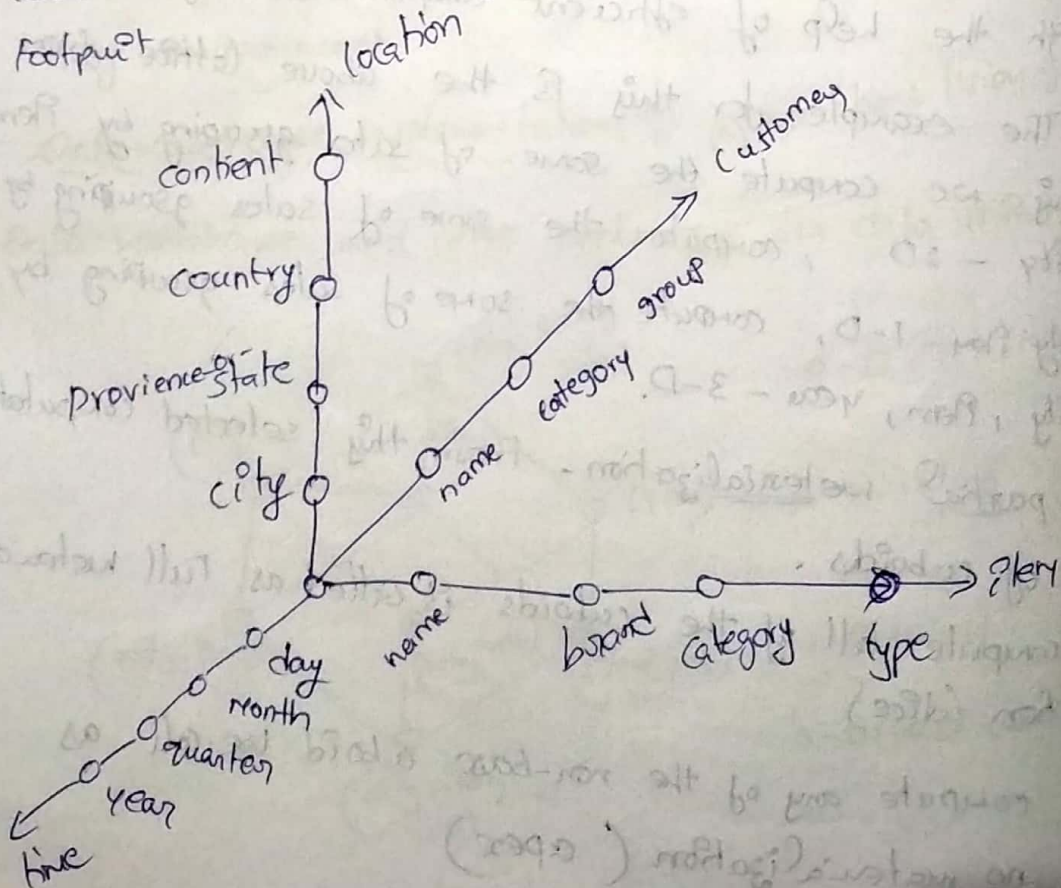
→ In the computation of the data cubes

note:- chunk - small amount of data.

→ star - Query model for querying multidimensional databases.

→ The Querying of multidimensional databases can be based on star model.

→ A star model consists of radial lines emanating from a central point, where each line represents a concept hierarchy for a dimension. each abstraction level in the hierarchy is called a footprint.



1.C

Unit - IV classification

Def: classification is nothing but process of mapping on attributeset 'x' to the 'y' its classification following nature

General approach

TTP	attribute 1	attribute 2	attribute 3	class
1	yes	large	125K	no
2	no	medium	100K	no
3	no	small	70K	no
4	yes	medium	125K	no
5	no	large	95K	yes
6	no	medium	60K	no
7	yes	large	220K	no
8	no	small	85K	yes
9	no	medium	75K	no
10	no	small	90K	yes
11	no	small	55K	?
12	yes	medium	80K	?
13	yes	large	110K	?
14	no	small	95K	?
15	no	large	67K	?

learning algorithm

Indecision

learn model

model

deduction
Apply model

confusion binary

		predicted class	
		class:1	class:0
Actual class	class:1 yes	+11	+10
	class:0 no	+01	+00

evaluation of the performances of the classification model is based on the counts of test records in

correctly predicted by the model the count or tabulated in table is known as confusion matrix

confusion matrix for a binary classification problem

$$\text{Accuracy} = \frac{\text{no. of correct production}}{\text{Total no. of production}}$$

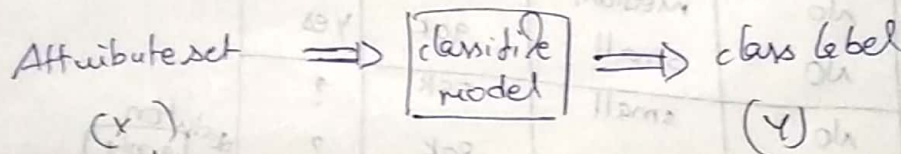
$$= \frac{t(11) + f(00)}{t(10) + t(01) + t(10) + t(11)}$$

Descriptive - induction (training)

predictive - deduction

$$\text{Error rate} = \frac{\text{no. of wrong production}}{\text{Total no. of production}}$$

$$= \frac{t(01) + t(10)}{t(01) + t(00) + t(10) + t(11)}$$



classification def: classification is a task of learning target function f that maps each attribute set x to one of the predefined class labels y .

$\rightarrow$ The target function is also known informally as a classification model.

$\rightarrow$ The input data for classification record.

$\rightarrow$ Each record also known in instances is characterise by a tuple (x, y) where x is attribute set Y

specialized attribute designated as the class label
also known as target attribute.

→ Description is nothing but training (a) induction,
training - having class labels.
induction - "

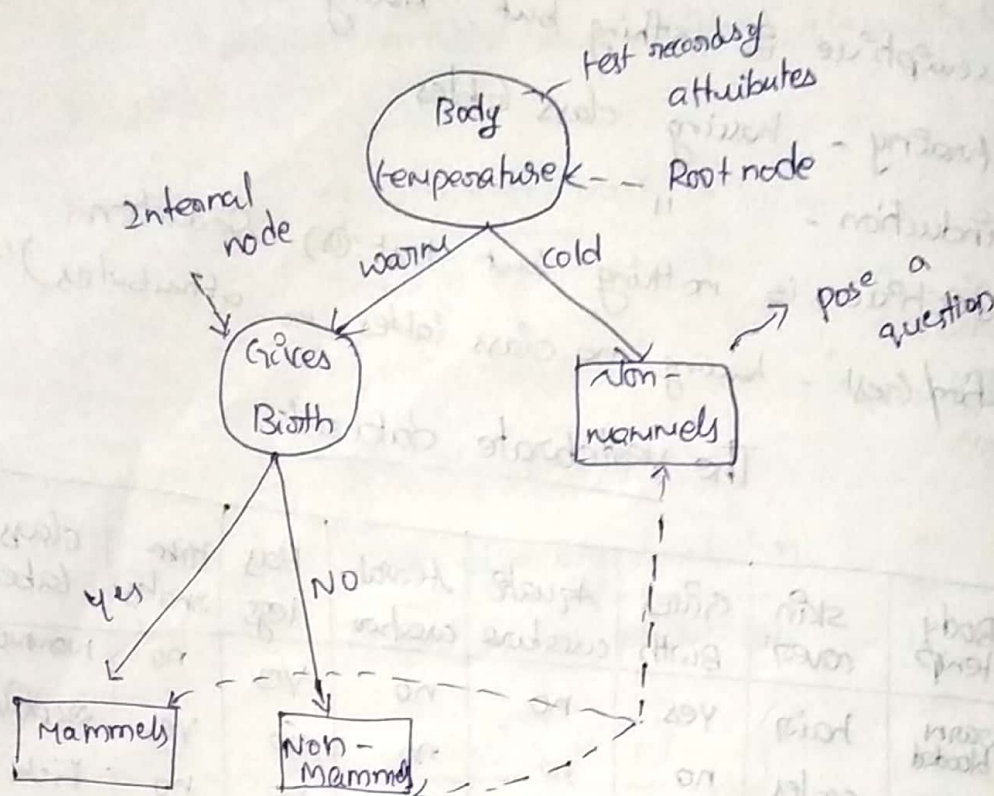
→ prediction is nothing but test (b) deduction
deduction/test - having no class labels, no attributes) / c

The vertebrate data set

C 2.9

name	Body temp	skin cover	Lives Birth	Aquatic creature	Aerial creature	Has legs	like snakes	class label
Human	warm blooded	hair	yes	no	no	yes	no	mammal
python	cold	scales	no	no	no	no	yes	reptile
salmon	cold	"	no	yes	no	no	no	fish
whale	warm	hair	yes	yes	no	yes	yes	mammal
Frog	cold	none	no	semi	no	yes	no	amphibian
comodo dragon	cold	scales	no	no	no	yes	yes	reptile
bat	warm	hair	yes	no	yes	yes	yes	mammal
pigeon	warm	feathers	no	no	yes	yes	no	bird
cat	warm	fur	yes	no	no	yes	no	mammal
eel	warm	fish	yes	yes	no	no	no	fish
eopard	cold	scales	yes	yes	no	no	no	reptile
shark	cold	scales	no	semi	no	yes	no	fish
turtle	cold	scales	no	semi	no	yes	no	reptile
bird	warm	feathers	no	semi	no	yes	no	bird
penguin	warm	quills	yes	no	no	yes	yes	mammal
caracine	warm	quills	yes	no	no	yes	yes	mammal
eel	warm	quills	yes	no	no	yes	yes	mammal
eel	cold	scales	no	yes	no	no	no	fish
salamander	cold	none	no	semi	no	yes	yes	amphibian

→ Best fit relationship b/w the attribute set and class labels of the P/P data.

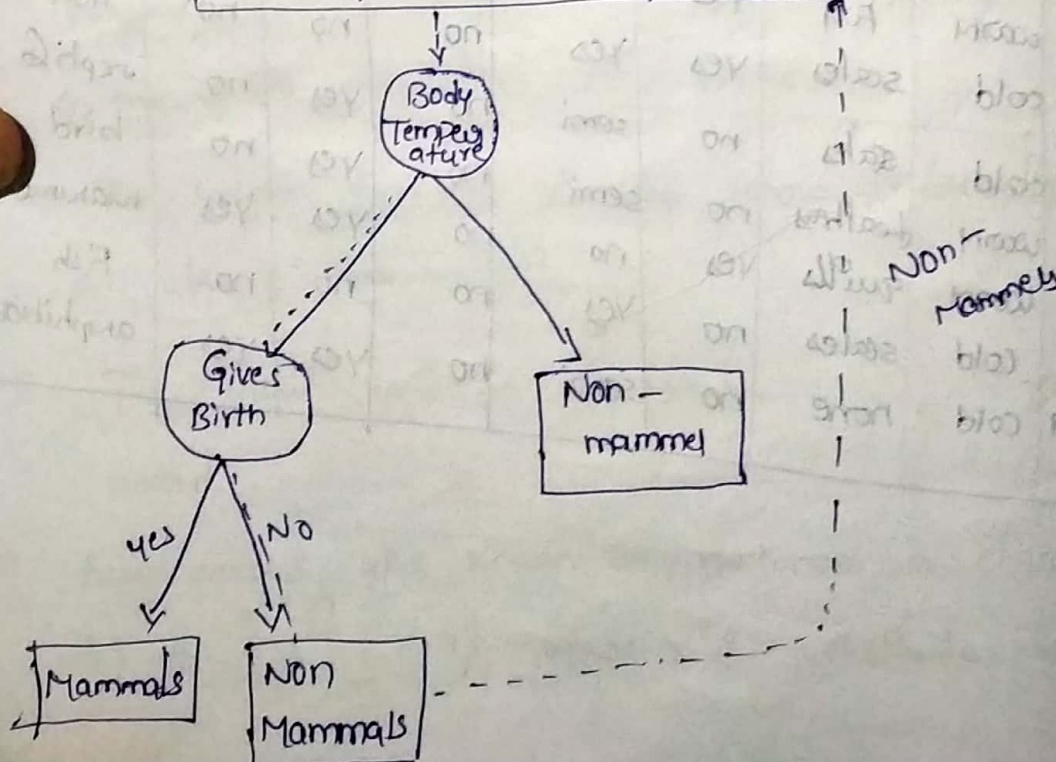


A decision tree for the Mammal classification problem.

How to build a decision tree

unabeled data

Name	Body Temperature	Gives Birth	Class
Flamingo	Warm	No	Non-Mammals



classifying an unlabeled vertebrate. The dashed lines represent the outcomes of applying various attributes test conditions on the unlabeled vertebrate. The Vertebrate is eventually assigned to the non-mammal class.

(2-b)
Hunt's Algorithm:- How to build a decision tree

With nearly we have the big dataset of construct more decision tree while some of the trees are more accurate than others. Finding the optimal tree is computationally infeasible. For this view the efficient algorithms of the exponential size of search space developed to include a reasonable accurate.

This algorithm usually employed greedy states that grows decision tree by making a series of locally optimum decisions about which attribute to use for partition the data. once such algorithm is hunt's algorithm. which is the basis of many existing decision tree induction algorithm.

inducing 2c-3

Hunt's algorithm:-

In hunt's algorithm decision tree is grown in recursive by partition by partitioning the training records in to successfully pure subset
→ Let D be the set of training records that are

associated with node 'T' is

$Y = \{Y_1, Y_2, \dots, Y_c\}$ be the class labels

step 1:- If all the records in D_T belong to the same class Y_c then 'T' is a leaf node labelled as Y_c

step 2:- If D_T contains records that belong to more than one class an attribute test condition is selected to partitioning the record into small subsets a child node is created for each outcome of the tests condition and record are distributed in the children based outcomes

step 3:- The algorithm is recursively applied to each child node.) 2.b

→ 3.a methods for expressing attribute test condition

Decision Induction algorithm must provide different methods for expressing an attribute test condition and its corresponding outcomes for different attribute types.

We have different attribute types. They are

1. Binary Attribute type

2. Nominal Attribute type

3. Ordinal Attribute type

4. Continuous attribute type

→ 1. Binary Attribute type:-

The test condition for a binary attribute type generates two potential outcomes.

2. Nominal Attribute type:-

The nominal attribute can have many values. Its test condition can be expressed in two ways.

- 1. multiway split
- 2. Binary split

3. Ordinal Attribute type:-

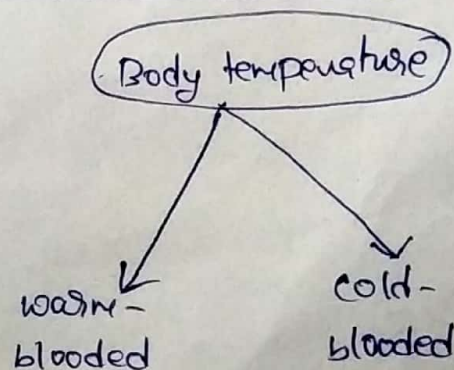
Ordinal Attributes can also produce binary or multiway splits.

Ordinal Attributes values can be grouped as long as the grouping does not violate the order property of the attribute values.

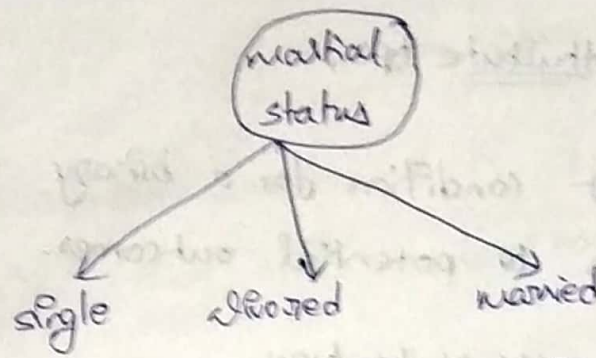
4. Continuous Attribute type:-

For a continuous attribute the test condition can be expressed as a comparison test ($<$, $>$) with binary outcomes or multiway outcomes.

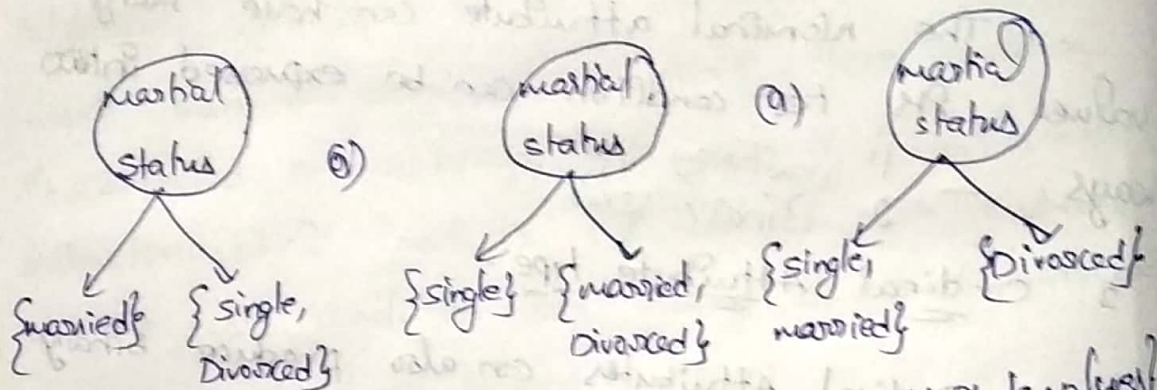
Test condition for binary attributes



Test conditions for nominal attributes

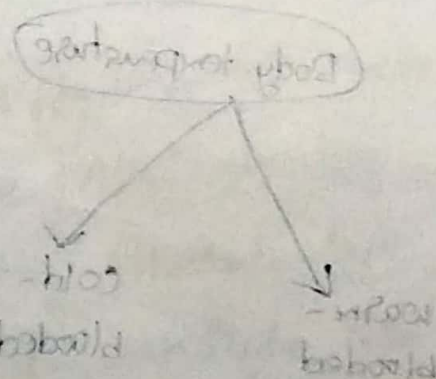
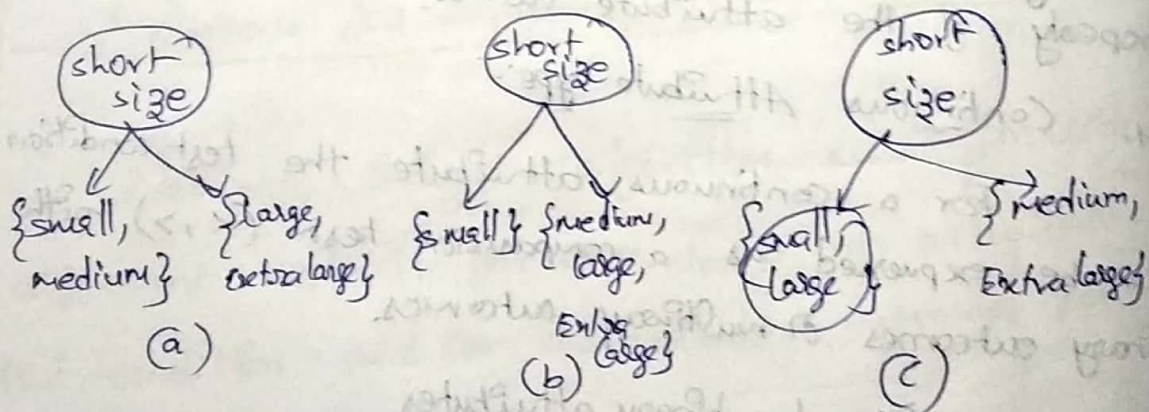


(a) multikway split

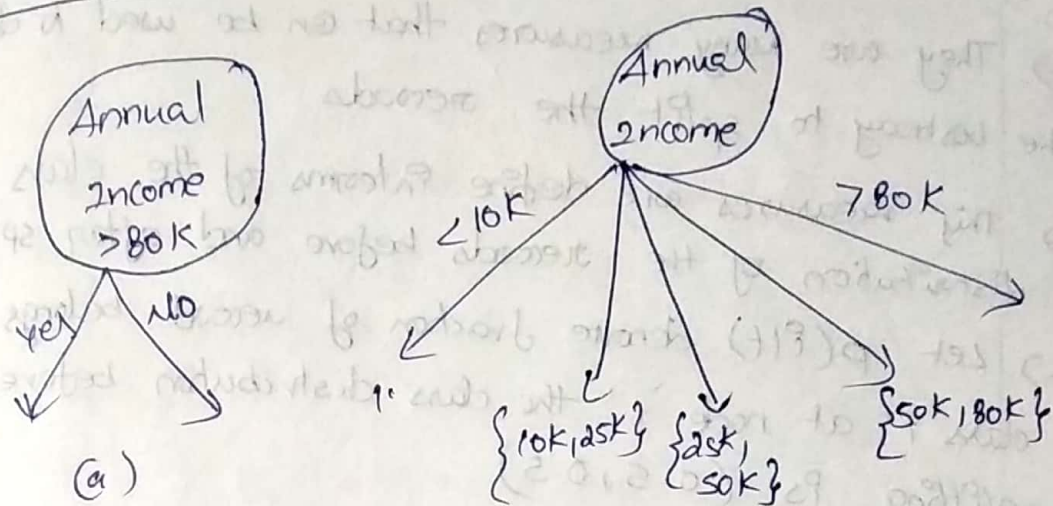


(b) Binary split - by grouping attribute values

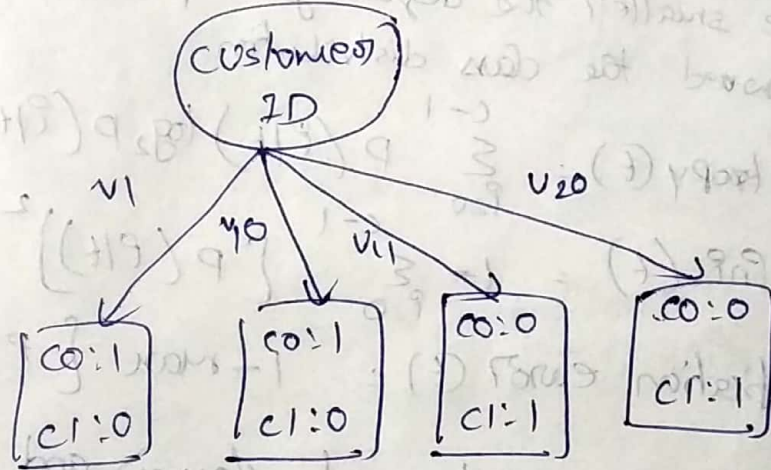
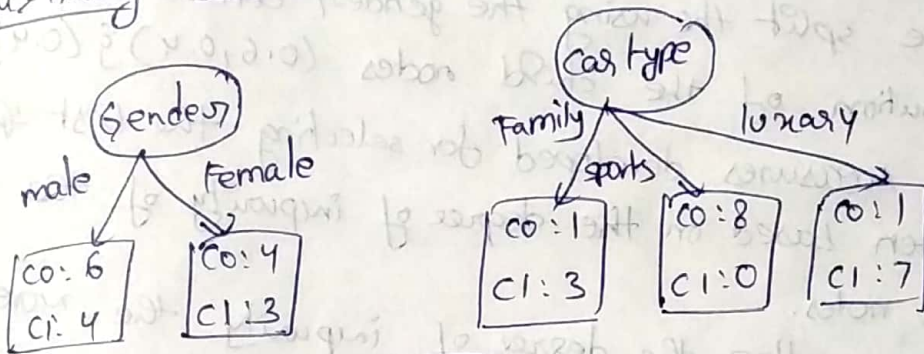
Different ways of grouping Ordinal attribute values



Test condition for continuous attributes



multinomial versus binary splits



3) Measures for selecting the best split

→ They are many measures that can be used to determine the best way to split the records

→ These measures are defined in terms of the class distribution of the records before and after splitting

→ Let $P(i|t)$ denote fraction of record belongs to class 'i' at node 't'. The class distribution before splitting is $(0.5, 0.5)$

They are an equal number of records from each class.

→ If we split using the gender attribute, the class distribution of the child nodes is $(0.6, 0.4)$ & $(0.4, 0.6)$

→ The measures developed for selecting the best split are often based on the degree of impurity of the child nodes.

→ The smaller the degree of impurity, the more skewed the class distribution.

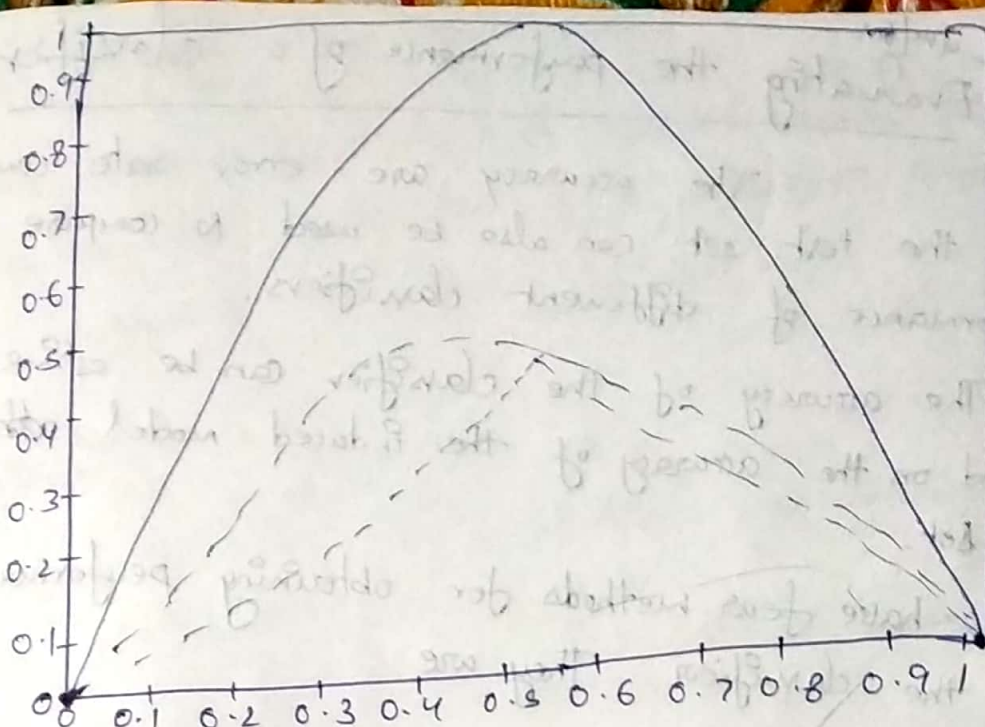
$$\text{Entropy}(t) = - \sum_{i=1}^{c-1} P(i|t) \log_2 P(i|t),$$

$$\text{Gini}(t) = 1 - \sum_{i=1}^{c-1} [P(i|t)]^2$$

$$\text{classification error}(t) = 1 - \max [P(i|t)]$$

where c is number of classes and

$0 \log_2 0 = 0$ in entropy calculations.



Node N1	count
class = 0	0
class = 1	6

$$Gini = 1 - (0/6) - (6/6)^2 = 0$$

$$Entropy = -(0/6) \log_2(0/6) - (6/6) \log_2(6/6) = 0$$

$$Error = 1 - \max[0/6, 6/6] = 0$$

Node N2	count
class = 0	1
class = 1	5

$$Gini = 1 - (1/6)^2 - (5/6)^2 = \frac{5}{6}$$

To determine how well a test condition performs we need to compare the degree of purity of the parent node.

3.6

Evaluating the performance of a classifier

u.a

The accuracy and error rate computed from the test set can also be used to compare the performance of different classifiers.

→ The accuracy of the classifier can be estimated based on the accuracy of the induced model on the test set.

→ We have four methods for obtaining performance of the classifier. They are

1. Holdout method
2. Random subsampling method
3. Cross-validation
4. Bootstrap

1. Holdout method:-

→ In the holdout method the original data with labeled examples is partitioned into two disjoint sets called the training & test sets.

→ A classification model is then induced from the training set and its performance is evaluated on the test set.

Limitations

→ The Holdout method has several limitations

1. Whether the training labeled examples are fewer and some of the further records are withheld for testing.

2. Random subsampling method

- The Holdout method can be repeated several times to improve the estimation of a classifier performance
- This approach is known as R.S.M

Limitations

Has no control over the no. of times each record is used for testing and training

3. Cross-validation

- An alternative to random subsampling method is the cross-validation.

→ In this method each record is used the same number of times for training and exactly one's for testing

→ In this method we partitioned the data into two equal sized subsets from they first we choose one of the subset for training and the other for testing.

→ We then swap the roles of subsets then the previous training set becomes the test set and the test set becomes training set. This approach is called a two fold cross validation.

4. Bootstrap Limitation

Here we have the chances of duplication of the records are possible.

4. Bootstrap:

→ In this method we remove ~~the~~ duplicate records in the training and test sets.

→ In the bootstrap approach the training records are sampled without replacement i.e. A record already

chosen for training is put back into the original pool of the records, so that it is equally likely to be selected again.

→ Bayes' theorem - Bayesian classification

Dataset

car-no	color	Type	origin	stolen
1	Red	sports	DOM	Yes
2	Red	"	"	No
3	Red	"	"	Yes
4	Yellow	"	2MP	No
5	"	"	"	Yes
6	"	SUV	"	No
7	"	"	"	Yes
8	"	"	"	No
9	Red	"	"	No
10	"	"	"	Yes

Previously we have the decision tree with the decision tree for the labeled records of a particular entry contained in the dataset is easily we can make a decision.

- Basically for every classification model we have the two classes only either Yes/No, true/false.
- Here we can't make a decision for the possibility of unlabeled entries of a record.
- To overcome this we can go and use Bayesian classification.
- With the Bayesian classification we can make a decision for the unlabeled entries also with the help of probabilities.

Probability

probability is nothing but getting the chance of outcome of possible events.

→ we have two types of probability. They are

1. Normal probability
2. Conditional probability

1. Normal probability - We have two boxes, named A, B. what is probability for selecting box A,

$$P(A) = \frac{1}{2}$$

probability for selecting box B is $P(B) = \frac{1}{2}$

ex-2 We have - 3 boxes named as A, B, C. Probability of getting box A = $P(A) = \frac{1}{3}$

probability of getting box B $P(B) = \frac{1}{3}$

probability of getting box C $P(C) = \frac{1}{3}$

2. Conditional probability - It is nothing but, Based upon the condition we can take the probability.

$$P(X|A) = R/A$$

We have box A it can have 5 Red balls, 6 black balls, 4 white balls. Here we can select the Red balls from the box A.

$$\frac{5}{15}$$

Bayes formula:

$$P(A|X) = \frac{P(X|A) P(A)}{P(X|A) P(A) + P(X|B) P(B)}$$

Normal probability:-

$P(\text{yes})$	$= 5$
$P(\text{no})$	$= 5$

} color
type
origin

conditional probability

color: only Red & yellow

$P(\text{Red} \text{yes}) = 3/5$	$P(\text{Red} \text{no}) = 2/5$
$P(\text{yellow} \text{yes}) = 2/5$	$P(\text{yellow} \text{no}) = 3/5$

type: suv & sports

$P(\text{suv} \text{yes}) = 3/5$	$P(\text{suv} \text{no}) = 2/5$
$P(\text{sports} \text{yes}) = 2/5$	$P(\text{sports} \text{no}) = 3/5$

origin: Domestic, imported

$P(\text{Dom} \text{yes}) = \frac{2}{3}$	$P(\text{Dom} \text{no}) = \frac{1}{3}$
$P(\text{Imp} \text{yes}) = \frac{3}{3}$	$P(\text{Imp} \text{no}) = \frac{2}{3}$

Now with the data set we can make a decision
we can make a decision regarding the first entry.
with the color used and type sports, origin is
Domestic, we can have the decision with the class
label stolen our answer is yes.

→ with this is also possible to make the decision

regarding the possible values of the sampling can't have the class label.

Ex: The sample

$P(X) : \langle \text{Red, suv, domestic} \rangle$

$$P(\text{yes}) = \frac{P(\text{Red/yes}) \cdot P(\text{suv/yes}) \cdot P(\text{dom/yes}) \cdot P(\text{yes})}{P(\text{yes})}$$

$$P(X/\text{yes}) \cdot P(\text{yes}) = \frac{P(\text{Red/yes}) \cdot P(\text{suv/yes}) \cdot P(\text{dom/yes}) \cdot P(\text{yes})}{P(\text{yes})}$$

$$= \frac{3}{5} \times \frac{2}{5} \times \frac{2}{5} = \frac{12}{125} = 0.096$$

$$P(X/\text{No}) \cdot P(\text{No}) = \frac{P(\text{Red/No}) \cdot P(\text{suv/No}) \cdot P(\text{dom/No}) \cdot P(\text{No})}{P(\text{No})}$$

$$= \frac{2}{5} \times \frac{3}{5} \times \frac{1}{5} = \frac{6}{125} = 0.048$$

The class label is stolen the belonged answer for the sample is yes.

u.b

Unit - V

Association Analysis : Basic Concepts

Association rule mining finds interesting Algorithms association correlation relationship.

→ Find association rule from large amount of data.

→ Example of association rule mining is market basket analysis.

→ It is a process of Analyses customer buying habits by find in Association b/w the different items that customer plays in their shopping basket.

→ To perform the association b/w the different items we have two measures they are

1. support

2. confidence

for example support is 2%. means all the transactions under analysis show that purchased milk product & bread product together.

Confidence = 60%. means 60% of the customers who purchased milk also bought the bread.

$$\text{support } A \Rightarrow B = P(A \cup B)$$

$$\text{confidence } A \Rightarrow B = P(B|A)$$

→ Rules that satisfy both minimum support threshold and min confidence threshold are called strong rules.

→ Support and confidence variables occur b/w 0% and 100%.

→ A set of items is referred to as an item set. An item set that contains 'K' items is a K item set.

$\{milk, bread\}$ - 2 item set

→ The occurrence frequency of an item set is the no. of transactions that contain the item set. This is known as frequency, Support Count, Count on item set.

If an itemset satisfies minimum support count then it is a frequent item.

→ Association rule mining:-

Association rule mining is a two process that it can have the two steps

Step 1: find all frequent itemsets; each of these items will occur atleast as frequently as a pre-determined min support count.

Step 2: Generate strong association rules from the frequent itemset.

These rules must satisfy min support & min confidence.

TID	List of item-ids
T100	11, 12, 15
T200	12, 14
T300	12, 13
T400	11, 12, 14
T500	11, 13
T600	12, 13
T700	11, 13
T800	11, 12, 13, 15
T900	11, 12, 13
T1000	

scan D for
count of
each candidate

itemset	sup. count
{1,2}	6
{1,3}	7
{1,4}	6
{1,5}	2
{2,3}	2

compare
candidate
support count
with
minimum
support count

itemset	sup count
{1,2}	6
{1,3}	7
{1,4}	6
{1,5}	2
{2,3}	2

generate L2
candidates
from
L1

itemset
{1,1,2}
{1,1,3}
{1,1,4}
{1,1,5}
{1,2,3}
{1,2,4}
{1,2,5}
{1,3,4}
{1,3,5}
{1,4,5}

scan D for
count of
each
candidate

itemset	sup. count
{1,1,2}	4
{1,1,3}	4
{1,1,4}	1
{1,1,5}	2
{1,2,3}	4
{1,2,4}	2
{1,2,5}	2
{1,3,4}	0
{1,3,5}	1
{1,4,5}	0

compare candidate
support count
with minimum
support count

itemset	sup count
{1,1,2}	4
{1,1,3}	4
{1,1,5}	2
{1,2,3}	4
{1,2,4}	2
{1,2,5}	2

generate
candidate
from L2

itemset
{1,1,2,3}
{1,1,2,5}

scan D
for count
of each
candidate

itemset	sup. count
{1,1,2,3}	2
{1,1,2,5}	2

compare
support-
count
with min
sup. count

itemset	sup count
{1,1,2,3}	2
{1,1,2,5}	2

Generate of candidate itemset and frequent itemset
where min. support count is 2.) 5

The Apriori algorithm uses two properties
for find out frequent item sets. we have the
join and prune

Join:- Is used to find out the $L_1, L_2, \dots, L_k$

The set of frequent one item set, two item set with join operation themselves the $L_1, L_2, \dots, L_k$ are join

Prune:- Prune using the Apriori property all non empty subsets of a frequent item sets must also be frequent

Process of Apriori algorithm:-

For performing the Apriori algorithm we have to follow two steps. They are the join step and prune step, with the help of the join step we prepare a set of candidate k item sets, is generated by joining L_{k-1} with itself ($\therefore$ For example we can find the L_k $k = 1, 2, \dots$)

In the prune step C_k is the super set of L_k ($k = 1, 2, \dots$) i.e., its members may or may not frequent but all of the frequent k item sets are included in C_k .

Process:-

Step-1 - In the 1st iteration of the algorithm each item is member of the set of candidate one-item set C_1 .

The algorithm simply scans the transactions in order to count the number of occurrences of each item.

Step 2:- Suppose that the minimum transaction support count required two ($\text{minimum sup} = \frac{2}{9} = 22\%$)

The set of frequent one-item sets L_1 can then be determined if it consists of the candidate one item sets satisfying minimum support.

step-3 - To discover the set of frequent two item sets L_2 the algorithm uses $L_1 \bowtie L_1$ to generate a candidate set of two item sets C_2 . C_2 consists of two item sets of $2^{|L_1|}$.

step-4 - The transaction in D (Data set) are scanned and the support count of each candidate item set in C_2 is accumulated. (combinations)

step-5 - From the accumulated we have the set of frequent two item sets L_2 is then determined consisting of those candidate two item sets in C_2 having minimum support.

step-6 After that the generation of the set of candidate 3 item sets C_3 . 1st we can take the C_2 then C_3 is calculated by performing $(L_2 \bowtie L_2)$ joining with itself. Finally the $C_3 = \{1, 12, 13\}, \{1, 12, 15\}, \{1, 13, 15\}, \{12, 13, 15\}, \{12, 13, 14\}$

step-7 Based on the apriori property that all subsets of frequent item sets must also be frequent. This can be possible with the help of a apriori prune property.

step-8 The transactions D are scanned in order to determine L_3 consisting of those candidate 3 item sets having minimum support.

step-9 The algorithm uses $L_3 \bowtie L_3$ to generate candidate set of 4 item sets C_4 . The join result in C_4 is $\{1, 12, 13, 15\}, \{1, 12, 13, 14\}, \{1, 12, 13, 15\}$. This

Itemset 3 pruned. Since its subsets $\{I_2, I_3, I_5\}$ is not frequent. For this $c_4 = \emptyset$.

And the algorithm terminates having found all of the frequent item sets.

~~Apriori A1~~

TID List of Items - IDS

T₁ I₂, I₃, I₄

T₂ I₁, I₃, I₅

T₃ I₂, I₄

T₄ I₁, I₂, I₃

T₅ I₅

T₆ I₃, I₄, I₅

T₇ I₁, I₂, I₃

T₈ I₂, I₄, I₅

T₉ I₂, I₄, I₅

c₁

item sup-count

I₁ 3

I₂ 6

I₃ 5

I₄ 5

I₅ 5

L₁

item

sup-count

I₁ 3

I₂ 6

I₃ 5

I₄ 5

I₅ 5

Apriori Algorithm:-

input: Database, D , of transactions; minimum support threshold, $\min\text{-sup}$

output: L = frequent itemsets in D

method

(1) L_1 = find-frequent-1-itemset(D);

(2) for ($k=2$; $L_{k-1} \neq \phi$; $k++$) {

(3) C_k = apriori-gen($L_{k-1}, \min\text{-sup}$);

(4) for each transaction $t \in D$ { // scan D for counts

(5) C_t = subset(C_k, t); // get the subsets of

(6) t that are candidates

(7) for each candidate $c \in C_t$

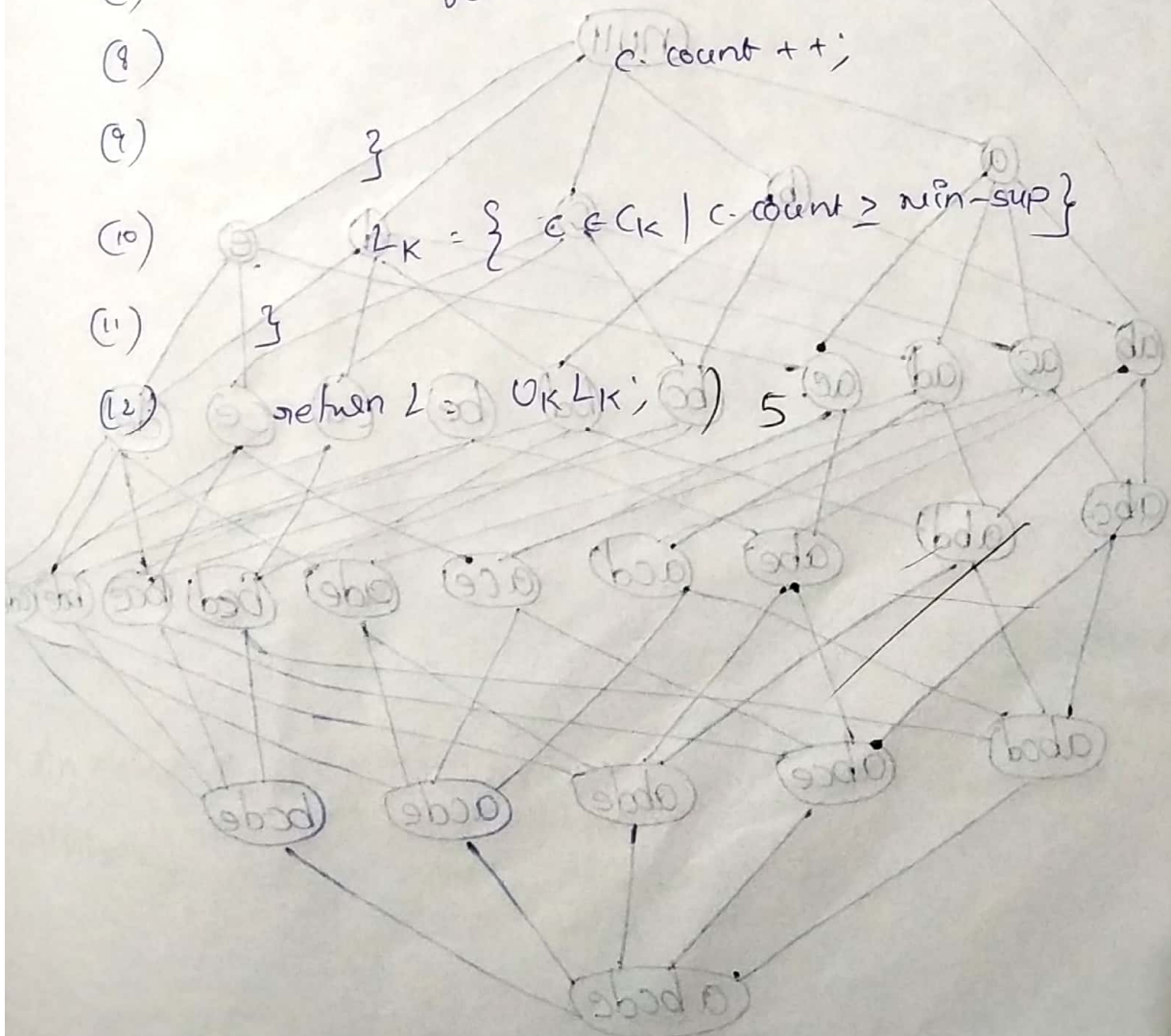
(8) $c.\text{count}++$;

(9) }

(10) $L_k = \{ c \in C_k \mid c.\text{count} \geq \min\text{-sup} \}$

(11) }

(12) return $L = \cup_k L_k$; }



→ Compact representation of frequent itemset

The no. of frequent itemset produce from a transaction dataset can be very large it is useful to identify a small representative set of item sets from which all other frequent item sets can be derived

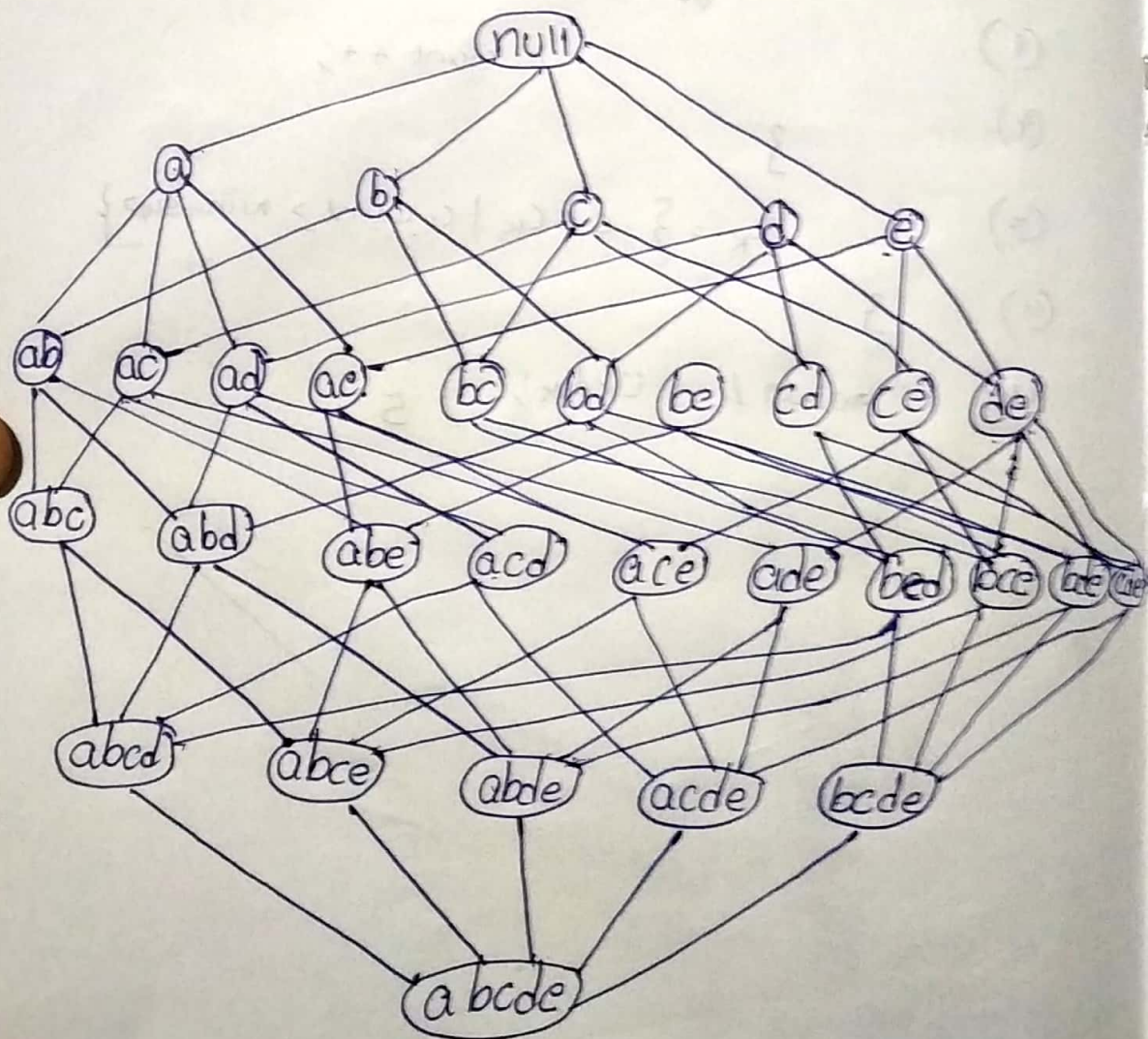
we have two methods for representing frequent item sets in term of compactness. They are

1. Maximum frequent item sets

2. closed frequent item sets

Maximal frequent item set:-

It is defined as frequent item sets for which none of its intermediate super set are frequent



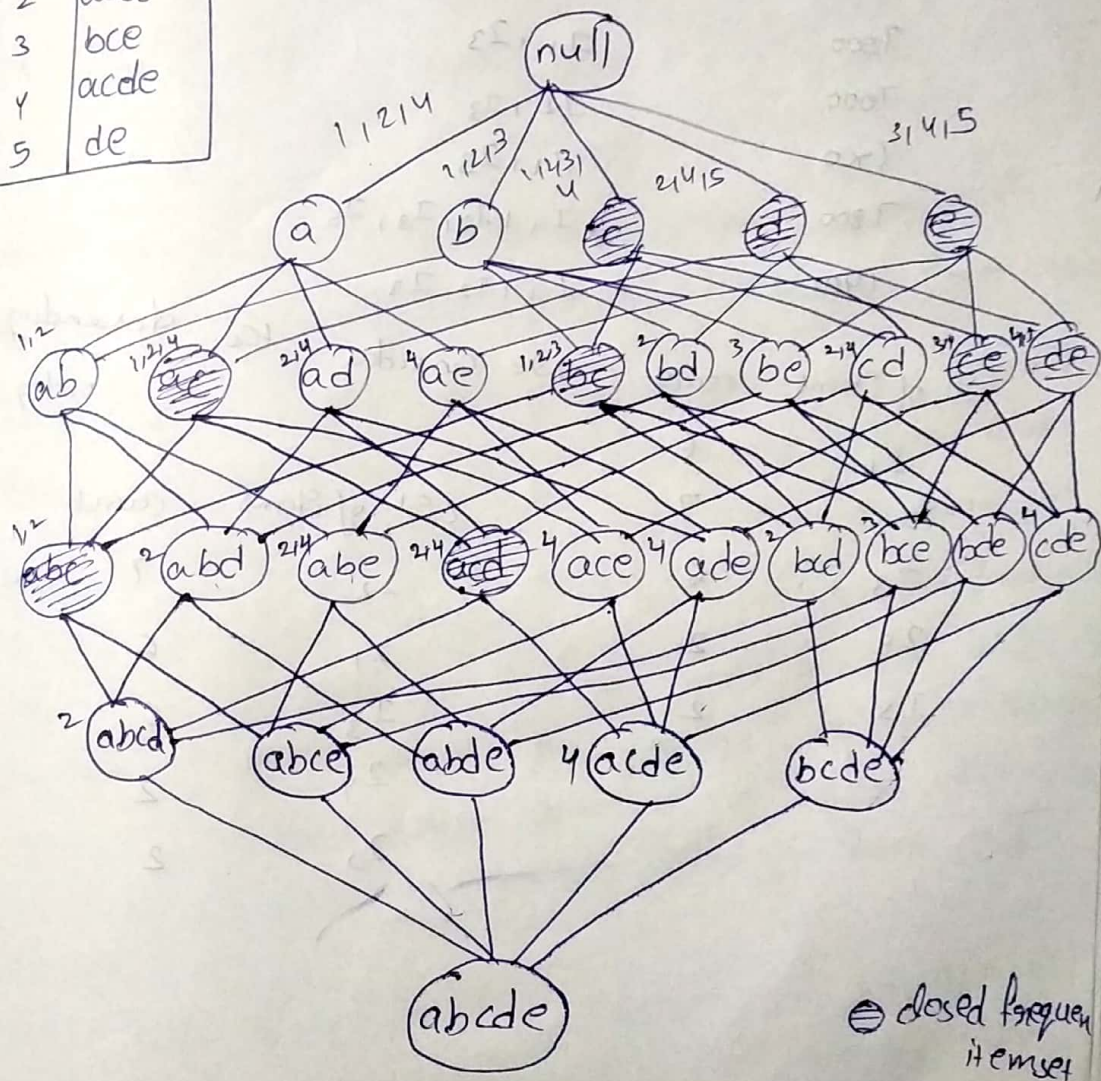
→ closed frequent item set

closed frequent item set provide minimal representation of item set without losing their support information.

Def An itemset X is closed if none of its immediate superset has exactly the same support count as X

TID	Items
1	abc
2	abcd
3	bce
4	acde
5	de

Min sup = 40%



An example of the closed frequent itemsets (with minimum support count equal to 40%.)

→ FP tree :- (Frequent pattern tree)

Apriori Algorithm example

Tid

List of item-ids

T₁₀₀

I₁, I₂, I₃

T₂₀₀

I₂, I₄

T₃₀₀

I₂, I₃

T₄₀₀

I₁, I₂, I₄

T₅₀₀

I₁, I₃

T₆₀₀

I₂, I₃

T₇₀₀

I₁, I₃

T₈₀₀

I₁, I₂, I₃, I₅

T₉₀₀

I₁, I₂, I₃

Item	Count
I ₁	6
I ₂	7
I ₃	6
I ₄	2
I ₅	2

List of Items

count

We consider

the descending order

I₁

6

I₂

7

I₃

6

I₄

2

I₅

2

List of Items

count

I₂

7

I₁

6

I₃

6

I₄

2

I₅

2

Here we use FP growth algorithm for finding
FP tree representation for the frequent-item sets
→ Originally in the Apriori we form the frequent
item sets with the help of join and prune steps
→ But with the Apriori we have the problem for
finding frequent itemsets for more of the transac-
tions can be happened to be complex for this
instead of the apriori that entire dataset data
using a compact data structure called a FP tree
and extracts frequent itemsets directly from this
structure

→ 1. Also for finding the pattern as usual of apriori
we can take the list of items and we can take
the support (minimum support threshold value = $0.2 = 20\%$)
then I can form a table with the list of items
with respect to minimum support of threshold
value.

List of Items	sup-count
21	6
22	7
23	6
24	2
25	2

2. Now this is a FP tree for this we can find out the pattern.

For finding the pattern we can use descending order (High to low) for the minimum support-count happens as equal or greater than happens item sets.

Now we have the pattern is

List of item	count
22	7
21	6
23	6
24	2
25	2

Then the pattern is $\{22, 21, 23, 24, 25\}$

3. Based upon the formed pattern we reorder the items of belonged q transactions because we have in our problem the dataset can have q transactions

Tid	List of Items (ordered)	List of items re-ordered
T100	I1, I2, I3	22, 21, I3
T200	I2, I4	I2, I4
T300	I2, I3	I2, I3
T400	I1, I2, I4	I2, I1, I4
T500	I1, I3	I1, I3
T600	I2, I3	I2, I3
T700	I1, I3	I1, I3
T800	I1, I2, I3, I5	I2, I1, I3, I5
T900	I1, I2, I3	I2, I1, I3

→ example problem:

Tid	itemssets
1	b , a, c, d, g, m, p
2	a, b, c, f, h, m, o
3	b, f, h, o
4	b, k, c, p
5	a, f, c, l, p, m, n

$$\min\text{-sup} = 3$$

item	count
a	3
b	3
c	4
d	4
m	3
p	3

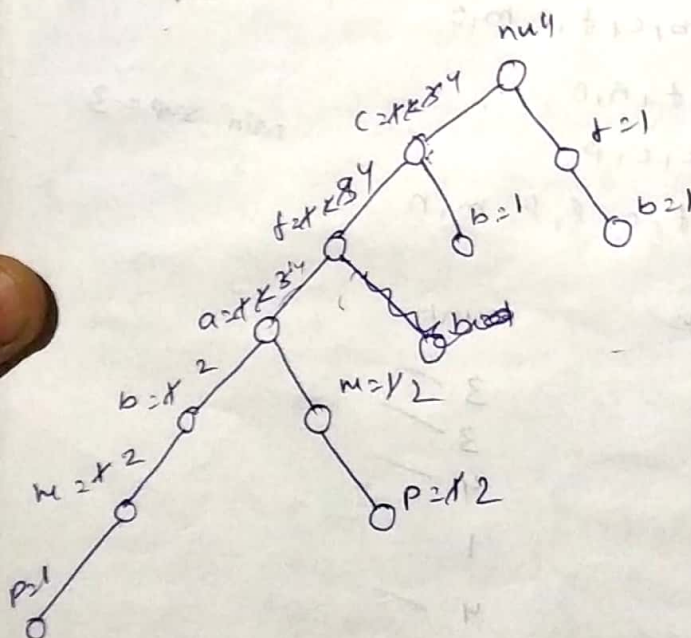
descending order

c	4
d	4
a	3
b	3
m	3
p	3

{c, d, a, b, m, p}

2nd list of item

- 1 c, d, a, m, p
- 2 c, d, a, b, m
- 3 b, d, b
- 4 c, b
- 5 c, d, a, m, p



Unit - VI

K-mean Alg

Cluster Analysis

DB scan "

Agglomerative

cluster is nothing but we form the group with the similar properties of the items are the data are object.

- Cluster means forming the groups.
- It was one of the data mining technique.
- Here we perform unsupervised learning, we use unsupervised learning actually in the classification. Here we have two types of learning supervised. supervised means we know the class label.
- unsupervised means we don't know the class label (deduction)
- Here we can take the data set. on that dataset we can apply the cluster algorithm. Then we obtain some of the clusters of the data can be happen.

Advantages

- Not effected by addition of new objects
- do need to worry about signs of the data
- computation process is very simple. Because here we are using the formulas are mean, ~~Eucledian~~ Eclidian distance.

Applications

- Medical - Establish taxonomy of diseases
- www - world wide web
- Sismology - trace out the earthquakes of the data
- Business marketing - find the targeted customers

we have the two types of clusters there are the central and hierarchical.

→ In the central we form the cluster with the main objects not to form the hierarchical manner.

Ex- Amazon

Books Electronic clothes

→ It's also the cluster of will be happened with the internality of the properties of the objects

electronic gadget

↓
phone

apple mi samsung

mi 2 mi 2

mi 2 3

K-mean algorithm :-

K-mean algorithm is the prototype based or the partitioned clustering technique that attempts to find a users specified number of clusters (K) which are represented by their centroids / mean values

Procedure -

Step 1 - Take mean values

Step 2 - Find nearest number of mean and put in the cluster

Step 3 - Repeat the above process i.e. 1, 2 steps until we get the same mean.

$K = \{ 2, 3, 4, 10, 11, 12, 20, 25, 30 \}$
In this we have to take two mean values randomly

$\mu_1 = 4$

$\mu_2 = 12$

$K_1 = \{ 2, 3, 4, 10 \}$

$K_2 = \{ 11, 12, 20, 25, 30 \}$

$\mu_1 = 19.1$

$$m_1 = 4$$

$$m_2 = 12$$

$$K_1 = \{2, 3, 4\}$$

$$K_1 = \{10, 11, 12, 20, 25, 30\}$$

$$m_1 = \frac{9}{3} = 3$$

$$m_2 = \frac{10+11+12+20+25+30}{6}$$

$$= \frac{19}{3} = 6.33$$

$$K_2 = \{2, 3, 4, 10\}$$

$$K_2 = \{11, 12, 20, 25, 30\}$$

$$m_1 = \frac{19}{4} = 4.75$$

$$m_2 = \frac{11+12+20+25+30}{5}$$

$$= 19.6 = 20$$

$$K_3 = \{2, 3, 4, 10, 11, 12\}$$

$$K_3 = \{20, 25, 30\}$$

$$m_1 = 7$$

$$m_2 = 25$$

$$K_4 = \{2, 3, 4, 10, 11, 12\}$$

$$K_4 = \{20, 25, 30\}$$

$$m_1 = 7$$

$$m_2 = 25$$

Then we stop the process, because got the same mean values finally $K_1 = 7$, $K_2 = 25$.

Basic k-mean algorithm:-

1. select k points as initial centroids
2. repeat
3. From k clusters by assigning each point to its closest centroid
4. Recompute the centroid of each cluster
5. Until centroids do not change.

19-3-20
→ DB Scan :- Density based spatial clustering of a
application with noise points

→ DB scan algorithm is based on highly the density.
→ Density based cluster locates regions of high density
that are separated from one another by regions
of low density.

→ DB scan is a simple and effective density based
clustering algorithm.

→ In the density based clustering algorithm we can
use the center based approach i.e., density is estima-
ted for a particular point in the dataset by counting
the number of points within a specified radius this
we call it as Epsilon of that point.

→ The number of points within a radius of EPS is
minimum points are three of that data point itself.

In the DB scan algorithm we have 3 types of points
are involved they are

1. In the interior of a dense region (a core point)
2. on the edge of dense region (border point)
3. In a sparsely occupied region (a noise or
background point)

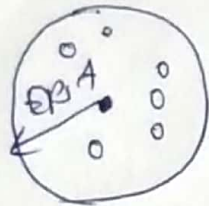
Corepoint - These points are in the interior of a
density based cluster.

A point is the core point if the number of points
within a given neighbourhood around the point is
determined by the distance function $EPS - EPS$

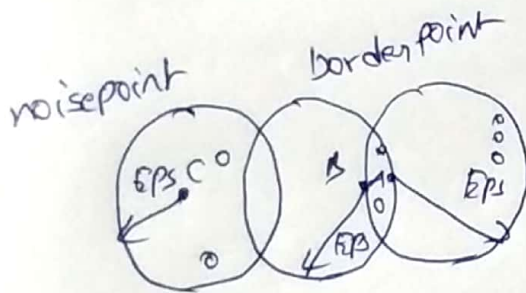
exceeds certain threshold minimum points.

Border point - A borderpoint is not a corepoint. But forms with in the neighbourhood of a corepoint.

Noise point - A noise point is any point i.e., neither a core point nor a borderpoint.



center-based density



DBSCAN algorithm

1. Label all points as core, border or noise points
2. Eliminate noise points
3. put an edge between all core points that are within Eps of each other.
4. make each group of connected core points into a separate cluster.
5. Assign each border point to one of the cluster of its associated core points

7-1-20

Unit-2

Data preprocessing

Dimensionality Reduction: Time/cost parameters

→ Attribute is nothing but a dimension.

Dimensionality reduction means we can reduce the attributes of the dataset.

From old attributes of the data we can follow the new attributes of the data. It is called dimensionality reduction.

We can take the dataset the dataset can have the attributes like sno, sname, saddr, smarks. These attributes are called old attributes. The old attributes are reduced to data by using following the data is in the new attributes.

The old attributes from the subset of the attributes the subset of data is called new attributes.

We can measure the dimension of a dataset with an attribute we can measure. The same will play by the dimension of attribute.

We can create the dataset by the new attributes with the combination of old attributes.