

Data base management system

data: Independent Component.
record: interrelated Component.

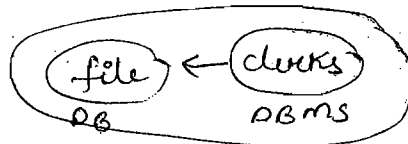
Data Base: Set of record.

DBMS: s/w used to create, manipulate

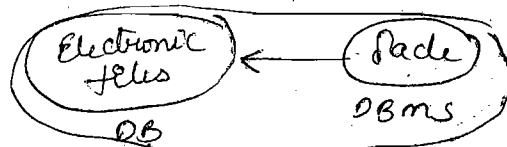
DB + DBMS → Data base system (DBS)

Paper - Data base
Pen - DBMS

↓
Data base system.



DBS



DBS

now a days wide varieties of DBMS are available from many vendors like Oracle.

How Construction

Requirements



2D-Plan



Prototype (3D-Plan)



Physical Construction.

Database Design:

Requirements.

↓ (conceptual)

Plan (ER diagrams.)



Logical design stage (Function dependencies, normalization, SQL)



Physical Construction. (B, B+, files, ...)

once the construction is over, we cannot change it.

→ in single user applications we may not think of concurrency and transaction management, but if we are using the multiuser application then they comes into picture.

Data Bases

OLTP

OLAP

on line transaction
processing

online analytical
processing

(Current DB)

(Historical DB)

It is simply called
as Data Base

It is called as
Data warehouse

← Data mining
algorithms

Size is MB, GB

Size in TB

R/W

R

Normalization ✓

no Normalization

Simple SQL

Complex SQL (join)

another classification

Commercial DB : Inventory control, library, student info system.

Multimedia DB : voice, video clips, images.

Temporal DB : attaching time aspect to data. Ex: Railway reservation.

GIS DB : (Geographical information system) Ex: To give ways to reach dest. (route map)

Distributed DB : web environment. (Consistency is very high) (Google earth)

Reductive DB : Rule based. ($A=B, B=C \Rightarrow A=C$)

Document DB

active DB

web DB.

reductive DB : If we buy one thing, then there is chance of
buying the other.

Document DB : Google, Content information system.

active DB : Triggers Flooded with Triggers

web DB : Any thing can become a web data base if
they are available in web environment.

Different types of DBMS

Relational DBMS

OR DBMS (Object Relational DBMS) & Object Oriented DBMS

OO DBMS (Object Oriented DBMS)

ORDBMS - For RDBMS we added object oriented principles.

OODBMS - For object oriented DBMS we added relational principles.

Except for multimedia, GIS, document and web we can use ORDBMS

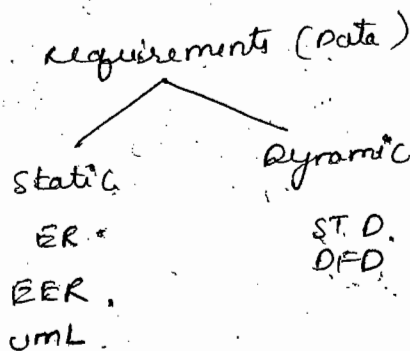
For the four we used OODBMS.

we are going to study only RDBMS.

ER-diagrams

model: imitation of reality.

different models: ER, UML, DFD, IPO etc.



At time of development of DBMS, the data is not going to be changed and so it is static data and we can go for any one of techniques available in this category. let us go for ER.

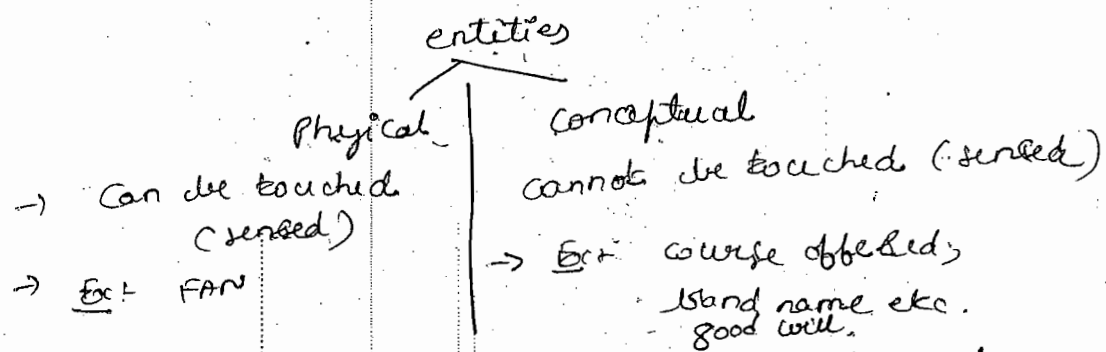
→ It gives the pictorial representation for all the requirements in the information system.

→ It is a data model, in which information stored in the data base is viewed as set of entities and their relationships.

→ The components of ER diagrams

- (i) Entities
- (ii) Relationships
- (iii) attributes.

A thing or object with an independent existence is called an entity.



→ Entities can have conceptual or physical.

→ storing conceptual entities is a big problem.

Ex: student

Roll no	name	age
1	a	5
2	b	6
3	c	7

} set

we have 3 entity occurrences, all these are called

entity set.

Rel
m
Ty
1
7

n

un

multi
w/w

other
deg

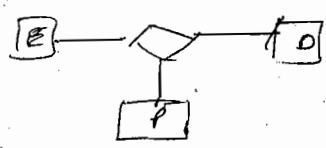
Relationships

Association among the entities is known as relationship
Types -

Binary



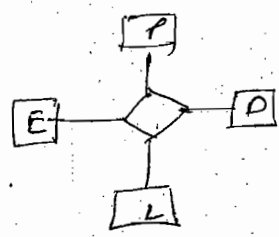
Three



E - employee
D - department
P - project

ER diagrams cannot take care of ternary relationship.

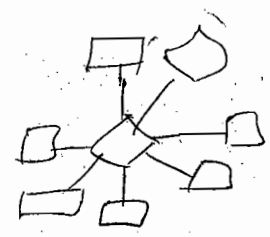
four



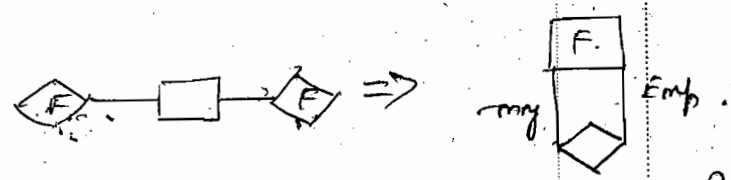
quaternary relationship.
L - location

ER cannot take care of four also.

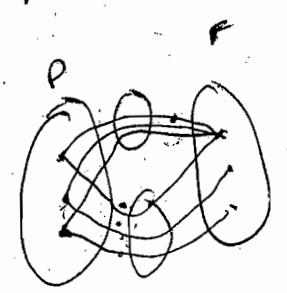
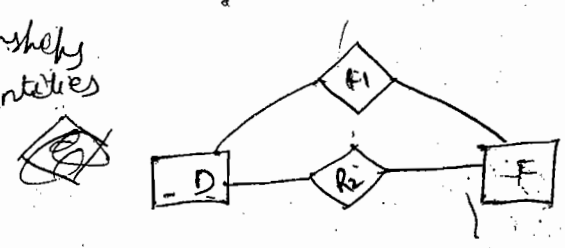
n-ary



unary



multiple relationships
b/w same entities



other type of classification -
degree of relationship

- 1:1
- 1: many
- many: many

Identify : unary, binary, ternary
non identify degree

total and partial participation.

Constraints on Relationships

(i) Structural Constraints: (applicable for binary relationships)

- Cardinality Ratio
- Participation Constraints

(ii) Overlap Constraints: } applicable to EER diagrams
↓
extended

(iii) Covering Constraints:

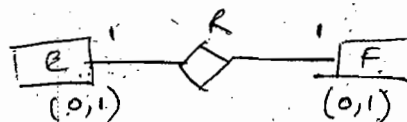
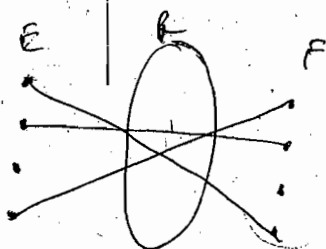
Cardinality Ratio: It describes the maximum number of possible relationship occurrences for an entity to an another entity through the relationship.

Participation: It describes whether all or only some entity occurrences are participating in a relationship. It is total if all entity occurrences are participating, otherwise, it is partial.

minimum cardinality = 0 for partial participation and equal to 1 for total participation.

maximum cardinality = 1. If one entity occurrence relates to another entity occurrence.

and max cardinality = n. If one entity occurrence relates to multiple entity occurrences.



min Card(E,R) = 0 Partial
max Card(E,R) = 1

min Card(F,R) = 0 Partial
max Card(F,R) = 1

min
max

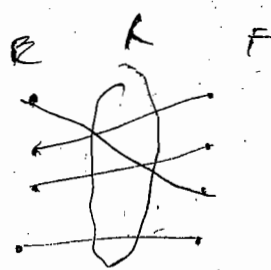
→

Describe
Er

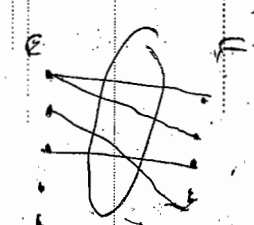
is

(

msheps)

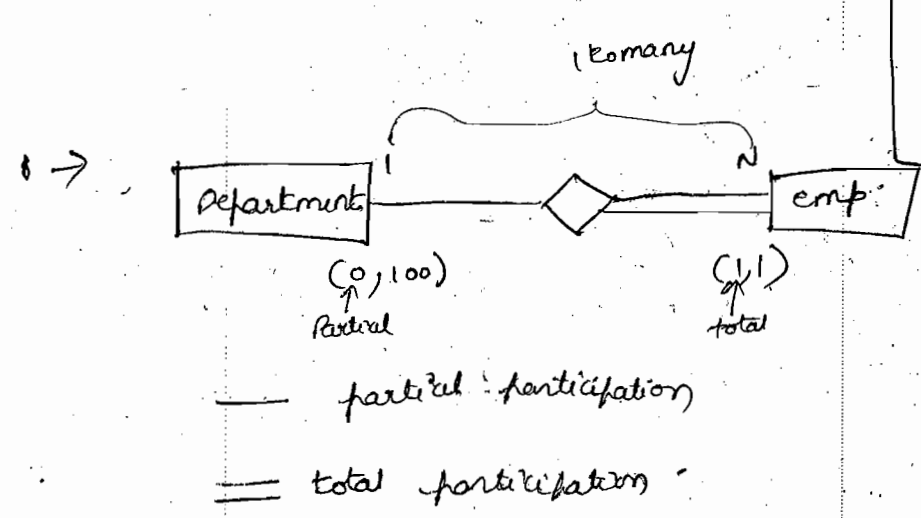


min 20 (par) min = 1 (Total)
max = 1 max = 1



min 20 (par) min = 1 (Total)
max = 2 (N) max = 1

rams



min Cardinality
give partial or total
max Cardinality
give the maximum
no of entities as
it relate to.

ly.

Agging 1 to many $\Rightarrow$ one dept can have many employees.

res are

m

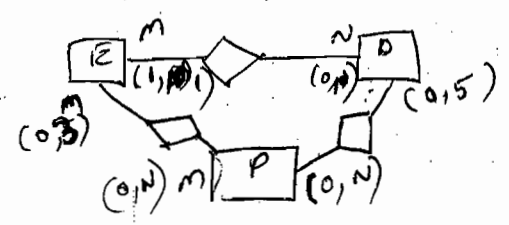
ce

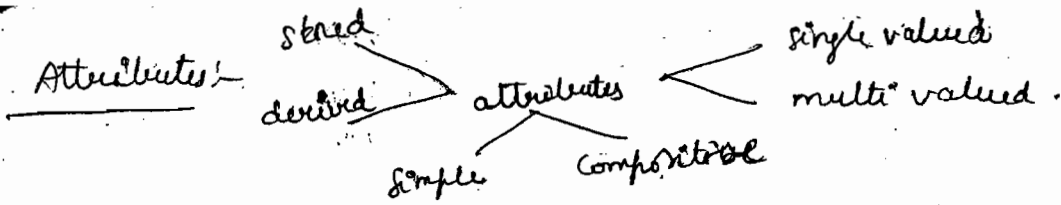
ce

]

Description

Employees may belong to different dept and participate in different projects. minimum number of employees in a project should be '100' and a single employee can participate max upto 3 projects and max of 5 depts can participate in project. Then every employee should belong to a dept. but they need not participate in the project; (note cardinality ratios are take only max value).

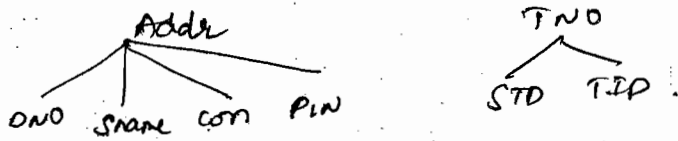




Simple: It will have a single atomic value and cannot be divided further.

Ex: Roll no, Company name, Job code etc.

Composite: It is the combination of many atomic values and will have complete meaning if we combine all the attributes.



Single valued: A attribute that have always a single value.
Roll no, Gender, Date of birth

multi-valued: An attribute that will have multiple values.

Ex: qualification of a person, Tel id of a person, Email id of a person

Stored: An attribute that supplies the value is called a stored attribute.

1)	DOB	AGE
2)	I/D	R/D duration fine.
3)	RNO	name
	1	a
	2	b
	3	c
	4	d
	5	e

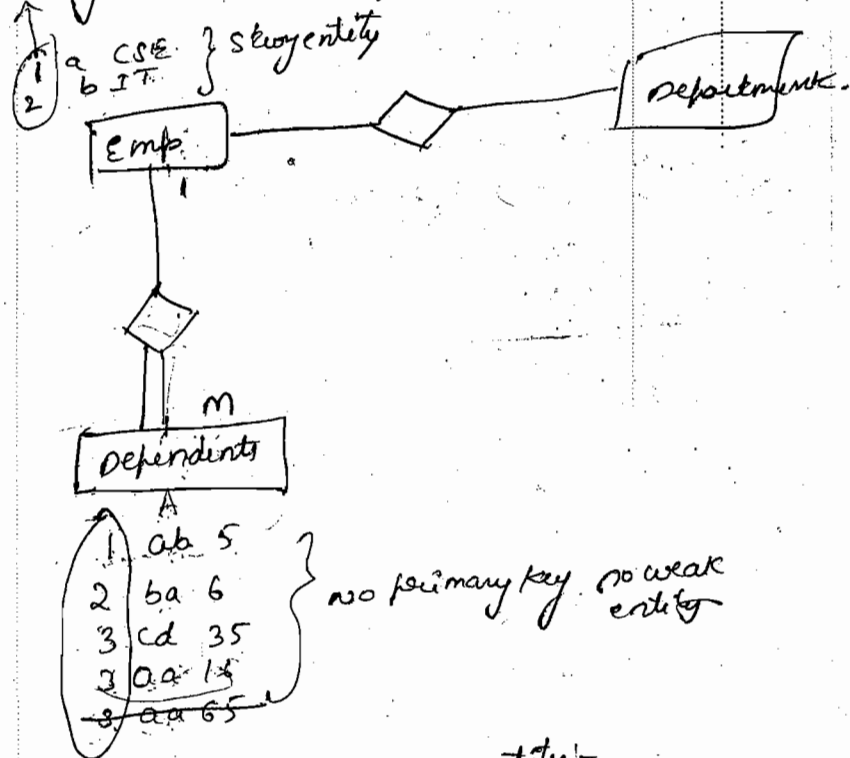
ONO	ename	OS
1	CSE	223
2	IT	4
3	ECE	1

- An attribute can be derived from a single attribute. ①
 An attribute can be derived from multiple attributes. ②
 An attribute can be derived from a separate table. ③.

special type of entities

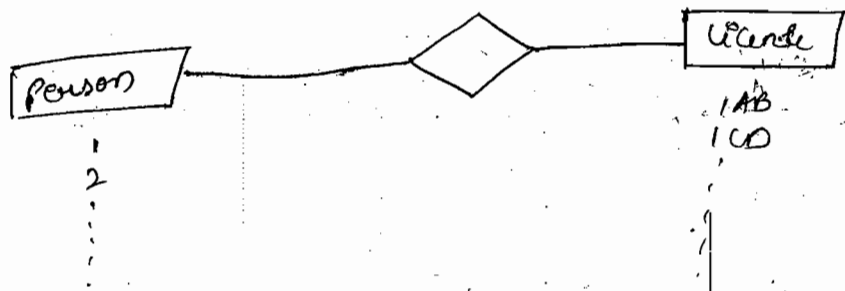
- 1) Strong entity
- 2) weak entity
- 3) Associative entity

Primary Key.



Characteristics of strong and weak entity:-

- (i) They will be identified by primary keys or dependencies.
- (ii) The relationship b/w strong and weak entity is one to many.
- (iii) The participation of strong entity is partial but weak entity is total.



here even the license is depending on person entity, license entity cannot be called as a weak entity because it will have its own primary key.
 ∴ classification of entities may be done by using

primary keys.

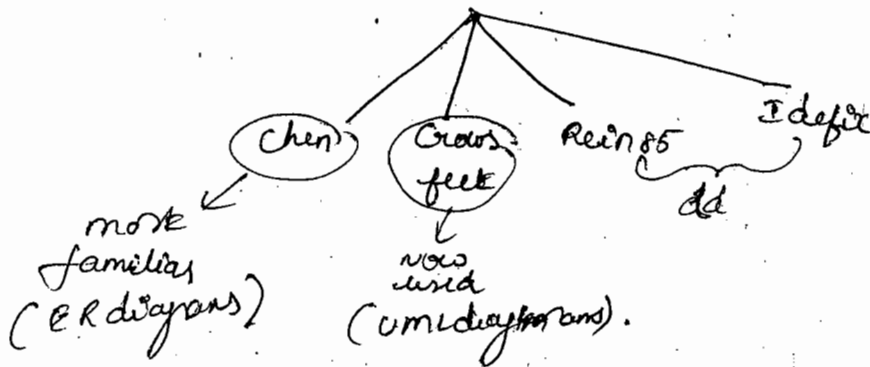
Associative entity:

The entity that looks like a Relationship but not a Relationship is known as associative entity.

Ex: Person, Product, Sale.

Person and product are entities but sale is looking like a Relationship b/w these two but it is a separate entity with its own attributes.

notations available to give pictorial representations



- strong entity.

- weak entity.

- associative entity.

- Recursive Relationship.

- multi-valued attribute.

- derived attribute.

- Relationship.

- identifying Relationship (strong & weak).

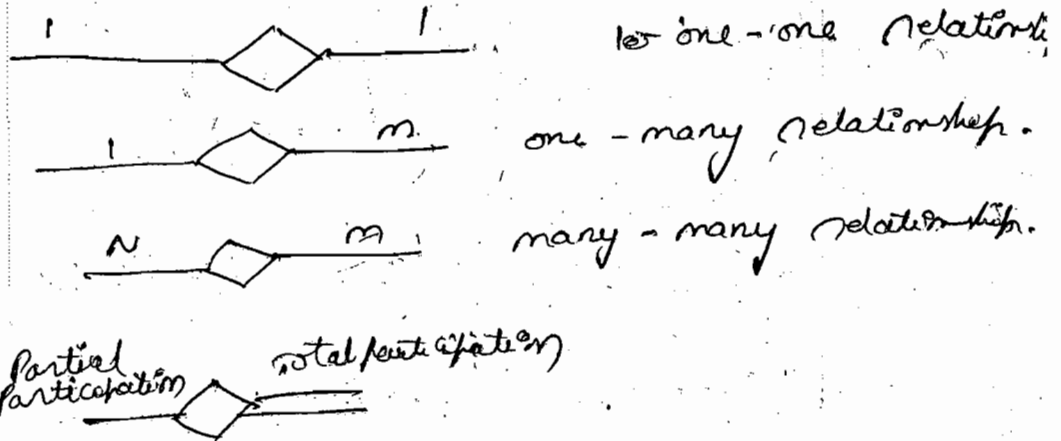
- attribute.

- primary key attribute.

- Composite attribute.

- discriminated attribute of weak entity.

The attribute of the weak entity that is added to the combined with primary key attribute of the strong entity to declare the combination as a primary key for the weak entity is known as discriminator attribute.



→ now ER diagrams are to be converted into tables & objects depending on RDBMS or ORDBMS respectively.

→ we have conversion of ER diagrams into tables (7 step criteria)

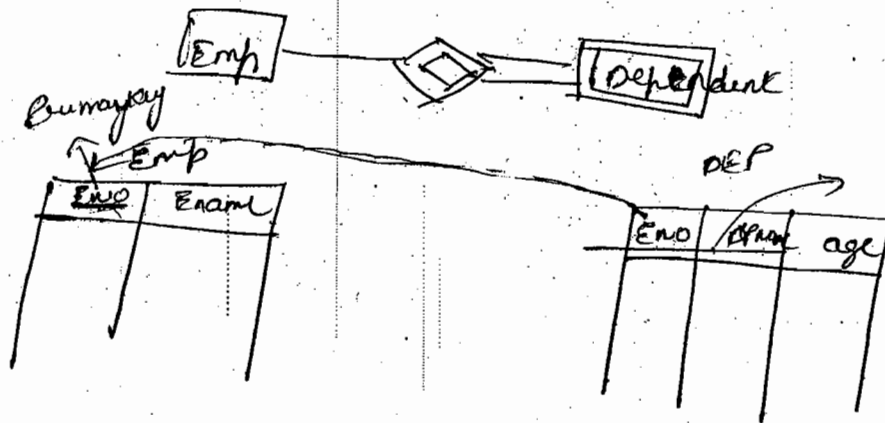
Step 1:- Conversion of strong entities:

- For each strong entity in ER diagram, create a separate table with the same name.
- Create all simple attributes.
- Break the composite attributes into simple attributes and create them.
- Choose a primary key for the table.

Step 2:- conversion of weak entities:

- For each weak entity create a separate table with the same name.
- Include primary key of the strong entity as a foreign key in this table.
- Select the primary key attribute of strong entity and discriminator attribute of weak entity.

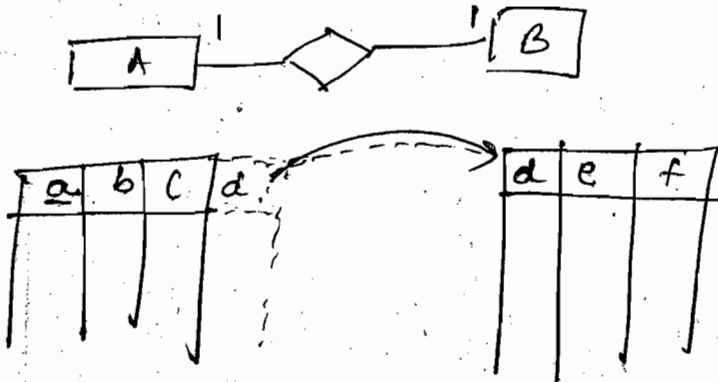
and declare them as Primary key.



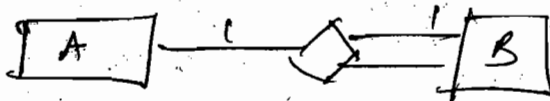
Cono

Step 3) Conversion of 1 to 1 relationships

- a) For each 1 to 1 relationship, say b/w. A and B tables, modify either A or B to include the primary key of other table as a foreign key.



- b) If one of these table either A or B has total participation, then modification can be done must be done only on that side.



else partial side null value - means waste

Step
Con
a)

in
stud
many
r

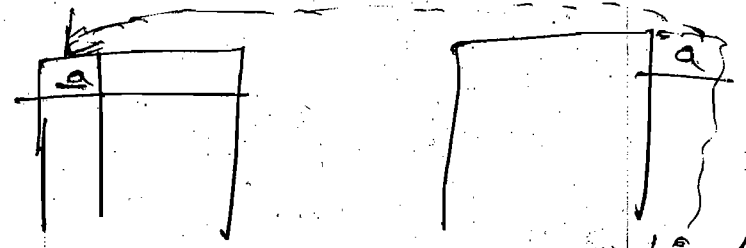
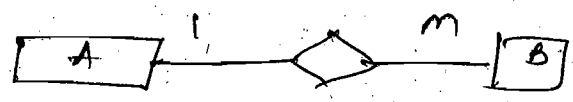
Step
a)

b)

Step 4:- Conversion 1 to many Relationships:

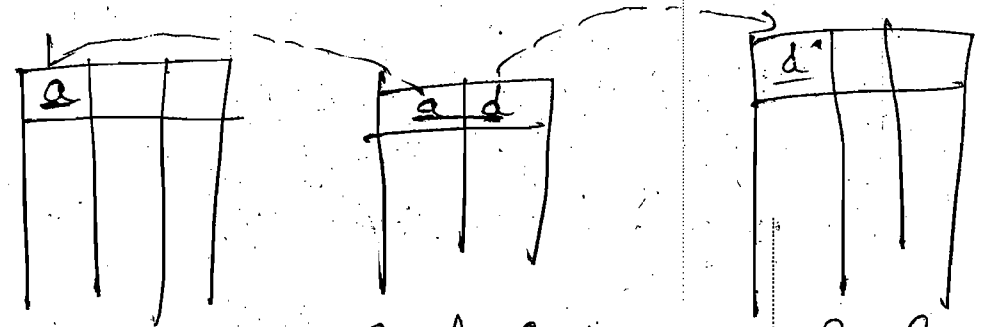
Correct
a) For each 1 to many relationship type A and B modify many side relation to include primary key from one side as a foreign key.

ex
stud many
dept no. 1

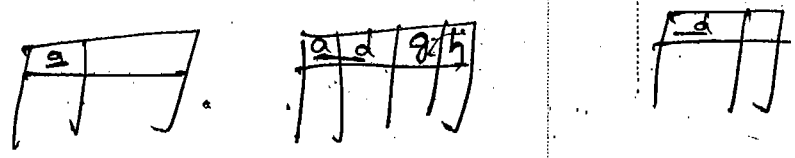
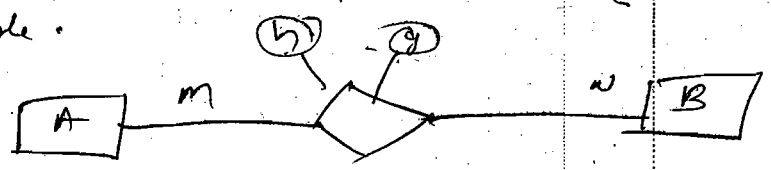


Step 5:- Conversion of many to many Relationships

a) For each many to many relationship type A and B, Create a separate table and include primary key of both the tables as foreign key.



b) If relationship is having one or more attributes, these must also be included in the table.



1 B
the
key.

total
rest de
de
- more waste

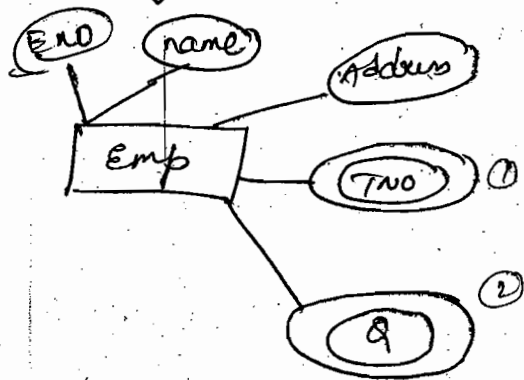
Step 6:

Conversion of n-ary relationship.

- For each n-ary relationship, create a separate table and include primary keys of all other entities as a foreign key.
- If the relationship has some attributes that must also be included.

Step 7: conversion of multivalued attributes:

- For each multivalued attribute, create a separate table, then include all of its simple attributes.
- Include the primary key of the original table as a primary foreign key.



Emp

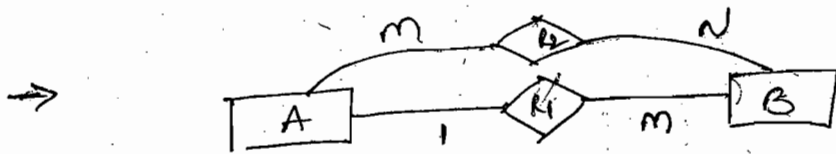
ENO	name	DNO	Sty	City

①

Empno	TNO

②

ENO	Q



min: ? 3 A B R₂
 max: ? A B R₁ R₂

A

a		

B

b		a

R₁

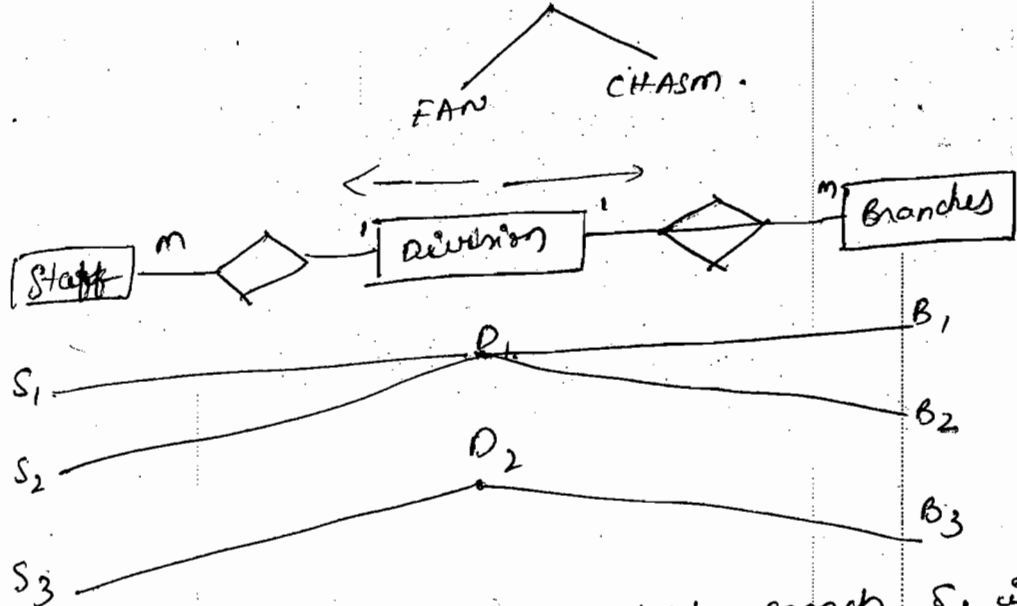
a	b

R₂

a	b

R₁ can also be represented as a separate table.

→ Traps in ER diagrams:



we cannot infer in which branch S₁ is working.
 9E is called FAN trap.

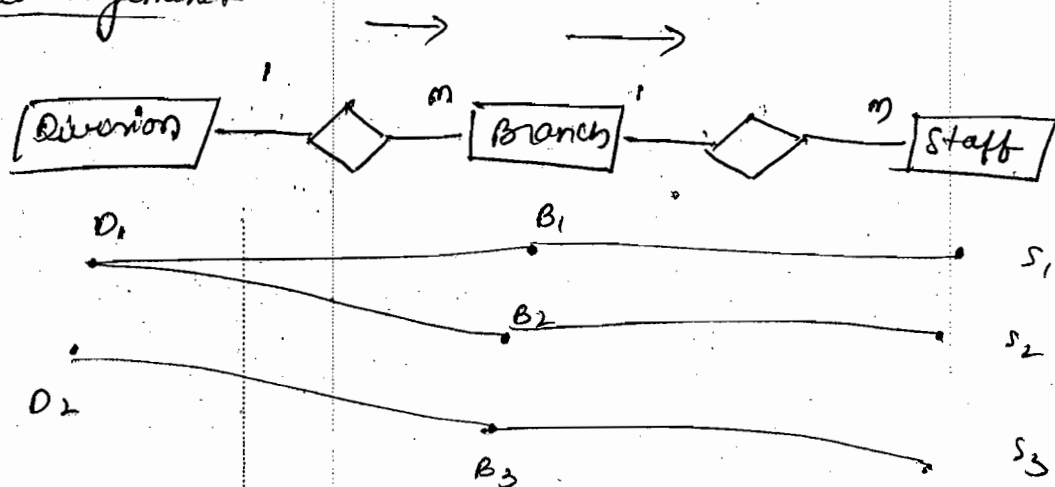
this problem is
Identification of the fan trap:
 A fan trap may exist if two or more one to
 many relationships fan out from the entity.

Information cannot be retrieved to the full
 extent if there is a fan trap.

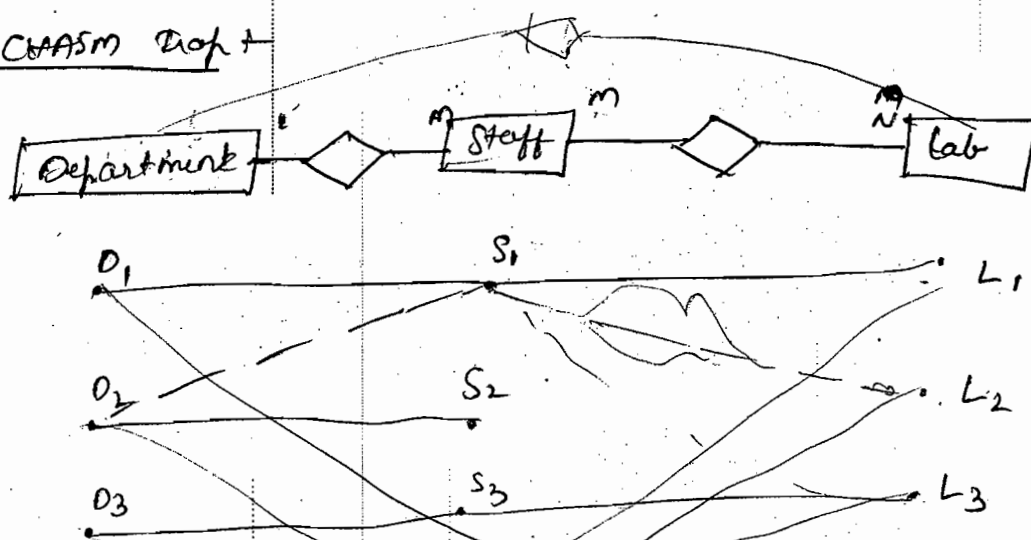
Solution 1

If there is a fan trap, rearrange the entities and redraw the ER diagrams.

Rearrangement



Chasm trap



Lab 2 comes under what dept? 4k is chasm

trap.

Identify chasm trap: A chasm trap may occur if one or more relationships with a cardinality of '0' forming a path b/w the related entities.

Solution Create an additional relationship b/w these related entities.

The advantages and of Relate ER diagrams :-

- (i) conceptually it is very simple
- (ii) better visual representation.
- (iii) it is an effective communication tool among users, domain experts and data base designers.
- (iv) it is tightly integrated with relational model, so converting ER diagrams to tables is very simple.

Disadvantages :-

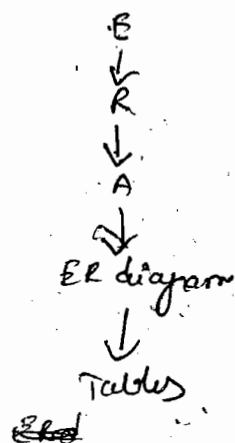
- (i) limited constraint & specification
Ex: min cardinality. $\text{max} = n$ 'even if more than'
- (ii) loss of information content.
Ex: FAN, CHASM trap (may be)
- (iii) limited relationship representation.
only binary

EER diagrams :- (Extended Entity Relationship)

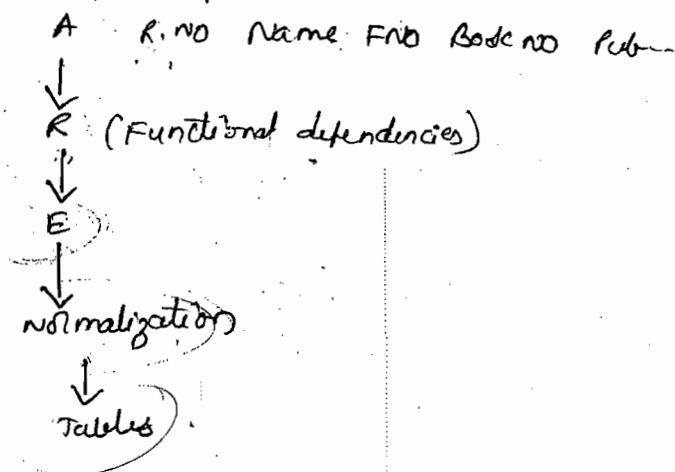
= ER diagrams + Extensions

These extensions are sub classes, super classes, Categories, attribute and relationship instances.

Top down



Bottom up



→ In BU, to identify the relationships among the attributes we use functional dependencies.

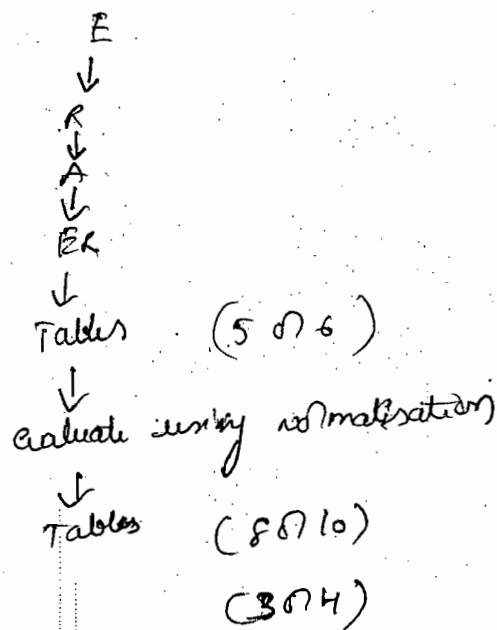
→ In TD, we use ER diagrams.

→ In BU, we use functional defns and normalization.

→ when project size is large, we use TD and for small we use BU.

→ BU is more powerful but in practice very difficult to apply.

→ In actual practice, we go for mixed approach.



Functional dependency:

deno Pub → we look at relationships among the attributes.

RNO name age branch.

RNO → name

name → RNO ⇒ RNO → name, age, branch.

RNO → age

RNO → branch

its we.

RNO name age branch FNO Fname DepNO Depname

① RNO → name, age, branch, depno.

② depno → d name.

③ F.NO → F.name. (faculty no.)

④ F.NO → depno, depname.

small

difficult to

~~From the test~~

	a	b
a → b	1	2
x b → a	2	2
	3	2
	4	5
	6	5

FD's are depending

The properties of FD's:-

→ The FD's will deal with one-one relationship among the attributes, sometimes it may be

one-many or many to one.

→ FD's must be defined on schema, not on instances of the schema.

ab → cd

a → b

ab →

→ FD must be a non trivial & completely non trivial.
A trivial dependency is the one where RHS is a subset of LHS

Ex for trivial: $abc \rightarrow bc$

• An FD is non trivial if atleast one of the RHS attributes are not part of LHS attributes.

An FD is completely non trivial if none of the RHS attributes are part of the LHS attributes.

Ex for non trivial: $abc \rightarrow cde$

Ex for complete NT: $abc \rightarrow def$.

→

A	B	C
a ₁	b ₁	c ₁
a ₁	b ₁	c ₂
a ₂	b ₁	c ₁
a ₂	b ₁	c ₃

Identify the FD's that holds good.

$A \rightarrow B$

$B \nrightarrow C$

$C \nrightarrow A$

$AC \rightarrow B$

(Because A itself can determine B and even we add c also it will no change).

$BC \nrightarrow A$

Interval
is a

RHS

u RHS

→

	A	B
t_1		
t_2		
①	If t_1 and t_2 agree here	Then they must agree here
②	If t_1 and t_2 disagree here	They may agree or disagree here.

A = roll no

B = name

1	a
2	b
3	a

1	a
1	a

Roll no → name

→

A	B	C
1	2	3
4	2	3
5	3	3

$A \rightarrow B$

$A \rightarrow C$

$B \rightarrow C$

$\left. \begin{matrix} AC \rightarrow B \\ AB \rightarrow C \end{matrix} \right\}$ A is already able to determine B and C

$B \nrightarrow A$

$C \nrightarrow A$

B and
u no

→

A	B	C
10	b ₁	c ₁
10	b ₂	c ₂
11	b ₄	c ₁
12	b ₃	c ₄
13	b ₁	c ₁
14	b ₃	c ₄

$$A \not\rightarrow B$$

$$B \rightarrow C$$

$$C \not\rightarrow B$$

$$B \not\rightarrow A$$

$$C \not\rightarrow A$$

$$(AB \rightarrow C)$$

$$BC \not\rightarrow A$$

$$AC \rightarrow B$$

→

A	B	C	D
a ₁	b ₁	c ₁	d ₁
a ₁	b ₁	c ₂	d ₂
a ₂	b ₁	c ₁	d ₃
a ₂	b ₁	c ₃	d ₄

$$A \rightarrow B$$

$$B \not\rightarrow C$$

$$A \not\rightarrow D$$

$$B \not\rightarrow A$$

$$C \rightarrow B$$

$$B \not\rightarrow C$$

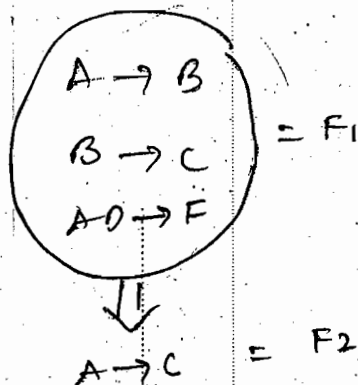
$$C \not\rightarrow A$$

$$D \rightarrow B$$

$$C \not\rightarrow D$$

$$D \rightarrow C$$

→



$$(F_1) + F_2 = F \rightarrow \text{no sim}$$

$$\downarrow$$
 semantics ?

Functional dependencies (F_1) can be identified by using semantics from the problem domain and it can be called as $\bullet F_1$ Set.

Additional FDS (F_2) can also be derived from F_1 .
Then total FDS $= F_1 + F_2$

Before making a move to the normalization process, we have to have both F_1 and F_2 .

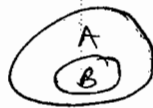
To identify F_2 , the following methods can be used.

(i) Inference rules,

(ii) closure set of attributes.

Inference rules -

Reflexive rule: If B is a subset of A then, $A \rightarrow B$.



$A \rightarrow B$

Transitive rule: If $A \rightarrow B$ and $B \rightarrow C$, then we can say $A \rightarrow C$.

Decomposition rule: If $A \rightarrow BC$, then $A \rightarrow B$ and $A \rightarrow C$.

Augmentation rule: If $A \rightarrow B$

$\Downarrow$

$AC \rightarrow BC$

Union rule: If $A \rightarrow B$ and $A \rightarrow C$ then $A \rightarrow BC$.
(applicable only if LHS are same).

Composition rule: If $A \rightarrow B$, and $C \rightarrow D$ then $AC \rightarrow BD$

Self determination: $A \rightarrow A$.

Ex 1 By semantics

$$F_1: \begin{aligned} A &\rightarrow B \\ B &\rightarrow C \\ C &\rightarrow D \\ D &\rightarrow E \\ E &\rightarrow F \\ F &\rightarrow G \\ G &\rightarrow H \end{aligned}$$

By using inference rules

F2:

By transitivity

$A \rightarrow B, B \rightarrow C \} A \rightarrow C$

$\begin{aligned} B &\rightarrow C \\ C &\rightarrow D \end{aligned} \} B \rightarrow D$

By union rule:

$\begin{aligned} D &\rightarrow E \\ D &\rightarrow H \end{aligned} \} D \rightarrow EH$

$B \rightarrow EH$ (if we mean
it it will cause
problems).

~~inference~~

Even though this method is very powerful, it is time consuming and error prone.

$\therefore$ we need an alternative.

Closure Set of attributes:

Algorithm to identify the closure sets:

Step 1: Let X be a set of attributes that will become the closure finally a closure.

Step 2: Repeatedly search for some FD such that LHS is available in X , then add RHS of that FD to X . (if it not available)

Step 3: Repeat Step 2 as many times as necessary until no more attributes can be added to X .

Step 4: The Set ' X ' after no more attributes can be added will become the closure of the attributes.

$\rightarrow A \rightarrow B$
 $BC \rightarrow DE$
 $AEG \rightarrow G$
 Compute AC^+

$$\begin{aligned}
 X &= AC \\
 &= ABC \\
 &= ABCDE
 \end{aligned}$$

$$AC^+ = ABCDE$$

$\rightarrow A \rightarrow BC$
 $CD \rightarrow E$
 $B \rightarrow D$
 $E \rightarrow A$
 Compute B^+

$$\begin{aligned}
 X &= B \\
 &= BD
 \end{aligned}$$

$$B^+ = BD$$

$\rightarrow AB \rightarrow C$
 $BC \rightarrow AD$
 $D \rightarrow E$
 $CF \rightarrow B$
 Compute AB^+

$$\begin{aligned}
 X &= AB \\
 &= ABC \\
 &= ABCD \\
 &= ABCDE
 \end{aligned}$$

$$\therefore AB^+ = ABCDE$$

→ The application of a closure set of attributes:

- (i) It is used to identify the additional FD's.
- (ii) It is used to identify keys (Candidate keys or primary keys)
- (iii) It is used to identify equivalences of FD's.
- (iv) It is used to identify irreducible set of FD's & canonical power of FD's.

(i)	F_1	F_2
	$AB \rightarrow C$	$AB \rightarrow DE$
	$BC \rightarrow AD$	
	$D \rightarrow E$	$AB^+ = ABCDE$
	$CF \rightarrow B$	

Identification of additional FD's by any closure set:-

→ To determine whether a particular FD is true or not, complete AB^+ from the original FD's. If AB^+ includes 'DE', then $AB \rightarrow DE$ is true, and will become additional FD.

→ Repeatedly check all additional FD's to find out F_2 set.

F_1
$A \rightarrow BC$
$CD \rightarrow E$
$E \rightarrow C$
$D \rightarrow AEH$
$ABH \rightarrow BD$
$DH \rightarrow BC$

Find whether $BCD \rightarrow H$ or not.
 $BCD^+ = BCD$
 $= BCDE$
 $= ABCDEH$
 $\therefore$ It is true.

whether $AB \rightarrow B$: $AED \rightarrow C$

~~$AB \rightarrow B$~~

$AED^+ = AED$

$= ABCDE$

Identification of keys with closure set of attributes:-

Different keys:

- (i) Super key
- (ii) Primary key
- (iii) Candidate key
- (iv) Foreign key
- (v) Surrogate key
- (vi) Composite primary key

Primary key (not null)

RNO Name age

Candidate keys

RNO ENO Type Xom work.
select any one as primary key.

Student table

RNO	ENO	I/date	R/date	Fine	S. Key
1	101	1/1	1/2		1
1	102	"	"		2
2	101	1/2	1/3		3
1	101	1/4	1/5		4

Surrogate key.
It is a system generated key.

if ~~I/date~~ one cannot take and return on the same day
then RNO ENO Idate can well form Composite key.

RNO	Name	age	Dept
1	a	16	101
2	b	17	101
3	c	15	101
		102	
		103	
		102	
		102	

Dept	Name
101	CSE
102	IT
103	EE

Foreign key is used for referential integrity.
 always Foreign key refers a primary key of some table.

FNO	name	age	deptno	submits to
-----	------	-----	--------	------------

There can be any no of Foreign keys in a table.

pkey	fkey
A	B
1	2
2	1
3	2
4	2
5	4
6	5
7	6

Superkey: Any Superset of primary key is a super key.

A	B	C	D	E

$A \rightarrow B$
 $B \rightarrow DE$
 $BC \rightarrow CE$

write all LHS attributes

whether $AB \rightarrow C$ $AED \rightarrow C$

~~$AB \rightarrow C$~~

$$AED^+ = AED$$

$$= ABCDE$$

Identification of keys with closure rule of attributes:-

Different keys:

- (i) Super key
- (ii) Primary key
- (iii) Candidate key
- (iv) Foreign key
- (v) Surrogate key
- (vi) Composite primary key

Primary key (not null)

RNO Name age

Candidate keys

VNO ENO Type Xom work.
select any one as primary key.

Student table

Book table

RNO	BNO	I/date	R/date	Fine	S. Key
1	101	1/1	1/2		1
1	102	"	"		2
2	101	1/2	1/3		3
1	101	1/4	1/5		4
					5

Surrogate key.
It is a system generated key.

of ~~I/date~~ one cannot take and return on the same day
then RNO BNO Idate can well form Composite key.

RNO	Name	age	Dept
1	a	16	101
2	b	17	101
3	c	15	101
		102	
		103	
		102	
		102	

Dept	Name
101	CSE
102	IT
103	EE

Foreign key is used for referential integrity.
 always foreign key refers a primary key of some table.

FNO	name	age	deptno	refers to

There can be any no of Foreign keys in a table.

PK	FK
A	B
1	2
2	1
3	2
4	2
5	4
6	5
7	6

Superkey: Any Superset of primary key is a super key.

A	B	C	D	E

A → B
 B → DE
 BC → CE

write all LHS attributes

- If we combine all the LHS attributes of the functional dependencies, then we will get a Super Key.
- Key for a table must be a subset of the Super Key.

Identification of keys with the help of closure set

- Compute closure for the attributes or combination of attributes on the LHS of functional dependencies.
- If any closure includes all the attributes, then that can be declared as a key for the table.

- - $A \rightarrow BC$
 - $CO \rightarrow E$
 - $E \rightarrow C$
 - $D \rightarrow AEH$
 - $ABH \rightarrow BD$
 - $DH \rightarrow BC$

Find the keys from the dependencies!

$$A^+ = ABC$$

$$E^+ = EC$$

$$D^+ = ADEH$$

$$= ACDEH$$

$$= ABCDEH$$

∴ D can be a key. Then we can have many superkeys

AD, BD, CD, —

$\rightarrow A \rightarrow B$
 $BC \rightarrow E$
 $ED \rightarrow A$

$$A^+ = AB$$

$$BC^+ = BCE$$

$$ED^+ = EDA$$

$$= EDAB$$

$$COE^+ = CDE$$

$$= ABCE$$

Candidate

$$ACD^+ = ACD$$

$$= ABCD$$

$$= ABCDE$$

$$BCD^+ = BCD$$

$$= BCDE$$

$$= ABCDE$$

$\rightarrow R = ABCDE$

$AB \rightarrow C$

$CD \rightarrow E$

$DE \rightarrow B$

$$AB^+ = ABC$$

$$DE^+ = DE$$

$$CD^+ = CDE$$

$$= CBDE$$

ABD

ADE

ACD

$$ACD^+ = ABCDE$$

$\rightarrow R = ABCDEFGHIJ$

$AB \rightarrow E$

$A \rightarrow OE$

$B \rightarrow F$

$F \rightarrow GH$

$AD \rightarrow IJ$

$$A^+ = ADE$$

$$= ADEIJ$$

$$B^+ = BF$$

$$= BFGH$$

$$D^+ = DIJ$$

$$AB^+ = ABC$$

$$= ABCDEFGHIJ$$

$\rightarrow 15)$

Ans:-
 17)

$\rightarrow$

$\rightarrow$ Eq

Ex 10)
 F=

Tak

→ (16) $R = ABCDEFGHIJ$

$AB \rightarrow C$

$BD \rightarrow EF$

$A \rightarrow GH$

$A \rightarrow \frac{I}{J}$

$H \rightarrow J$

ABD

Ans:-

(17) AB
BC
BD

(18) A
CD

(19) ABCD

→ $R = ABCDEFG$ marked as FD.

$A \rightarrow BC$

$B \rightarrow DE$

$C \rightarrow BC$

A is primary key according to attributes of F
now AFG will be the actual key of the table.

→ Equivalences of functional dependencies:-

Ex1

(10)

$F: A \rightarrow C$

$AC \rightarrow D$

$E \rightarrow AD$

$E \rightarrow H$

Req

$G: A \rightarrow CD$

$E \rightarrow AH$

Take set F and check 'G' is covered or not

$A^+ = A$
 $= AC$
 $= ADD$

$E^+ = E$
 $= EAD$
 $= EADH$

both $A \rightarrow CD$ and $E \rightarrow AH$ are covered.

∴ G can be derived from 'F'; Hence G is covered by F.

Ans:-

Step 1: Take $\text{Set } G'$ and check whether F is covered from G' .

$$A^+ = ACD \quad (A \rightarrow C)$$

$$AC^+ = ACD \quad (AC \rightarrow D)$$

$$\begin{aligned} E^+ &= E \\ &= AH \\ &= AHCD \quad (E \rightarrow AD, E \rightarrow H) \end{aligned}$$

Since F is ^{covered by} ~~equivalent~~ G' $\rightarrow$ ②.
From ① and ②, F and G are equivalent.
 $\therefore G$ is preferred one.

$$F_1 : G : \begin{aligned} A &\rightarrow CD \\ E &\rightarrow AH \end{aligned}$$

E and F_2 $F_1 + F_2 \rightarrow$ normalize.

$$\rightarrow F : \begin{aligned} B &\rightarrow CD \\ AD &\rightarrow E \\ B &\rightarrow A \end{aligned}$$

$$G : \begin{aligned} B &\rightarrow CDE \\ B &\rightarrow ABC \\ AD &\rightarrow E \end{aligned}$$

$$\begin{aligned} \text{Take } F : \quad B^+ &= AB CDE \quad (B \rightarrow CDE, B \rightarrow ABC) \\ AD^+ &= ADE \quad (AD \rightarrow E) \end{aligned}$$

$$\begin{aligned} \text{Take } G : \quad B^+ &= AB CDE \quad (B \rightarrow CD) \\ AD &= ADE \quad (\end{aligned}$$

Prefer : F .

$\rightarrow$ the
d
n
f

one
func
neo
these

a
An
step

step I
d
al

step II
FI

$\rightarrow$ Find
(i) ① A
②
③
④ A
⑤ B
⑥ C
⑦ C
⑧

covered

→ Irreducible set of FDs & minimal set of FDs
in standard form of FDs in canonical form with
no redundancies:

From remaining F_1 then F_2 then normalization.

once F_1 and F_2 are identified, some of the
functional dependencies among them may have
redundant attributes either on LHS or on RHS.
These attributes must be eliminated before making
a move to the normalization process.

And it is a four step procedure:

Step I: have single attributes on RHS

Step II: evaluate attributes on LHS by using some
of the irreducible premises from evaluation
attributes on RHS by any closure set of attributes

Step III: Apply the union rule to get the original
FDs.

→ Find the irreducible set from the following FDs:

- (i) ① $AB \rightarrow C$
② $C \rightarrow A$
③ $BC \rightarrow D$
④ $ACD \rightarrow B$
⑤ $BE \rightarrow C$
⑥ $CE \rightarrow AF$
⑦ $CF \rightarrow BD$
⑧ $D \rightarrow EF$

Step I:

$AB \rightarrow C$
 $C \rightarrow A$
 $BC \rightarrow D$
 $ACD \rightarrow B$
 $BE \rightarrow C$
 $\times$ $CE \rightarrow A$
 $CE \rightarrow F$
 $CF \rightarrow B$
 $CF \rightarrow D$
 $D \rightarrow E$
 $D \rightarrow F$

$C \rightarrow A$
By augmentation rule
 $CE \rightarrow AE$
By decomposition
 $CE \rightarrow A$
 $CE \rightarrow E$
1. $CE \rightarrow A$ is redundant
in given set, because
 $C \rightarrow A \Rightarrow CE \rightarrow A$

$$CF \rightarrow B \xRightarrow{\text{augmentation}} CF \rightarrow BC$$
 ~~$CF \rightarrow B$~~
 ~~$CF \rightarrow B$~~

Remo

Rem

$$BC \rightarrow D \quad \therefore \quad CF \rightarrow D$$

$CF \rightarrow D$ is redundant

Re

But using inference rules
if this is the way we go, it is a time consuming
process and error prone.

$\therefore$ Instead of using inference rules, we can make use of
closure sets of attributes to identify the duplicates.

Rem

Ex: 20:

$$\begin{aligned} A &\rightarrow B \\ C &\rightarrow B \\ D &\rightarrow ABC \\ AC &\rightarrow D \end{aligned}$$

(There can be several minimal covers for
a set of dependencies).

(Fourth - multivalued dependency
fifth - join dependency)

Remo

Step 1:

$$\begin{aligned} ① & A \rightarrow B \\ ② & C \rightarrow B \\ ③ & D \rightarrow A \\ \times ④ & D \rightarrow B \\ ⑤ & D \rightarrow C \\ ⑥ & AC \rightarrow D \end{aligned}$$

Step

Step 2:

Remove $A \rightarrow B$ and Compute A^+ from the ~~rem~~ ②, ③,
④, ⑤, ⑥ if we get B we can say that A determines
 B without having FD $A \rightarrow B$ and so $A \rightarrow B$ otherwise
it is not redundant.

C^+

$A^+ =$

$$A^+ = A$$

$\therefore A \rightarrow B$ is not redundant.

Remove $C \rightarrow B$ and compute C^+ from ① ③ ④ ⑤ ⑥

$C^+ = C$. It is not redundant.

Remove $D \rightarrow A$ and compute D^+ from ① ② ④ ⑤ ⑥

$$D^+ = DBC$$

It is not redundant.

Remove $D \rightarrow B$ and compute D^+ from ① ② ③ ⑤ ⑥.

$$D^+ = DACB$$

$B \Rightarrow D \rightarrow B$ is redundant.

Remove $D \rightarrow C$ and compute D^+ from ① ② ③ ⑤ ⑥.

$$D^+ = DAB$$

It is not redundant.

Remove $AC \rightarrow D$ and compute AC^+ from ① ② ③ ⑤ ⑥

$$AC^+ = AC$$

$$= ABC$$

It is not redundant.

~~Step 3~~
Step ③

$$A \rightarrow B$$

$$C \rightarrow B$$

$$D \rightarrow A$$

$$D \rightarrow C$$

$$AC \rightarrow D$$

$$C^+ = DCB$$

$$A^+ = AB$$

remove AC in $AC \rightarrow D$

$$A \rightarrow B$$

$$C \rightarrow B$$

$$D \rightarrow A$$

$$D \rightarrow C$$

$$C \rightarrow D$$

$$C^+ = CD$$

$$= ACD$$

$$= ABCD$$

x

Remove C in FD

$$A \rightarrow B$$

$$C \rightarrow B$$

$$D \rightarrow A$$

$$D \rightarrow C$$

$$A \rightarrow D$$

$$A^+ = ABD = ABCD$$

x

we need both hands on LHS.

assuming

the use of
ites.

in ②, ③,

unifies

else

step 4:-

$A \rightarrow B$
 $C \rightarrow B$
 ~~$A \rightarrow A$~~
 $D \rightarrow A$
 $D \rightarrow C$
 $AC \rightarrow D$

apply union for the common attributes

$A \rightarrow B$
 $C \rightarrow B$
 $D \rightarrow AC$
 $AC \rightarrow D$

Ex 121

$AB \rightarrow C$
 $C \rightarrow B$
 $A \rightarrow B$

Step 1:-

① $AB \rightarrow C$
 ② $C \rightarrow B$
 ③ $A \rightarrow B$

Step 2:-

① $AB \rightarrow C$
 ② $C \rightarrow B$
 ③ $A \rightarrow B$

Remove ① and compute AB^+ from 2 and 3.

$$AB^+ = AB$$

It is not redundant.

Remove ② and compute C^+ from ① and ③.

$$C^+ = C.$$

It is not redundant.

Remove ③ and compute A^+ from ① and ②.

$$A^+ = A.$$

It is not redundant.

step 2

step 1

22)

$\begin{array}{r} 5 \\ \times 1) 1 \\ \times 2) 1 \\ \times 3) 1 \\ \times 4) 1 \\ \times 5) 1 \\ \times 6) 1 \\ \times 7) 1 \end{array}$

Step 3:

$AB \rightarrow C$
 $C \rightarrow B$
 $A \rightarrow B$

$A^+ = ABC$
 $B^+ = B$

Remove A

$B \rightarrow C$
 $C \rightarrow B$
 $A \rightarrow B$

$B^+ = BC$
 not equal

Remove B

$A \rightarrow C$
 $C \rightarrow B$
 $A \rightarrow B$

~~Compute~~ $A^+ = ABC$

equal
 so we can remove B.

Step 4:

$A \rightarrow C$
 $C \rightarrow B$
 $A \rightarrow B$

Transitive
 $\Rightarrow$

$A \rightarrow C$
 $C \rightarrow B$

is irreducible set.

22) FD's are $ABD \rightarrow AC$
 $C \rightarrow BE$
 $AD \rightarrow BF$
 $B \rightarrow E$
 Find minimal set.

Ans

$AD \rightarrow CF$
 $C \rightarrow B$
 $B \rightarrow E$

Step 1:

- X 1) $ABD \rightarrow A$
- ✓ 2) $ABD \rightarrow C$
- ✓ 3) $C \rightarrow B$
- X 4) $C \rightarrow E$
- ✓ 5) $AD \rightarrow B$
- ✓ 6) $AD \rightarrow F$
- X 7) $B \rightarrow E$

Step 2:

Remove A ①

$ABD^+ = ABD$
~~AB~~

Step 3:

$ABD \rightarrow C$
 $C \rightarrow B$
 ~~$AD \rightarrow B$~~
 $AD \rightarrow F$
 $B \rightarrow E$

Remove ②

$BD^+ = BDE$
 $AD^+ = ADBFEC$
 $AB^+ = ABE$

Remove A X

$BD^+ = BDE$

B ✓

$AD^+ = ADBFEC$

D X

$AB^+ = ABCE$

Remove B from 1 FD.

- 1) $AD \rightarrow C$
- 2) $C \rightarrow B$
- 3) $AD \rightarrow B$
- 4) $AD \rightarrow F$
- 5) $B \rightarrow E$



- $AD \rightarrow C$
- $C \rightarrow B$
- $AD \rightarrow F$
- $B \rightarrow E$

$D^+ = D$
 $A^+ = A$



step IV

- $AD \rightarrow C$
- $C \rightarrow B$
- $AD \rightarrow F$
- $B \rightarrow E$



- $AD \rightarrow CF$
- $C \rightarrow B$
- $B \rightarrow E$



~~$AD \rightarrow CF$~~
 ~~$C \rightarrow B$~~
 ~~$B \rightarrow E$~~

- remove ① $AD^+ = ADBFE$
 remove ② $C^+ = C$
 remove ③ $AD^+ = ADCFB$
 remove ④ $AD^+ =$

remove A X
 $D^+ = DCB$
 not equal
 remove D X
 $A^+ = ACB$
 not equal

remove A X
 $D^+ = BDF$
 remove D X
 $A^+ = AF$

23) (iii)

- $A \rightarrow BC$
- $AE \rightarrow H$
- $C \rightarrow D$
- $D \rightarrow G$
- $E \rightarrow F$

24) (iii)

- $BC \rightarrow AE$
- $C \rightarrow D$
- $A \rightarrow F$
- $F \rightarrow G$

cl

(i)

(ii) I

(iii) 4

cl

1

2

3

4

5

on

re

Par

A

ac

b

bc

Trans

1

re

d

Classification of FD's :

(i) Partial FD :

(ii) Transitive FD

(iii) Full FD

Classification of Dependencies :

1) Functional dependency. (2NF, 3NF, BCNF)

2) multivalued dependency (4NF)

3) Join dependency (5NF, PJNF)

4) Template dependency

5) Inclusion dependency (dependencies among the data bases).

Once ER diagrams are converted to tables 4NF and 5NF are not needed.

Partial FD

A FD in which one or more non key attributes are functionally depending on a part of the primary key.

Ex

$R = ABCD$

$AB \rightarrow C$

$B \rightarrow D$

Key = AB

$\therefore B \rightarrow D$ is a partial dependency.

Transitive dependency

A dependency in which there is a relationship among non key attributes i.e. one non key attribute is determined by another non key attribute.

Ex-

$R = ABCD$

$AB \rightarrow C$

$B \rightarrow D$ (PD)

$C \rightarrow D$ (TD)

25)

Full functional dependency

A functional dependency $X \rightarrow Y$ is a full functional dependency, if removal of any attribute from X the dependency does not hold any more.

Ex: If $AD \rightarrow C$ is full functional dependency, then we cannot remove A or D . i.e. C is fully dependent on AD . If we are able to remove A or D , then it is not full functional dependency.

Normalization

It is the process of evaluating logical design by applying the series of rules of normalization as a formal process based on primary key and FD's to decide what attributes should be grouped together in different tables.

1st NF: All columns in table with single value, no multiple values, all entries in col of same kind.

2nd normal form:

A table is said to be in the second normal form, if it is already in the 1st normal form and it is free from partial functional dependencies.

The table is automatically in second normal form.

If primary key consists only one attribute or all attributes one part of primary key of table consists only 2 attributes.

or

Identify

Identify

1)

2)

3)

4)

Find

If

+

+

a

25) $R = ABCDEFGHIJ$

$AB \rightarrow C$

$A \rightarrow DE$

$B \rightarrow F$

$F \rightarrow GH$

$D \rightarrow IJ$

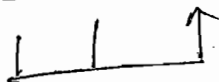
Decompose the table upto II normal form.

Identify Key: $AB^+ = ABCDEFGHIJ$

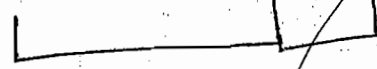
Identify partial dependency

A B C D E F G H I J

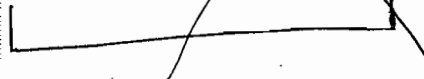
1)



2)



3)



4)

Confusing and error prone.

Find $A^+ = ADEIJ$
 $B^+ = BFGH$

~~A B C D E F G H I J~~

~~A B C D E F G H I J~~

$R_1 = ADEIJ$

$R_1 = ADEIJ$

$R_2 = BFGH$

$R_2 = BFGH$

$R_3 = ABC$

$R_3 = ABC$

If there is any partial dependency, remove partially dependent attributes from original table place them in a separate table along with a copy of its determinant.

26) $\begin{matrix} A & B & C & D & E \\ \text{Book}(\text{Book_title}, \text{Author_name}, \text{book_type}, \text{Price}, \text{author_affiliation}, \\ & \text{Publisher}) \end{matrix}$

$$A \rightarrow FC$$

$$C \rightarrow D$$

$$B \rightarrow E$$

decompose upto II normal form.

Identify the key:

$$AB^+ = ABCDEF$$

$\therefore AB$ is key.

$$A^+ = ACDF$$

$$B^+ = BE$$

$$AB \not\rightarrow C, D, E, F$$

$R_1: \underline{A}CDF$ (in 2nd normal form because only one key attribute)

$R_2: \underline{B}E$ (only two attributes)

$R_3: \underline{AB}$ (all are key attributes)

27) $R = ABCDE$ Primary Key = AC.

$$B \rightarrow E$$

$$C \rightarrow D$$

$$A \rightarrow B$$

decompose upto II normalization.

$$A^+ = ABE$$

$$C^+ = CD$$

$R_1: \underline{A}BE$

$R_2: \underline{C}D$

$R_3: AC$

is
rather difficult

$R = ABCDEFGHIJ$

- $AB \rightarrow C$
- $BD \rightarrow EF$
- $AD \rightarrow GH$
- $A \rightarrow I$
- $H \rightarrow J$

Primary Key: ABD

~~$A^+ = AI$~~
 ~~$B^+ = B$~~
 ~~$D^+ = D$~~

$AB^+ = ABCI$
 $ABD^+ = BDEF$
 $AD^+ = ADGHJI$

- ~~R_1~~ : $ABCI$
- ~~R_2~~ : $BDEF$
- ~~R_3~~ : $ADGHJI$
- ABD (II normal form because all attributes in primary key)

key factor)

$ABCI$
 $\wedge$
 $A^+ = AI, B^+ = B$
 AI ABC

$BDEF$
 $\wedge$
 $B^+ = B, D^+ = D$
no partial dependencies
 $BDEF$

$ADGHJI$
 $\wedge$
 $A^+ = AI, D^+ = D$
 AI $ADGHJ$

ABD

<u>AI</u>
<u>ABC</u>
<u>$BDEF$</u>
<u>$ADGHJ$</u>
<u>ABD</u>

III normal form:-

26)

→ A table is said to be in III normal form, if it is already in the second normal form and free from transitive dependencies.

→ A table is automatically in III normal form if one of following hold:

- (i) If table consists of two attributes.
- (ii) If II normal form table consists only one non key attributes.

25) ABCDEFGHIJ

$AB \rightarrow C$

$A \rightarrow DE$

$B \rightarrow F$

$F \rightarrow GH$

$D \rightarrow IJ$

II normal

$R_1: \underline{A} DEIJ$

$R_2: \underline{B} FGH$

$R_3: \underline{ABC}$

$\underline{DIJ} \quad \underline{ADE}$
 $\underline{FGH} \quad \underline{BF}$

$\underline{A} DEIJ$

$D^+ = DIJ$

$E^+ = E$

$I^+ = I$

$J^+ = J$

$\underline{B} FGH$

$F^+ = FGH$

$G^+ = G$

$H^+ = H$

$\underline{DIJ}$
 $\underline{ADE}$

$\underline{EGH}$
 $\underline{BF}$

$\underline{ABC}$

$C^+ = C$

$\underline{ABC}$

$R_1: \underline{DIJ}$

$R_2: \underline{ADE}$

$R_3: \underline{FGH}$

$R_4: \underline{BF}$

$R_5: \underline{ABC}$

↑
3NF

27)

26)

ABCDEF

$A \rightarrow FC$

$C \rightarrow D$

$B \rightarrow E$

II normal table:

$R_1: \underline{A} C D F$

$R_2: \underline{B} E$

$R_2: \underline{A B}$

III normal form because only two attributes.

$\underline{A} C D F$

$C^+ = C D$

$D^+ = D$

$F^+ = F$

<u>C</u> D
<u>A</u> C F

<u>C</u> D
<u>A</u> C F
<u>B</u> E
<u>A B</u>

III normalized tables.

27)

$R = A B C D E$

$B \rightarrow E$

$C \rightarrow D$

$A \rightarrow B$

$R_1: \underline{A} B E$

$R_2: \underline{C} D$

$R_2: \underline{A C}$

$\begin{matrix} \underline{B E} \\ \underline{A B} \end{matrix}$

II normal

ABC

DIJ

ADE

FGH

BF

ABC

28) R = ABCDEFGHIJ

FDs :
 $AB \rightarrow C$
 $BD \rightarrow EF$
 $AD \rightarrow GH$
 $A \rightarrow I$
 $H \rightarrow J$

2NF

AI ——— AI
ABC ——— ABC
ODEF ———
ADGHJ
ABD

29) a) $B \rightarrow C$
 $B \rightarrow D$
 $A \rightarrow C$
 $A \rightarrow D$
 b) $C \rightarrow D$ (1)
 $D \rightarrow C$

→ If $x \rightarrow A$ is a dependency, then the table is in the III normal form if one of the following conditions exists:

- (i) If x is a super key (because if x is should not be non key attribute and not also not be part of a key, if it is a part, then it is partial dependency).
- (ii) If A is a part of key.

→ If ~~not~~ $x \rightarrow A$ is a dependency, then the table is said to be not in III normal form if the following

- (i) if x is a proper subset of some key (partial dependency)
- (ii) if x is not a proper subset of key. (non key)

→ De

9

7

at
 (d
 in
 de

(i)

(ii)

(i) loss

→

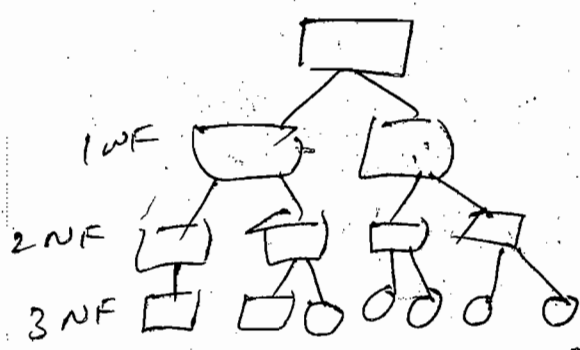
A
a ₁
a ₂
a ₃

A
 diff
 an
 if
 of

either

su
 1

→ dependency preserving and lossless join properties
 If we apply the normal forms, to the universal table, it may be splitted up into different fragments.



at any stage, if we combine the fragments (denormalization) it should give the original table in terms of columns and rows, and it will be described as the following properties:

- (i) lossless join property - (mandatory)
- (ii) dependency preserving property (optional)

(i) lossless :-

→

R		
A	B	C
a ₁	b ₁	c ₁
a ₂	b ₂	c ₂
a ₃	b ₁	c ₃

R₁

A	B
a ₁	b ₁
a ₂	b ₂
a ₃	b ₁

R₂

B	C
b ₁	c ₁
b ₂	c ₂
b ₁	c ₃

A decomposition of a relational schema R into different relations like R_i where i=1, 2, ..., n and it is said to be lossless join decomposition if natural joins of all decompositions gives the original relation.

$$R = \pi_{R_1}(R) \bowtie \pi_{R_2}(R) \dots \dots \pi_{R_n}(R) \rightarrow \text{①}$$

or R = $\pi_{R_1}(R) \bowtie \pi_{R_2}(R)$

Suppose the natural join of all decompositions look like $R \subseteq \pi_{R_1}(R) \bowtie \pi_{R_2}(R)$, we have spurious tuples.

non key part of a is partial
 key
 (partial key)

$$\pi_{R_1}^{(R)} \bowtie \pi_{R_2}^{(R)} =$$

A	B	C
a ₁	b ₁	c ₁
a ₂	b ₂	c ₂
a ₃	b ₁	c ₃
a ₁	b ₁	c ₃
a ₃	b ₁	c ₁

Formula ① is time consuming and even people do use another simple checking. It is as follows.

$$R_1 \cap R_2 \rightarrow R_1 \text{ all can now determine losses}$$

If R is decomposed into R_1 and R_2 , then we don't have any losses if

$$(R_1 \cap R_2) \rightarrow R_1$$

$$(R_1 \cap R_2) \rightarrow R_2$$

i.e. common attribute is a key for atleast one table.

Ex: $R = ABC$

$$A \rightarrow B$$

$$R_1 (A, B)$$

$$R_2 (A, C)$$

Check which decomposition is lossless & lossy.

A is primary key for R which is common attribute and so it is lossless.

Defend

Defend

if b

then

FO

5

Oct

note

for

at

to

7

BCNF

A table is said to be in BCNF, if it is already in 3NF and every FD will have a candidate key on the LHS & all determinants are the candidate keys.

The differences b/w 3NF and BCNF:

3NF

- 1) It concentrates on primary key
- 2) Redundancy is high compared to BCNF
- 3) It may preserve all the dependencies
- 4) A dependency $x \rightarrow y$ is allowed in 3NF if x is a super key or y is a part of some key.

BCNF

- 1) It concentrates on candidate keys.
- 2) Redundancy is low compared to 3NF.
- 3) It may not preserve the dependencies.
- 4) A dependency $x \rightarrow y$ is allowed if x is a superkey.

→ The Third normal form may not deal the following cases in a proper way.

- (i) If table consists two or more candidate keys
- (ii) If candidate keys overlap, i.e. contain one common attribute among the keys.

31) 1

2)

last
3

31) 1) $C \rightarrow D, C \rightarrow A, B \rightarrow C$

a) Key: B.

B A C D

b) 2NF

c) $\left. \begin{matrix} C \rightarrow D \\ C \rightarrow A \end{matrix} \right\} T(D).$

3NF Tables: C D

C A

B C.

all Determinants are keys $\therefore$ it is in BCNF.

2) $B \rightarrow C, D \rightarrow A$

a) Key: B D

b) B D A C INF

Partial dependencies are there $\therefore$ it is only in INF.

2NF: B D B C D A.

it is also in 3NF

all determinants are super key $\therefore$ it is in BCNF.

$ABC \rightarrow D, D \rightarrow A$

$ABC^+ = ABCD$

BCD D A

a) Key: ABC and BCD

b) ABC D, ~~it is in 3NF~~ it is in 3NF.

c) it is in BCNF because D is not a super key.

ABCD, A D

even though this decomposition satisfies BCNF it causes lot of redundancies.

normalization will effectively identify redundancies within the table. but it may not identify among the tables.

→ Identification of Redundancy among the tables

non key attributes should not be repeated in the decomposed tables.

H) $A \rightarrow B$,
 $BC \rightarrow D$ 2NF.
 $A \rightarrow C$

a) Key: A

b) A B C D
 2NF

c) A B C, BC D

A and BC are determinants and are superkey
 ∴ this is in BCNF.

5) $AB \rightarrow C$
 $AB \rightarrow D$
 $C \rightarrow A$
 $D \rightarrow B$

a) Key: (AB), CD, BC, AD

b) AB C D
 3NF

determinator AB, is superkey but C, D are not.
 ∴ this is not in BCNF.

~~ABCD, CA, DB~~

$C \rightarrow A$ and $D \rightarrow B$ are not allowed on ABCD

we among
 32) BCNF as per the rules.
 To satisfy BCNF, we can have
AB CD, CA, DB as tables
 but it will have lot of redundancy among the
 tables, $\therefore$ it causes problems.

32) $R = ABCDEFGH$

- $AB \rightarrow CEFGH$
- $A \rightarrow D$ (FD)
- $F \rightarrow G$ (FD)
- $FB \rightarrow H$
- $HBC \rightarrow ADEFG$
- $FBC \rightarrow ADE$

AB is key

superkey

2nd normal form

$R_1: \underline{AD}$

$R_2: \underline{ABCEFGH}$

3rd normal form

$R_1: \underline{AD}$
 $A \rightarrow D$

$R_2: \underline{ABCEFGH}$
 $AB \rightarrow CEFGH$
 $FB \rightarrow H$ (BCNF)
 ~~$HBC \rightarrow ADEFG$~~

$R_3: \underline{EG}$
 $F \rightarrow G$

BCNF

$R_1: \underline{AD}$

$R_2: \underline{ABCEFG}$

are not

on ABCD

35) 1) a) Key BD

b) BC and AD

Since there is no common attribute, the natural join of these two tables leads to multiple spurious tuples.

∴ It is not a good decomposition.

2) $AB \rightarrow C, C \rightarrow A, C \rightarrow D$; decompose into ACD, BC .

a) Keys: AB, BC, CB

b) (C is a key in ACD and so it is lossless decomposition)

There is a common attribute C and it can be a key in 1 table hence it is lossless, but it is not preserving FDs.

3) $A \rightarrow BC, C \rightarrow AD, ABC, AD$.

a) Key: A, C

b) The common attribute A can be key for both tables and so it is lossless.

4) lossless Key: A

5) lossy. Key: A

36) $ABCD E$; FDs: $AB \rightarrow CDE, A \rightarrow C, D \rightarrow E$.

Key: AB

II NF+

$AB \rightarrow CDE, A \rightarrow C, D \rightarrow E$

III NF

$AB, ABDE, AC, D \rightarrow E, A \rightarrow C$

It is upto BCNF.

37) 1.

II ND

III

It is

38)

The dim

→ not

am

→ not

is

→ the

is

Adva

→ It

m

→ mai

⇒ not

denormal

The ten

is R

we C

(1) of

(1) of

For

37) $R = ABCDE$

key: AB.

ABDE

AC

$A \rightarrow C$

III Normalization

ABE

ACD

$A^+ = ACD$

(II)

ABE

AC

CD

DE is in BCNE

38)

The limitations of normalization

- Normalization cannot identify the ~~limitations~~ redundancies among the tables, it can do it with in the tables.
- Normalization leads to the so many table fragments.
 - ∴ During the query retrieval from, the response of system is slow.
- The table fragments may not have the real world meaning.
 - ∴ difficult to understand and maintain the system.

Advantages

- It reduces anomalies like insertion, deletion and modification.
- maintaining data integrity is very easy
- Normalization increases the speed of query retrieval.

Denormalization

The process of combining multiple table fragments selectively is known as denormalization.

- We will go for denormalization under following circumstances
- If majority of the queries are data retrieval queries.
 - If a particular query requires join of multiple tables.
 - Ex: while preparing balance sheet.

- (iii) If update problems are not important but performance is important.

SQL

- 1) Select * from departments where deptno = 101 Group by AGE.
- 2) Select RNO, max(marks) from student.
- 3) select min(distinct AGE) from student.
- 4) select * from orders where max(price) > 1000
- 5) select idno, idname, idquantity from orders having count(idno) > 10.
- 6) select Rollno, name, Age from student order by 1, 2.
- 7) select Rollno, name, Age from student order by name & ascending age descending.
- 8) select name A, Age B, Rollno C from student having C > 15 & order by A, B.
- 9) select idno, Agent no, Agent name from orders where Customer name between C and B.
- 10) select * from students S1 where 2 = (select count(Rollno) from students S2 where S1.marks ≤ S2.marks).

Ans

- 1) Group Group by and '*' should not be used.
- 2) Add 'Group by Rollno'.
- 3) distinct is not necessary.
- 4) Aggregate cannot ~~appear~~ appear
- 5) Group by needed.
- 6) always cannot be used in having.
- 7) lower ~~value~~ value should be for first.
- 8) It gives the second max marks student.

Query

(i) To

Ex.

or

2) Gr

Ex.

3) min

Ex.

4) F

The go

SQL

SQL 1

SQL 2

SQL 3

informal

Classification of database languages

1) Transform Oriented languages

Ex: SQL.

Input are tables and o/p are tables.

2) Graphical Oriented languages

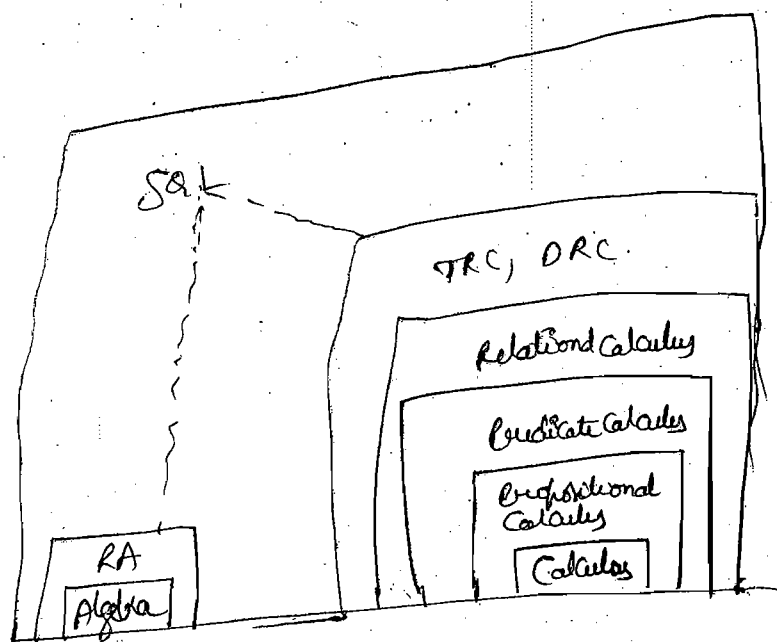
Ex: QBE.

3) menu driven & fourth generation languages.

Ex: commercial websites.

4) Fourth generation languages. (English like languages).

The roots of SQL:



SQL Standards:

SQL1 also SQL86

SQL2 also SQL92

SQL3 also SQL99

Components of SQL:

(i) DDL (ii) DML (iii) DCL.

where clause

Based on the filtering activity, the SQL queries are divided into 5 groups.

(i) Queries that uses basic search operations.

(ii) Queries that uses Range operations.

(iii) Queries that uses Comparison operations.

(iv) Queries that uses pattern matching.

(v) Queries that uses nulls and not nulls.

✶

Ex 1

(i) select * from student ~~from~~ ^{where} Roll no = 101.

(ii) select * from student where age > 50 and branch = IT

(iii) select * from student where age > 50 and branch = IT and year = 2.

Find the details of students who belongs to CSE, IT, ECE, EEE and MECH.

select * from student where branch ∈ (CSE ∪ ^{branch} IT ∪ ECE ∪ EEE ∪ MECH).

more efficient:

✓ select * from student where branch in (IT, CSE, EEE, ECE, MECH).

select * from student where branch not in (CSE).

Comparison

select * from student where age > 15;

select * from student where city = Hyderabad and branch = IT and year = 2 and age > 15.

Pattern m

% 95 d

- : 95

A%: 9

%A: 6

%A%:

%ACE%

A---

searching

search

→ Find it with

• select

→ Find all

select

→ Find the

IV b

select

→ Find

Comm

etc.

Pattern matching

% indicates sequence of zero or more characters.

- indicates a single character.

are

A%: First letter is A followed by any no. of characters.

%A: Last character is A preceded by 0 or more characters.

%A%: A character A may appear anywhere in the word.

%ACE%: A word 'ACE' may appear anywhere in the pattern.

A---: A four letter word starting with 'A'.

Searching special characters using escape symbol (\)

searching for % in any word.

% | % %

↓
Pattern we are searching for.

→ Find the names of student details whose name starts with 'S' and from branch 'CSE'.

Select * from student where name like 'S%' and branch = 'CSE'.

→ Find all the student details who opted network protocol subject.

Select * from student where subject like '% network protocol'.

→ Find the details of students whose 1st letter is 'A' and 4th letter is 'C' and last letter is 'D'.

Select * from student where name like 'A--C%D'.

→ Find the details of students who are in CSE branch coming from hyd and whose 3rd letter is 'A'.

Select * from student where name like '--A%' and city = 'hyd' and branch = 'CSE'.

IT

ITD

EEE and

DEED

ECE, MECH

branch IT

15

→ Range queries (b/w and not b/w operator).

→ Find details of Student whose marks b/w 500 and 600.

Select * from Student where marks between 500 and 600.

→ select * from Student where marks between 350 and 400.
104 and 107. (Here b/w is inclusive)

→ select * from Student where marks not between 350 and 400.
101, 102, 103, 105, 106, 108, 109, 110

→ Still select * from Student where name between 'A' and 'K'
'A' and 'K'.

103, 104, 108, 110, 111

111 K 19 M IT 3 500.

(Here b/w accept from A, AB, ABins, ..., K)
KA, KB, K- X are not included.

→ while executing the b/w query with characters, it
will take lower boundary and stops exactly at the
upper boundary.

∴ 102 and 106 will not appear as an o/p.

but they will appear in not b/w case.

→ If any name consists only a single letter 'K', that will
be included in the b/w case.

→ select * from Student where marks between 400 and 350
results in syntax error because lower bound mark
is mentioned first in the b/w query.

Null

→ Find

Select

→ Find the

address

Select

Agg

max

Star

→ Select

→ Select

max

non

→ Select

→ Select

Sum

qu

→ Select

→ Select

dis

will

→ Select

→ Select

→ Select

→ Select

Null queries:-

- Find the details of students who are free from fine.
- Select * from library where fine is null.
- Find the details of all students who are having email addresses
- Select * from students where email is not null.

Aggregate functions +

maximum, minimum, sum, average, count, count(*), standard deviation (stdev), variance (var)

- select max(marks) from student ✓
- select max(marks) from student ✓
max and min can be used on both numeric and non numeric columns.
- select sum(marks) from student X
- select sum(marks) from student ✓
sum and avg must be applied only on numeric queries.
- select distinct sum(distinct marks) from students.
- select min(distinct name) from students.

distinct has no effect with min and max, but will have a considerable impact on sum and average.

- select count(roll no) from students ✓
- select sum(marks * 10) from student ✓
- select avg(marks₁ + marks₂) from student ✓
- select avg(marks₁, marks₂) from student X

Aggregate functions operate on a single column of a table and on expressions ~~for~~ and returns a single value; but cannot work on multiple columns simultaneously.

Q: select * from Student where max(marks) = 500 X
 → select * from Student where order by (max(marks)) X
 → select * from Student Group by count (roll no)

aggregate functions can be used only in the select and having clauses, but not in where, order by and Group by clauses.

→ select roll no, max(marks) from Student X

If select clause consists of aggregate function along with ordinary column, then it must be associated with group by clause.

Group by:

→ Find the maximum marks among students?

Sol: select max(marks) from Student.

→ Find the details of student who got maximum marks.

select * from Student where marks = (select max(marks) from Student);

→ Find the student details who got second maximum & 3rd maximum & in general nth maximum?

(used to generate the group aggregate and works on non aggregate columns in the select columns).

<u>Table</u>	<u>roll no</u>	<u>marks</u>
	101	450
	102	300
	103	500
	104	480
	105	490
	106	300

Sol: select

Si
10
10
10

Q: Fi
less

As: d

Ex:

Two

Q: Fi
W
Side

Q:

Restor

1) of
col
in
of
of

Single
 turns
 = 500 X
 ac(marks) X
 (100) X

Group by
 7 of
 k de

rows.
 rows) from
 student);
 um
 7?
 up aggregate
 rate columns
)

select * from student S₁ where ~~marks~~
 (for second maximum) 2 = (select Count (S₂.marks) from student
 S₂ where S₁.marks <= S₂.marks)

S₁ (alias of students)

101	450
102	500
103	480

S₂ (another alias of student)

101	450
102	500
103	480

Q: Find the details of the student who got less than the average marks of the class.

Ans: select * from student where marks < (select Av (marks) from student)

Ex: Rollno:
 100
 100
 —
 200
 Count (*) returns '6' (SE counts null values)
 Count(Rollno) - 4 (SE counts the duplicate values).

Two ways for average calculations: Avg(marks), $\frac{\text{sum(marks)}}{\text{count(Rollno)}}$
 This is more accurate.

Q: Find the total no of students in each and every branch?

select Branch-id, Count(Roll no) from student group by Branch-id;

Q:

101	a	IT
102	b	CSE
103	c	I

Output:

IT	60
CSE	60
ECE	120
EEE	40

Restriction for Group By

1) If SELECT list consist a non-aggregate column with an aggregate function, then the query must be associated with a group by clause and it should consist all non aggregate columns in the group by list.

Ex: select age, bid, count (roll no) from student group by bid, age

4. The clause select

1) we cannot use a column in group by clause that is not available in select list.

select roll no, count (bid) from student group by age X

ORDER

2) we cannot use aggregate function in group by clause.

→ It is

any

→ It is a

group

Q: Find out the branch details whose total strength is more than 50 students.

select bid, count (roll no) ^{from student} group by bid having count (roll no) > 50

Ex: It

It: S

note: Group by is applied on independent rows. Having is applied among groups.

Ex: It

It

Q: Find the details of student who paid more than 20 rupees as a total fine?

Roll	Date	Fine
1	1-1-06	100
1	2-1-06	200
2		

select roll no, sum (fine) from library group by roll no having sum (fine) > 20

Ex: select roll

correct +

So

How to know the execution time of a query?

Ex: For

the d

at

select

to

1. before execution of a query at command prompt write "set timer on".

2. execute query

3. then at the command prompt write "set timer off"

note

Ex: It

It

where clause

having clause

1. It is used to filter rows
2. Aggregate functions cannot be used
3. It is mandatory (i.e. no alternative for where)
1. It is used to filter groups
2. Aggregate functions can be used
3. It is optional (i.e. the results given by having can also be achieved without having clause)

4. The columns appear in where clause need not appear in select list.

4. The columns in where clause must either be aggregated in select list or appear in group by clause.

ORDER BY

→ It is used to sort out the output in a particular way like descending order or in ascending order.
→ It is applied on independent columns but not on group of columns.

Ex:- select * from student order by name;

Ex:- select roll no, name, age from student where age > 15 order by (1, 2) → place order

Ex:- select roll no, name, age from student where age > 15 order by roll no descending, name ascending.

Ex:- select branch id, count (roll no) group by branch id order by count (roll no) X

Correct form:

select branch id, count (roll no) as group by branch id order by a;

Ex:- Find total strength of each branch and display the details of branches where the total no of students is more than 45 and display them in ascending order
select br id, count (roll no) as from student group by br having > 50 order by a;

Note aliases are not allowed in having clause.

Ex:- select roll no, max (marks) from student group by roll no having roll no in (104, 110, 115);

Sub Query :- A query embedded in another query is called sub query.

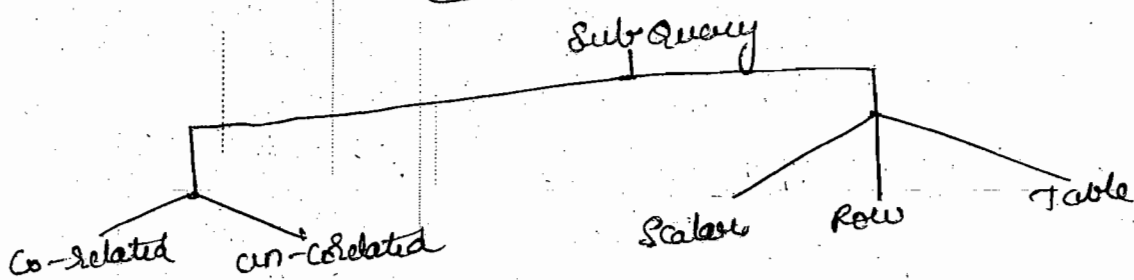
Characteristics

- The inner query will run first and substituted results in outer query.
- The variable from outer query can be used in the inner query but reverse is not true.
- The subquery can be used in the where and having clause and sometimes in from clause also.
- The inner query will run always only once. If it is independent of outer query, then it is called uncorrelated sub query.
- Inner query will run sometime as many times as the no of rows available in the outer query. If it refers a variable in outer query. Then it is called correlated sub query.

Ex: Select * from student where marks = (select max(marks) from student)
(un-correlated sub query).

select * from student S₁ where S₁ >= (select count(S₂. marks) from student S₂ where S₁. marks <= S₂. marks)
(correlated sub query).

Classification of Sub query



Scalar Sub Query:

It returns only one row and one column.

Row sub Query:

It returns multiple columns but only one row.

Table sub Query:

It returns multiple rows and multiple columns.

The form

→ By -

see 1

→ By -

→ By -

Ex -

1

Sol: 2

Ex: 3

2

Sol: 3

Ex: 4

3

4

Sol:

Ex: 5

5

6

Sol:

7

Ex: 6

8

9

The possible ways of writing Sub-queries:

- By using In-predicate (Between, like, is null cannot be used with the sub queries).
- By using quantifier Comparison predicates (Any, All)
- By using exist and not exist.

Ex determine the information about agents who ~~leave~~ lived in Bombay & Delhi.

Sol select * from agents where city in ('Bombay', 'Delhi');

Ex determine the details of customers who placed orders with agents in Bombay & Delhi?

Sol select Cid from orders where aid in (select aid from agents where city in ('Bombay', 'Delhi'));

Ex Find the complete details of customers who place orders with the agents who are living in Bombay & Delhi.

Sol select * from customers where cid in (select cid from orders where aid in (select aid from agents where city in ('Bombay', 'Delhi')));

Ex Find details of student from the IT and CE branch where student table having (Rollno, name, BTD) and Dept table having (Dept, BID, Bname).

Sol select * from student where BID in (select BID from dept where Bname in ('IT', 'CE'));

Ex Find out the total strength.

Sol select BID, count(Rollno) where BID in (select BID from dept where Bname in ('IT', 'CE')) group by BID having count(Rollno) > 60 order by 1.

select Customer id, count (Order id) where Order id from Orders where aid in (select aid from agents where rank < 26) group by customer id having count (Order id) < (select avg (Order id) from Orders)

→ Find
max

select

Restrictions for sub-queries

The following operators cannot be used b/w main query and sub query.

- b/w and not b/w op.
- like and not like op.
- is null, is not null.

The following operators are normally used b/w main query and sub query.

- in predicate: in & not in.
- quantified comparison predicate: Any, all, some.
- Exist predicate: Exist and not exist.

→ select
the

with

with

sel

Any, all, some:-

> any ()	
5 ✓	1
6 ✓	2
7 ✓	3
8 ✓	10
0 X	

> all ()	
5 X	1
6 X	2
7 X	3
8 X	10
20 ✓	

→ ider
Pol

sel

ik

sel

= any $\Leftrightarrow$ in

If subquery is prefixed with 'any' key word, then the condition will be true if it is satisfied by any of the values produced by inner queries.

If subquery is prefixed with 'all' key word, then the condition will be true if it is satisfied by all the values produced by inner queries.

- Show agents where count (order) is less than the CID value of customers with discount smaller than those of any customers who live in Bombay.

select CID from customer where discount < ~~any~~ ^{all} (select discount from customer where city = 'Bombay')

X	10	10
X	12	8
X	8	
X	8	
	0 ✓	

only one row.

query

→ select AID as value of agents who receive more than the minimum % commission.

without subquery

select aid from agents where percent > 5; (Hard coded query)

with subquery

select aid from agents where percent > ~~any~~ ^{any} (select percent from agents)

6 ✓	6
6 ✓	6
7 ✓	7
6 ✓	6
5	5
5	5

→ identify the customer number who purchases both products P01 and P05.

select cid from orders where pid in (P01, P02);

it give all cid who purchase P01 & P02 & both.

select cid from orders where pno = P01 and ~~and~~

cid in (select customer number ~~is~~ from orders where pno = P05);

the
the
to
us

→ Exists and not exists

If the subquery is prefixed with exists, then the condition is treated as true, if inner query returns non empty rows.

If the inner query is prefixed with not exists then the condition is treated as true, if inner query returns empty rows.

→ always exists and not exists should be have select

Retrieve all customer names, where customer places an order through agent.a05.

1) select cname from customers where cno in (select cno from orders where ano=a05).

2) or = any in place of 'in'.

3) select c.cname from customer c where exists (select * from orders where c.cid=o.cid and o.cid=a05);
(time complex).

4) select c.cname from customer c, orders o where c.cid=o.cid and o.cid=a05;
(space complex).

select cno from cno where both p01 and p02

→ 1) select cno from orders o1 where exists o1.pno=p01 and

exists (select * from orders o2 where o1.cno=o2.cno and o2.pno=p05).

exists

→ Any query formulated in all and any can also be formulated in exists.

→ exists queries always leads to correlated sub queries.

→ They are less efficient.

All and any

→ Any query formulated in exists may not be formulated in all and any.

→ They need not be correlated queries.

→ more efficient.

Joins

→ It is multi
(i)

(ii) natural
join

(iii)

(iv)

as

and

all

exists

(i)

Various

①

②

③

Joins :-

→ It is an alternative to subqueries to retrieve the results from multiple tables.

(i) self join: The table is joined with itself to get appropriate results.

Ex: To get the manager ids of all the employees.

(ii) natural join: where two tables are combined based on the available common columns.

(iii) inner join: Join operation is performed on matching columns and result is retrieved if there is a match.

(iv) outer join: Join is performed on matching column and results are retrieved based on matching rows. If there is no matching row, appropriate results are displayed with null values.

Ex: (a) left join: (b)

Student left join library.
It displays all the rows from student and rows from the library if they have a match with the student, otherwise, library values will be null for a student who has not performed any transaction with the library.

(b) right join:

Student right join library.

Here all the rows from library are listed out and rest of the results are same as left join.

(c) full join:

Student full join library.

Here all the rows from library and student are listed out and rest of the results are same as left and right outer joins.

Various ways of specifying join condition:-

- ① from customer c, join orders o on c.cid = o.cid
- ② from customer join orders using old.
- ③ from customer natural join orders.

4) ----- from Customer, also where $ECID = 0-12$.

→ Consider the following two tables to execute the queries.

Branch		Flight	
S.NO	name	Fno	Station from
1	chennai	A4	delhi
2	Bang	B4	hyd
3	hyd	C4	chennai

① select B*, F* from Branch B, Flight F where B.name = F.Station from

1	chennai	C4	chennai	(inner join)
3	hyd	B4	hyd	

2) select B*, F* from Branch B left join Flight F where

B.name = F.Station from

1	chennai	C4	chennai	
2	Bang	X	X	(left outer join)
3	hyd	B4	hyd	

3) select B*, F* from Branch B right join Flight F where

B.name = F.Station from

X	X	A4	delhi	
2	hyd	B4	hyd	
1	chennai	C4	chennai	

4)

Branch B full Flight F.

1	chennai	C4	chennai	
3	hyd	B4	hyd	(left outer join)
2	Bang	X	X	(right outer join)
X	X	A4	delhi	

Relational algebra

Different operators in RA,

(i) ~~basic~~ basic relational algebra operations:

σ , π , $\div$, rename(ρ), $\leftarrow$ (assignment),
 $\bowtie$ (join).

(ii) Set theoretic operators:

$\cup$, $\cap$, $-$, $\times$ (Cross product).

(iii) Extended relational algebra operators

max, min, Avg, sum, count, Count(*)

select(σ) :-

Get ~~select~~ details of students whose age is greater than 15.

Syntax $\sigma_{\langle \text{condition} \rangle}(\text{relation})$

$\sigma_{\text{age} > 15}(\text{student})$

Characteristics :-

- (i) It is unary operator.
- (ii) operation is applied to each tuple individually.
- (iii) degree of the relation from select operation is same as input relation.
- (iv) It is commutative in nature.

$$\sigma_{\langle C_1 \rangle}(\sigma_{\langle C_2 \rangle}(R)) = \sigma_{\langle C_2 \rangle}(\sigma_{\langle C_1 \rangle}(R)).$$

Ex:

$$\sigma_{\text{age} > 15}(\sigma_{\text{branch} = \text{'CSE'}}(\text{student})) = \sigma_{\text{branch} = \text{'CSE'}}(\sigma_{\text{age} > 15}(\text{student})).$$

$$\sigma_{\langle C_1 \rangle}(\sigma_{\langle C_2 \rangle}(\sigma_{\langle C_3 \rangle}(R))) = \sigma_{\langle C_1 \rangle \text{ and } \langle C_2 \rangle \text{ and } \langle C_3 \rangle}(R).$$

The no of rows returned by selection operation is less than or equal to no of rows in original table.

$$|\sigma_{\langle C \rangle}(R)| \leq |R|.$$

Projection - It is used to filter the columns

Syntax $\pi_{\langle \text{attribute list} \rangle} (R)$

Ex: $\pi_{\text{name, age}}(R)$

The degree of ~~off~~ relation = no of attributes mentioned in the attribute list.

* Projection eliminates duplicates automatically

* Commutative property cannot be applied for projection.

$$\pi_{\text{name}}(\pi_{\text{name, age}}(\text{student})) \neq \pi_{\text{name, age}}(\pi_{\text{name}}(\text{student}))$$

→ Find the student name and age who belong to CSE and marks > 500.

SQL: select name, age from student where branch='CSE' and marks > 500;

RA: $\pi_{\text{name, age}}(\sigma_{\text{branch='CSE' and marks > 500}}(\text{student}))$

→ SQL: select * from student

RA: $\sigma(\text{student})$

Rename operator

It is used to store the intermediate result of any query.

$$\rho(\text{branch-cse}) \leftarrow \sigma_{\text{branch='CSE' and marks > 500}}(\text{student})$$

Intermediate results are stored in a temporary table using argument operator and it is renamed to

branch-cse

The rename operator can also be used to rename attributes

$\rho(\text{st})$

Division

It is used

for

"C"

make

list

Condition

The

Ex

a

A
a
a
a
a
c

R

$P(\text{student name, student age}) \leftarrow T_{\text{name, age}} (\text{branch - CSE})$

Division ($\div$) :-

It is well suited for the queries that consists the word "for all".

This operation defines a relation and the attributes "C" that consists set of tuples from "R", that matches the combination of every tuple in the attribute list "C".

Condition for $\div$ operation :-

The denominator must be subset of numerator.

Ex: $\frac{abc}{c} = ab$

and result is numerator - denominator.

A	B	C
a ₁	b ₁	c ₁
a ₂	b ₁	c ₁
a ₁	b ₂	c ₁
a ₁	b ₂	c ₂
a ₂	b ₁	c ₂
a ₁	b ₂	c ₃
a ₁	b ₂	c ₄
a ₁	b ₁	c ₅

$S = \frac{C}{c_1}$

$R \div S =$

A	B
a ₁	b ₁
a ₂	b ₁
a ₁	b ₂

$S =$

C
c ₁
c ₂

$R \div S =$

A	B
a ₁	b ₂
a ₂	b ₁

all combinations of A and B having both c₁ and c₂.

$R \div S$

$S =$

C
c ₁
c ₂
c ₃
c ₄

$R \div S$

A	B
a ₁	b ₂

$S =$

B	C
b ₁	c ₁

$R \div S =$

A
a ₁
a ₂

→ Set the following operations:

A	B	C
a ₁	b ₁	c ₁
a ₂	b ₂	c ₃
a ₂	b ₁	c ₂

A	B	C
a ₁	b ₁	c ₁
a ₁	b ₁	c ₂
a ₁	b ₂	c ₃
a ₃	b ₂	c ₃

a ₁	b ₁	c ₁
a ₁	b ₂	c ₃
a ₂	b ₁	c ₂
a ₁	b ₁	c ₂
a ₃	b ₂	c ₃

a ₁	b ₁	c ₁
a ₁	b ₂	c ₃

a ₂	b ₁	c ₂
----------------	----------------	----------------

a ₁	b ₁	c ₂
a ₃	b ₂	c ₃

$\sigma_{C=C_1}(R_1)$

$\pi_{\text{attributes}}$

$\frac{R_2}{R_1}$

A Compl

1) inter
and

2) Join
ca

3) de

Ex The name of two relations are R₁ and R₂, where R₁ contains n₁ tuples and R₂ contains n₂ tuples. $n_2 \geq n_1 \geq 0$. Give the minimum and max size in terms of tuples, for the result relation produced by each of the following relational algebra expression, and mention assumption if any.

	O/P	assumption
R ₁ ∪ R ₂		union compatible
R ₁ ∩ R ₂		union compatible
R ₁ - R ₂		1)
R ₁ × R ₂		

min	max
n ₂	n ₁ + n ₂
0	n ₁
0	n ₁
n ₁ × n ₂	n ₁ × n ₂

→ The +
H

L

$$\sigma_{\langle C_1 \rangle}(R_1)$$

$\pi_{\langle \text{attribute list} \rangle} R_2$ attribute list contains primary key

0

N_1

N_1

N_2

$$\frac{R_2}{R_1}$$

$$R_2 \supseteq R_1$$

0

N_2

A complete set of relational algebra operators:

$$\sigma, \pi, \cup, \cap, \times, \div, \setminus$$

1) Intersection can be expressed in terms of union and difference.

$$R \cap S = R \cup S - ((R - S) \cup (S - R))$$

2) Join can be expressed with help of σ and cartesian product.

$$R \bowtie_{\langle C_1 \rangle} S = \sigma_{\langle C_1 \rangle}(R \times S)$$

3) division can be expressed with $\pi, \times, -$

$$T_1 \leftarrow \pi_{\text{attribute}}(R)$$

$$T_2 \leftarrow \pi_{\text{attribute}}(S \times T_1) - R$$

$$T_{\text{result}} \leftarrow T_1 - T_2$$

$$R \div S = T_{\text{result}}$$

T_1
A B
 a_1, b_1
 a_2, b_1
 a_1, b_2
 a_2, b_2
SXT₁

a_1, b_1, c_1
 a_2, b_1, c_1
 a_1, b_2, c_1
 a_2, b_2, c_1

→ The precedence of RA operators

$$\begin{array}{c} \pi \\ \sigma \\ \times \\ \div \\ \cap \\ \cup \end{array}$$

max

$$N_1 + N_2$$

$$N_1$$

$$N_1$$

$$N_1 \times N_2$$

Transactions

TRANSACTIONS

A Transaction is a logical unit of work and consists of multiple read and write operations on the data base.

It is associated with the following 3 problems:

- (i) It may create an inconsistent result.
- (ii) It may create problem in concurrent execution.
- (iii) It may create an uncertainty to decide when to make changes permanent.

→ To solve the above 3 problems, we have defined some properties for the transactions and they are called ACID properties.

→ If any transaction satisfies these properties, then we can say it will be free from above problems.

Atomicity - Either all updates of transactions occur in the data base or none of them otherwise there is no partial transactions.

→ It is the responsibility of transaction recovery manager of DBMS to ensure this property.

Consistency - Operations of the transactions must bring data base from one consistent state to another consistent state.

To ensure the consistency we will use integrity constraints and they are depending on policies of organization.

Isolation

Interfer
to ens
must

Serial

by a
person
per
along

The prob

A big

A transc

Start

Isolation

Nai
man
Scott
John

(i) Dis
con

At

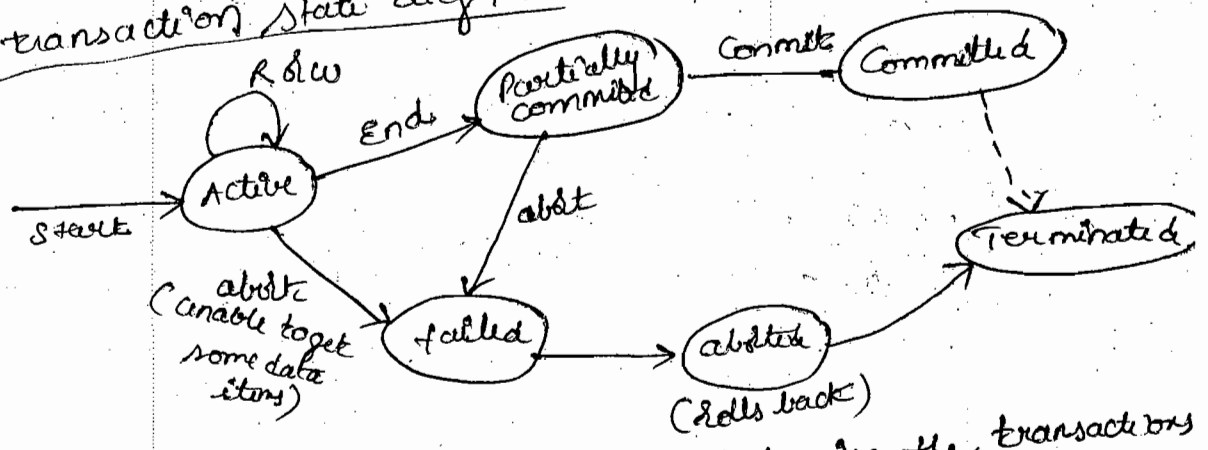
Isolation: The execution of one transaction should not interfere with the execution of another transaction. To ensure the isolation, concurrency control protocols must be implemented.

Availability: The changes applied to a data base by a committed transaction must be made permanent. To ensure this logs in the dbms will be used along with the recovery system.

The primitives in the transaction:

A ~~begin~~ begin Begin, Read, Write, End, Commit, Roll back, abort

A transaction state diagram



Isolation problems & problems in interleaving the transactions

<u>name</u>	<u>sal</u>	<u>dept</u>
mary	300	sales
scott	500	mkty
John	400	production

(i) Dirty read or reading uncommitted data or WR conflict:

At 10 am ST TX₁
 10:10 am set marys Sal to 500
 10:20 am Rollback TX₁
 Commit TX₁

At 10:15 AM St TX₂
 10:15 AM Read all sal (here uncommonly 500 will be read).

→ lost

→ Not Repeatable Reads (values):

At 10 am St TX₁
 At 10:05 am Read all salaries (300, 500, 400)
 At 10:10 am Read all salaries (500, 500, 400)

At 10:05 am St TX₂
 At 10:06 am set may's salary to 500
 10:07 Commit.

→ Unst

→ Not Repeatable Reads (phantom row) of RW conflict

At 10:00 AM St TX₁
 10:05 AM Read all salaries (300, 500, 400)
 10:10 AM Read all salaries (300, 500, 400, 1000)

At 10:05 AM St TX₂
 10:06 AM Insert ('Sam', 1000, 'Personal')
 10:07 AM C₂.

There a
 (i) Phy

→ Incorrect Summary problem:

At 10:00 AM St TX₁
 10:05 AM Sum (Sal) 1200
 10:10 AM (Sum ? (Sal) 1400

At 10:05 AM St TX₁
 10:06 AM Set may's salary to 500
 10:07 Commit.

→ Ex

(ii) Log

at

mary 500 will

→ lost update problem (WU conflict)

10:00 AM ST TX₁
10:10 AM SET mary sal to 500
10:20 AM Commit

10:10 AM ST TX₂
10:15 AM SET mary sal to 600
10:21 AM C₂.

→ Unstable Cursors:

10:00 AM ST TX₁
10:05 AM SET mary sal to ~~1000~~ 500
10:15 AM Read all sal

AT 10:10 AM ST TX₂
10:11 AM drop sal column
10:12 AM C₂.

There are two ways to drop a column:

(i) Physical dropping:

Alter table tablename drop column column-name
Ex: Alter table student drop column sal.

(ii) Logical drop:

Alter table tablename set unused (column-name)
afterwards, if you want to have it again set use.

Schedule

The order of the execution of the transactions is known as schedule.

Two types of schedules:

(i) serial (ii) non serial.

The no serial schedules = $n!$ where n = no of transactions.

Assume there are m transactions and no of operations in these transactions are $n_1, n_2, n_3, \dots, n_m$.

The no non serial schedules =
$$\frac{(n_1 + n_2 + \dots + n_m)!}{n_1! \times n_2! \times \dots \times n_m!}$$

Ex:- Assume there are two transactions and
~~Transaction~~ T_1 contains 2 and T_2 contains 5 operations.
 Find the no. of serial and non serial schedules.

Sol) serial schedules = 2.

non serial schedules =
$$\frac{7!}{2! \times 5!} = \frac{6 \times 7}{2} = 21$$

Serializability

A schedule S of n transactions is serializable if it is equivalent to some serial schedule of the same ' n ' transactions.

There are ~~two~~ three types of serializability:

(i) Result serializable.

(ii) ~~view~~ conflict serializable.

(iii) view serializable.

$x = 100$

$R(x)$

$x = x + 10$

$w(x)$

$= 110$

If $x \leq 200$

$x = 100$

$R(x)$

$x = x * 1.1$

$w(x)$

$= 110$

then they are not equal.

The end
 output
 Δ 9

Conflict

ok

Two

all

(i) ∇

(ii) ∇

(iii) ∇

Ex: 9

and

are

(i)

(2)

(3)

→ Consider

are

then

if

now

→ side

the

on

The end results of schedules are heavily depend on input of schedules
 Δ result equivalences are not generally used.

Conflict Serializability

It is based on the conflict operations.

Two operations are conflict, if they satisfy all three of the following conditions.

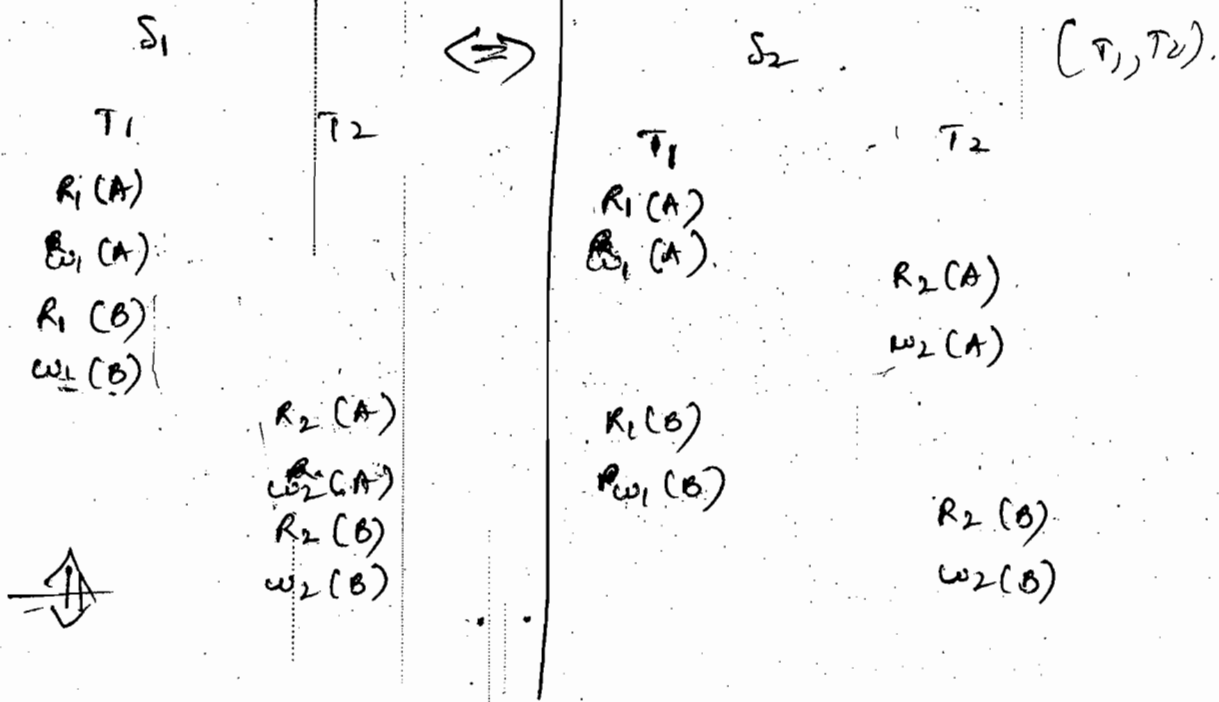
- (i) They belong to different transactions
- (ii) they access the same data item
- (iii) At least one of the operations is a write operation.

Ex: If we have $X_i(A)$ and $Y_j(B)$ mention under what conditions these two operations are confliction.

- (1) $A = B$
- (2) $i \neq j$
- (3) x & y must be write operation.

→ Consider T_i and T_j are two transactions. I_i and I_j are two consecutive instructions of T_i and T_j . Then operations are said to be ^{conflicts} if they refer same data item and ^{at least} one of them must be write operation.

→ Identify non conflicting operations and we change the order of non conflicting operations to get a non serial schedule, that is equal to serial schedule.



product

T_1 For

as if

per

S_1

T_1 T

$R(A)$

$A = A + 50$

$w(A)$

$R(B)$

$B = B + 50$

$w(B)$

R

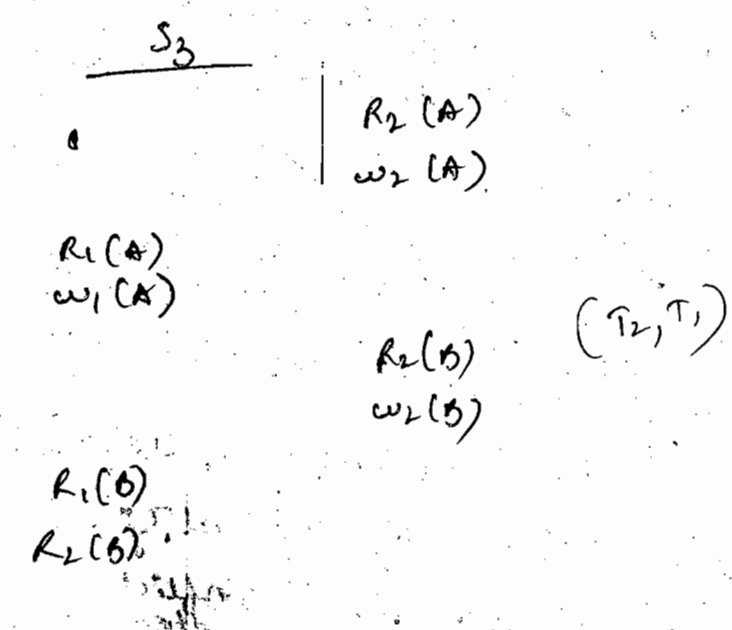
t

A

w

B

S_1 and S_2 are
Conflict



view equivalent

Consider the two schedules S_1 and S_2 , they are said to be view equivalent if the following conditions are satisfied.

(i) For each data item Q , if T_1 reads Q in S_1 , then T_1 in S_2 must also read the initial value of Q .

(ii) For each data item Q , if T_1 in S_1 reads a data item produced by a transaction T_2 , then T_1 in S_2 must also read the data item

S_1 =

→ Differ

(i) Co

w

(ii) f

to

(cu)

(T_1, T_2) .

produced by T_2 .

(iii) For each data item 'Q', if T_1 in S_1 performs a final write operation, then T_2 in S_2 must also perform a final write operation on data item

S_1		S_2		T_1	T_2
T_1	T_2	T_1	T_2		
$R(A)$		$R(A)$	$R(A)$	$R(A)$	$R(A)$
$A = A - 50$			$E = E \times 0.1$	$A = A - 50$	$E = A \times 0.1$
$w(A)$			$A = A - E$	$w(A)$	$A = A - E$
$R(B)$			$w(A)$		$w(A)$
$B = B + 50$			$R(B)$		
$w(B)$			$B = B + T$		
	$R(A)$	$R(A)$	$w(B)$	$R(B)$	$R(B)$
	$E = A \times 0.1$	$A = A - 50$		$B = B + 50$	$B = B + E$
	$A = A - E$	$w(A)$		$w(B)$	$w(B)$
	$w(A)$	$R(B)$			
	$R(B)$	$B = B + 50$			
	$B = B + T$	$w(B)$			
	$w(B)$				

S_1 and S_3 are view equivalent.

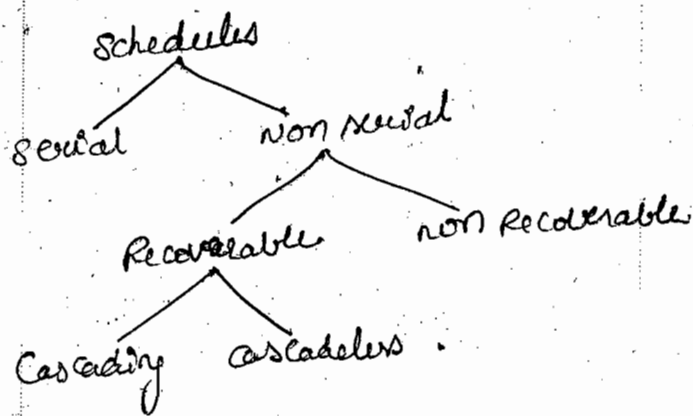
→ Differences b/w conflict and view:

- (i) conflict serializability is easy to achieve but view serializability is difficult to achieve.
- (ii) Every conflict serializable is a view serializable but reverse is not true.
- (iii) It is easy to test conflict serializability but expensive to test view serializability.

(v) most of the concurrency control schemes used in practice are based on conflict serializability.

Schedules can also

→ the recoverable schedules can also be



S1: $R_1(X), R_2(X), w_1(X), R_1(Y), w_2(X), C_2, w_1(Y), C_1$

T1
 $R_1(X)$

$w_1(X)$
 $R_1(Y)$

$w_1(Y)$
 C_1

T2
 $R_2(X)$

$w_2(X)$
 C_2

S2: $R_1(X), w_1(X), R_2(X), R_1(Y), w_2(X), C_2, C_1$

T1
 $R_1(X)$

$w_1(X)$

$R_1(Y)$

C_1

T2

$R_2(X)$

$w_2(X)$
 C_2

→ if these trans schedule is schedule S_1 is

→ Cas Ca
TO be in Con here but if not TO be

but
→ The Re cit

used
 serializability → If there is no necessity to roll back a committed transaction, then it is called a recoverable schedule otherwise it is a non recoverable schedule.

S_1 is recoverable and S_2 is non recoverable.

→ Cascading schedules

To make a non recoverable schedule to recoverable like S_2 , the following modification is needed i.e. Commit T_2 only after T_1 is committed & aborted.

Here if T_2 is reading uncommitted data from T_1 , but T_2 should commit only after T_1 is committed. If you roll back T_1 , T_2 must also be rolled back and it is called cascading roll back.

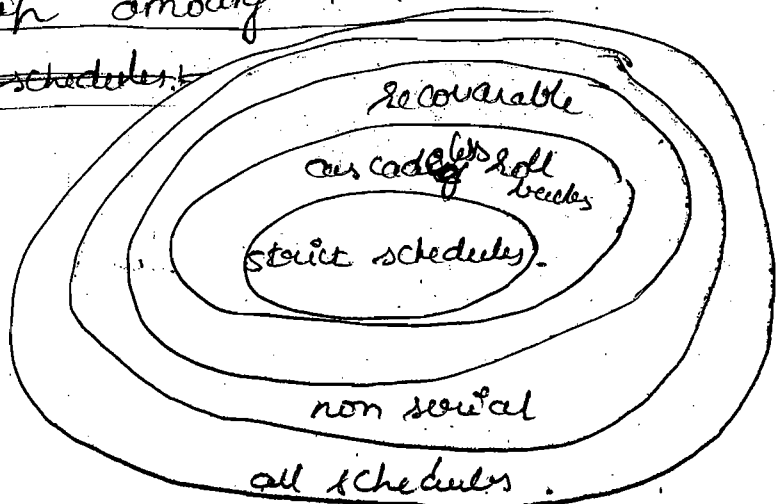
To avoid the cascading roll backs, modification is provided to S_2 like

T_1	T_2
$R_1(x)$	
$W_1(x)$	
$R_1(x)$	
a_1	
	$R_2(x)$
	$W_2(x)$
	c_2

but it leads to a serial schedule.

→ The relationship among the schedules

(i) strict schedules



→ Recoverable schedule (p 419)

T_2 commits after T_1 , if T_2 has read any data item written by T_1 . Ex: $r_1(x); r_2(x); w_1(x); r_1(y); w_2(x); w_1(y)$

Cascades schedule (p 419)

T_2 reads a data item only after T_1 has written to it and terminated.

Strict schedule

T_2 reads a data item after T_1 has written to it and terminated, in addition T_2 writes a data item after T_1 has written to it and terminated.

Page 98:

Problems

1) which of following is conflict serializable schedule, determine the equivalent serial order.

(i) $r_1(x), r_3(x), w_1(x), r_2(x), w_3(x)$.

$T_3 < T_2$ $T_1 < T_2$

(ii) $r_1(x), r_3(x), w_1(x), w_2(x), r_2(x)$.

$3 \rightarrow 1 \rightarrow 2$

(iii) $r_1(x), r_3(x), w_2(x), w_1(x), r_2(x)$.

$T_3 < T_2$

$T_2 < T_1$

$3 \rightarrow 2 \rightarrow 1$

(iv) $r_3(x), r_2(x), w_3(x), r_1(x), w_1(x)$

$2 < 3$ $3 < 1$

T_2, T_3, T_1

→ $r_3(x)$

2) a) do

b) R

c)

3) (a) T

(b) T

(b)

→ $r_3(x), r_2(x), r_1(x), w_3(x), w_1(x)$
 $\uparrow \quad \uparrow$
 $T_1 < T_3$

my data
 $w_2(x), w_4(y)$

1. ~~write~~

2. a) dirty read

$R_1(x), R_2(y), w_1(x), R_2(x)$ - - - - -
 $\uparrow \quad \uparrow$

b) Repeatable reads.

$R_1(x), R_2(y), w_2(x), R_2(y), R_1(x)$
 $\uparrow \quad \uparrow$
 $T_2 < T_1$

c) $R_1(x), R_2(x), R_2(y), w_1(x), w_2(x)$ - - -

$T_3 \leftarrow T_1$
 $\downarrow$
 T_2
 the equivalent
 of schedule.
 → 2

3) (a). $T_1: R(x), T_2: R(x), T_1: w(x), T_2: w(x)$.
 $\uparrow \quad \uparrow \quad \uparrow \quad \uparrow$
 $T_2 < T_1 \quad T_1 < T_2$

not conflict serializable.

(b) ~~$T_1: w(x), T_2: R(y), T_1: R(y), T_2: R(x)$~~

$T_1 \quad T_2$
 $R_1(x) \quad R_2(x)$
 $w_1(x) \quad w_2(x)$

Cascades schedule, recoverable,
 view serializable (T_1, T_2)

(b) $T_1: w(x), T_2: R(y), T_1: R(y), T_2: R(x)$

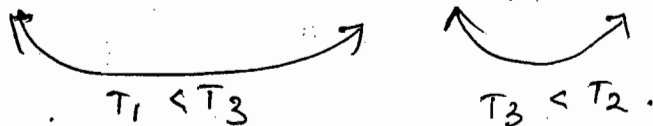
$\uparrow \quad \uparrow$
 $T_1 < T_2$

It is conflict serializable, so view serializable.
 To decide about recoverability, common statement
 should be given.

T_1 T_2
 $w_1(x)$ $R_2(y)$
 $R_1(y)$ $R_2(x)$

It is cascading schedule, if T_1 is rolled back, then T_2 must also be rolled back. If T_1 commits first then T_2 , it is recoverable if T_2 commits first then T_1 , it is not recoverable.

c) $T_1: R(x), T_2: R(y), T_3: w(x), T_2: R(x)$



It is conflict and view serializable.

It is cascading

T_1 T_2 T_3
 $R_1(x)$ $R_2(y)$ $w_1(x)$
 $R_2(x)$

If T_3 commits before T_2 , it is recoverable else it is not.

d) ~~$T_1: R(x)$~~ not conflict, ~~not view~~, ~~not~~ cascadeless.

$T_1: R(x), T_1: R(y); T_1: w(x), T_2: R(y), T_3: w(x)$

$T_1: w(x), T_2: R(y)$

T_1
 $R_1(x)$
 $R_1(y)$
 $w_1(x)$

$w_1(x)$

e) It is an

$T_1: R$

f) not casc

g) Conf

h) not

i) T_1

j) co

k)

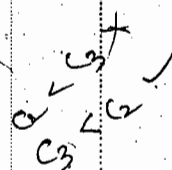
l) not

T_1 T_2 T_3

$R_1(X)$
 $R_1(Y)$
 $w_1(X)$

$R_2(Y)$ $w_3(Y)$

$w_1(X)$ $R_2(Y)$



e) ~~is~~ ~~Conflict~~ ~~here~~ view serializable, Cascadeless and recoverable.

$T_1: R(X)$, $T_2: W(X)$, $T_1: W(X)$, $T_2: \text{abort}$, $T_1: \text{Commit}$

$\uparrow$ $\uparrow$ $\uparrow$ $\uparrow$
 $T_1 < T_2$ $T_2 < T_1$

f) not conflict serializable, not view, recoverable and Cascadeless

g) Conflict, view, recoverable, Cascadeless.

h) not Conflict, not view, not recoverable, not Cascadeless

i) $T_1: W(X)$, $T_2: R(X)$, $T_1: W(X)$, $T_2: C$, $T_1: \text{Abort}$, $T_1: \text{Abort}$.

Conflict, view, not Cascadeless, not recoverable.

j) Conflict and view, not recoverable, Cascadeless.

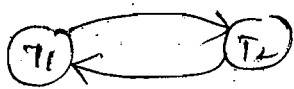
k) ~~same as j~~ not Conflict, not view, Cascadeless, recoverable.

l) not Conflict, not view, recoverable, not Cascadeless.

Testing for Conflict Serializability:-

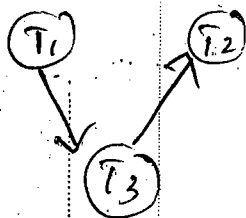
To test the conflict serializability we can draw a graph and it consists set of nodes and set of edges. no of nodes = no of transactions. Each edge is of the form $T_a \rightarrow T_b$ where T_a is starting node and T_b is ending node such an edge must be created if one of the operations in T_a appears in schedule before some conflicting operation in T_b .
Once the graph is drawn, check for the loops. If there is a cycle, the schedule is not in conflict serializability.

4) a) $R_1(x), R_2(x), w_1(x), w_2(x), C_1, C_2$.



not ~~serial~~ conflict serializable

b) $R_1(x), R_2(y), w_3(x), R_2(x), r_1(y), C_1, C_2$.



Conflict serializable.

c) $R_1(x), w_2(x), w_1(x), a_2, C_1$

node T_2 is removed because of a_2 .



d) $w_1(x)$

e) R_1

f) $f)$

g) $g)$

h) $h)$

i) $i)$

Implement

1) Access mode

read and

2) Diagona

that can

executed

3) Isolation

transacti

onecting

Concurrent

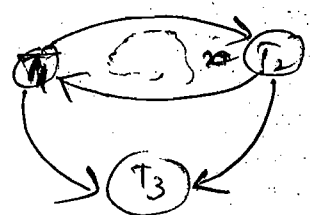
4

we a
 set of
 set
 Ta is
 an
 in
 before
 loops.
 in

c) $w_1(x), R_2(x), w_3(x), C_2, C_1$

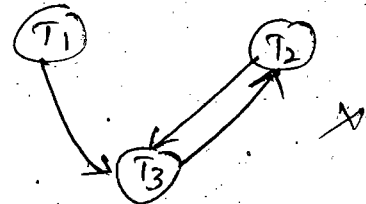


e) $R_1(x), w_2(x), w_3(x), R_3(x), C_1, C_2, C_3$



not conflict serializable.

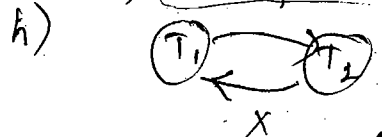
f) $R_1(x), R_1(y), w_1(x), R_2(y), w_3(y), w_1(x), R_2(y)$



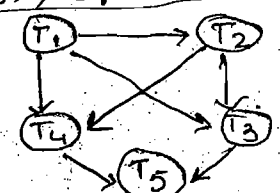
g) Draw procedure graph for $w_3(A), R_1(A), w_1(Z), R_2(B)$

g) $T_2 \rightarrow T_3 \rightarrow T_1$ ✓ $w_2(y), w_3(B), C_1, C_2, C_3$

h) Draw procedure graph for $R_1(A), w_1(A), R_2(A), R_2(B), C_2, R_1(B), w_1(B), C_1$



i) Consider the procedure graph, is the corresponding schedule Conflict Serializable



(i) no cycles,

Implementation of transaction SQL

- 1) Access mode - There are 2 types of access mode read only and read and write.
- 2) Diagnostic size - It determines the no of errors conditions that can be recorded on the user on the most recently executed SQL Statement.
- 3) Isolation level - It controls the extent to which a given transaction is exposed to the actions of other transactions executing concurrently.

Concurrency	Dirty read	unrepeatable	phantom
Level 1	✓	✓	✓
Level 2	✓	✓	✗
Level 3	✗	✗	✗

phantom
 {
 test
 w17-p

✓

Test * Classification of Concurrency protocols

(i) Lock based protocol

- (a) 2 phase locking protocol → Basic 2PL
→ Conservative 2PL.
→ Strict 2PL
→ Heuristic 2PL.

(b) Graph based protocol.

(ii) Time stamp based protocol

- (a) Timestamp ordering protocol
(b) Thomas write rule.

(iii) multiple granularity protocol.

(iv) multi version protocols.

T_i
 $R(A)$

① $RTS(T_i) < WTS(A)$ X
900 < 910

② $TS(T_i) \geq WTS(A)$ ✓

T_{i2}
 $W(A)$

① $TS(T_{i2}) < RTS(A)$ X
900 < 910

② $TS(T_{i2}) \leq WTS(A)$ X

③ $TS(T_{i2}) > WTS(A)$ ✓

→ Consider a data base, that is spanned in terms of the following hierarchy.

- The database itself is an object and it contains 2 files, each file contains 1000 pages each page contains 100 records and these records are identified as p.i where p is a page no and i is a record number.

For each of the following operations indicate the sequence of lock request that must be generated by a transaction manager.

① Read



② Read

IS
S on
IS
S on

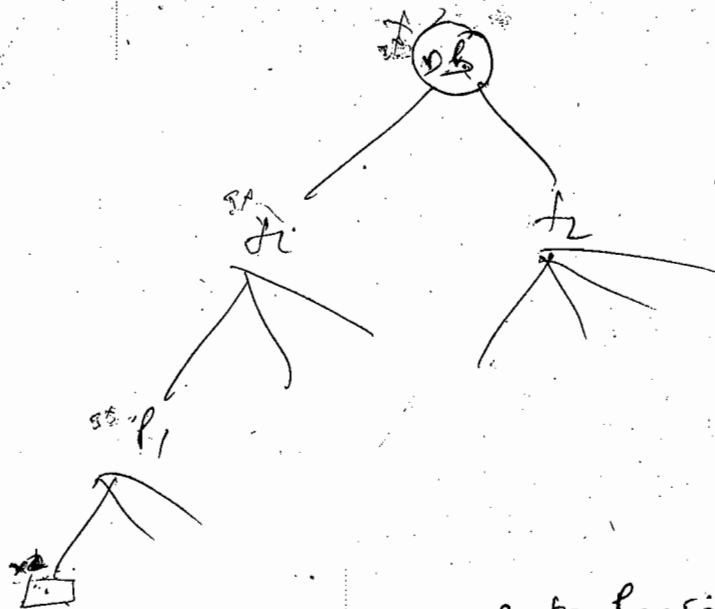
S on

IS

IS

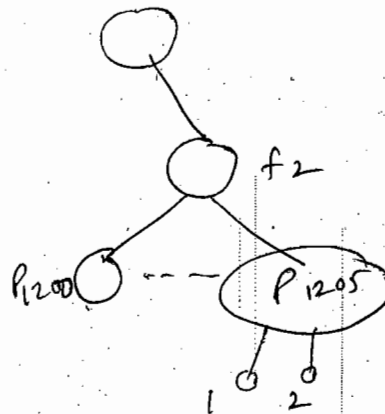
③ Read

① Read rec'd $P_{1200.5}$



IS on DB
IS on f2
IS on P_{1200}
Son $P_{1200.5}$

② Read records $P_{1200.98}$ to $P_{1205.2}$



IS on P_{1205}
S on $P_{1205} : Y_2$
IS on P_{1200}
Son $P_{1200} : 99/99/100$
Son $P_{1201}, P_{1202}, P_{1203}, P_{1204}$

AS
13
01
00

IS on f2
IS on DB

③ Read P_{500} to P_{600}

IS on DB
S on f1

IS on DB
IS on f1
Son $P_{500}, 501 \dots 600$

→ Read all pages in F_1 and modify about 10 pages which can be identified only after reading F_1 .

5 IX on DB
8 on F_1

→ delete the 1 record from each page.

1 IX on DB
X on F_1 and F_2

→ delete all records

1 IX on D
X on F_1 and F_2

Indexes

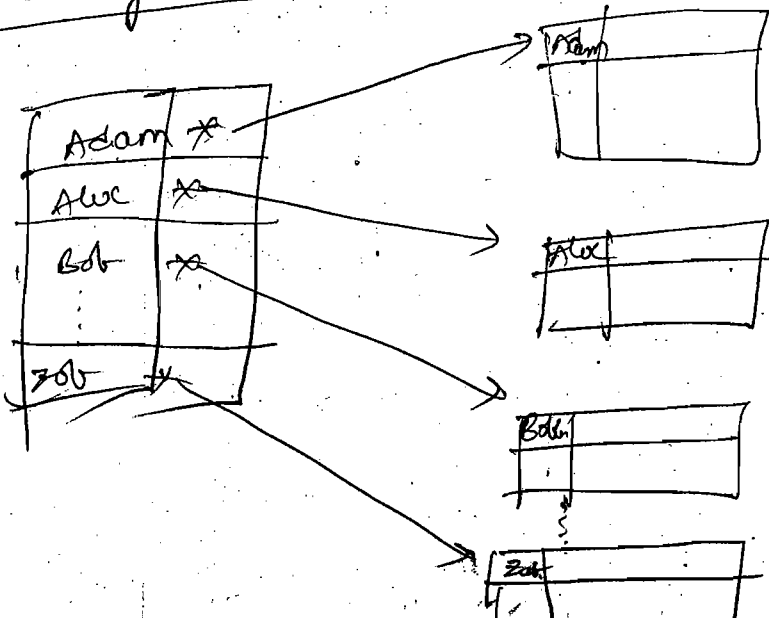
→ Types of indexes:

- (i) Primary indexes.
- (ii) Secondary indexes.
- (iii) clustering indexes.

Other classification:

- (i) sparse index.
- (ii) dense index.

Primary index



Characteristics

- (i) SE is
- (ii) SE G
- (iii) Tle
- (iv) SE

Consider

→ Consider

∴ n

Assume

of the

Contai

Size of

Search

(a) u

(b) 2

out
after

Characteristics

- (i) It is on secondary file
- (ii) It consists 2 entries.
(a) Key (b) a ~~block~~ block pointer
- (iii) The no of index records ~~no~~ = no of files & pages or blocks.
- (iv) It is an example for sparse index.

Example

→ Consider a block size = 1024 B, then record length = 100

$$\therefore \text{no of records/block} = \frac{1024}{100} = 10 \text{ records}$$

Assume there are 30,000 records in data base

$$\therefore \text{no of blocks} = \frac{30000}{10} = 3000 \text{ blocks.}$$

If the search key ~~no~~ length = 9 bytes and block pointer contains 6 bytes then each index record length = 15 bytes

$$\text{Size of total } \underline{\text{no of index}}^{\text{file}} \text{ records} = 3000 \times 15 \text{ Bytes.}$$

$$\therefore \text{no of index blocks} = \frac{3000 \times 15}{1024} = 45 \text{ blocks.}$$

Search derivation

(a) without index

$$\log_2 3000 = 12$$

(b) with index

$$\log_2 45 = 6$$

$$6 + 1 = 7$$

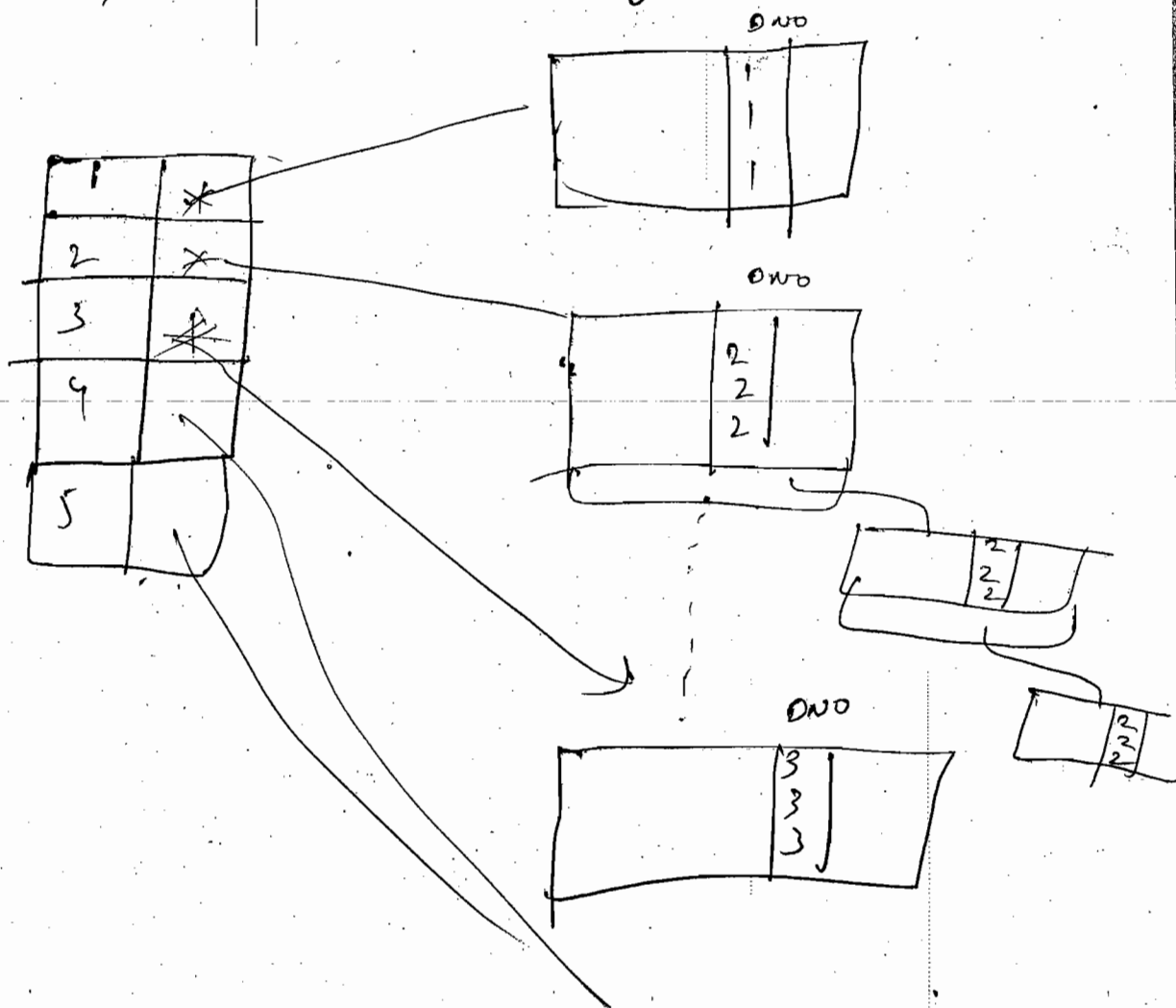
→ Problems with primary index:-

The problem of with primary key index is insertions and deletions. If we insert a record in its exact position, then there is a possibility of moving anchor records. Then it leads to changes in index file structure.

Solution:- To avoid this, we will use a fill factor normally from 0.5 to 1

Clustering index:-

Primary index must be created on the field that must have a unique value for each record, then if you want to create an index on a field that does not have a unique value, then we use clustering index.



characteristic

It is cluster two block

Secondary

A also create unique are the



It is

→ Block

is

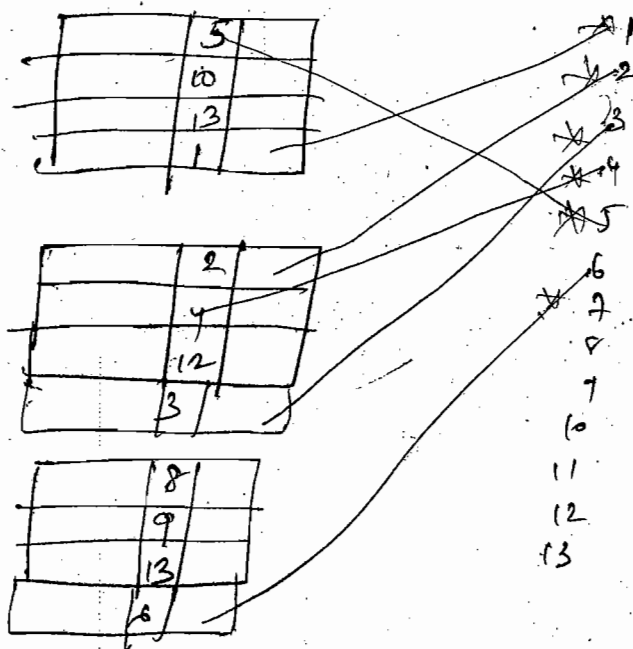
Characteristics:

It is a
new B+ tree
structure.

It is sparse index, we can have only one clustered index per table and it will have two entries, one is search key and other is block pointer.

Secondary indexes:

A secondary index is on ordered file and it is created on the fields and they need not have unique values and multiple secondary indexes are possible as they need not require physical reordering of the records.



It is a dense index.

→ Block size = 1024, record length = 100 bytes.

no records/block = $\frac{1024}{100} = 10$ records

index record length = 9 + 6 = 15

no of block index records/block = $\frac{1024}{15} = 68$

no of secondary index records = 30,000

no of blocks required = $\frac{30000}{68} = 442$ blocks

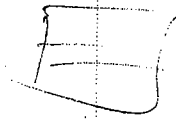
Search

(a) without Secondary Index?

$$\frac{30,000}{2} = 15000$$

(b) with secondary index.

$$\log_{2.44} = 9 + \log_{2.44} 68$$



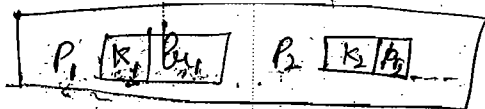
→ Consider 1.

root
level 1
level 2
level 3

→ multi level index:

They are implemented with B and B⁺ trees and these trees grow horizontally instead of vertically.

The structure of B tree node will look like



P_i is a tree pointer, pointing to another node.

P_{i,j} is a data pointer whose key value is K_j.

→ B⁺ tree
m B
data
ptr
root
of
type
(1) -

Structure



P_i

→ Consider a p is order of the tree, then no of node pointers = p, then the no of data pointers and search key values = p-1.

→ Consider a node size & block size = 512 B, a search

key length (K_i) = 9 bytes, data pointer = 7 bytes, block pointer (P_i) = 6 bytes, then find the size of node

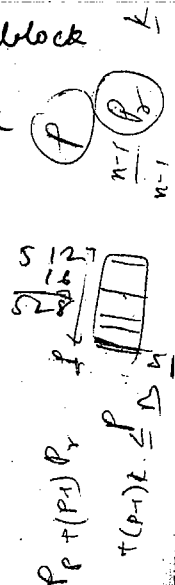
$$p \times 6 + (p-1)(16) = 512$$

$$22p = 512 + 16$$

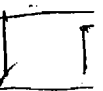
$$p = \frac{528}{22} = 24$$

Let p be the order of the tree, then

$$p \times P_1 + (p-1)K_1 + (p-1)P_2 \leq B$$



Structure



Structure

leaf

→ Consider the fill factor for the nodes 69%.

1. Size of the node = $24 \times 0.69 = 16$.

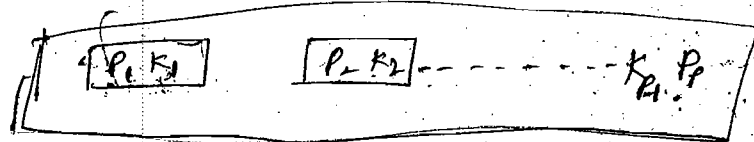
root	1 node	15 data entries	16 ptrs.
level 1	16 nodes	$(16 \times 15) = 240$	$16 \times 16 = 256$
level 2	256	3840 (256×15)	$256 \times 16 = 4096$
level 3	4096	61440 (4096×15)	4096×16

→ B⁺ tree

In B tree every node consists search key value, data pointer along with the block pointer, so redundancy is high. But in B⁺ tree, data points are stored only at the leaf nodes of the tree. ∴ B⁺ tree will have two types of the nodes.

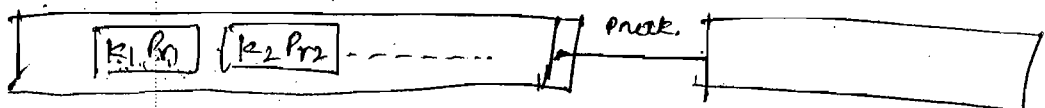
(i) internal nodes (ii) leaf node.

Structure of internal node



P_i - blocks ptr K_i - search key value, p - size of the tree

Structure of leaf node



Internal

$$(p \times p) + (p-1) K_i \leq B.$$

Leaf

$$(p_{leaf} \times p_r) + (p_{leaf} \times K_r) + P_{next} \leq B.$$

is and
- vertically

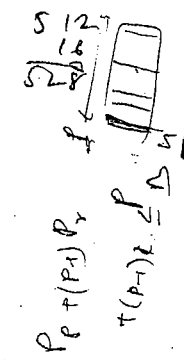
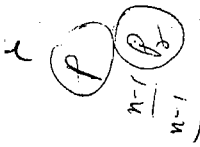
key node.

key node.

no of
pointers

x search

block



Prob 1

$$B = 512, p_{\text{row}} = 7, p_{\text{col}} = 6, p_{\text{col}} \geq 9$$

$$p_{\text{inter}} = 34$$

$$p_{\text{leaf}} = 31$$

→ Consider fill factor as 64% $\therefore p = \frac{34 \times 0.64}{\text{link}} = 223$

Root	1 node	22 data entries	23 node per
Level 1	23 node	506 (23×22)	$23 \times 23 = 529$
Level 2	529	11,638 (529×22)	$529 \times 23 = 12,167$
Level 3	12,167	2,55,507 ($12,167 \times 22$)	- - -

So no of level in B+ are less than B for the same data size.

Ans.

strict

any

com.

read

Dec

Basic

i.

phas

Corne

befor

also

items

Rigor

relc

unt

&

item

tran

when

of it

so

pha

strict 2PL: A transactⁿ T doesn't release any of its exclusive (write) locks until after it commits or aborts. Hence no other transactⁿ can read or write on that is written by T until T commits.
Dead lock.

Basic 2PL: All locks once then all unlocks
i.e. expanding & growing phase & shrinking phase.

Conservative 2PL: Lock all items it accesses before the transactⁿ begins executⁿ. If any item already locked then don't lock rest of items also & wait. Deadlock free.

Rigorous 2PL: In this, a transactⁿ T doesn't release any of its locks (exclusive or shared) until after it commits or aborts.

Diff b/w conservative & Rigorous:

In conservative must lock all its items before it starts so once the transactⁿ starts it is in its shrinking phase whereas the latter does not unlock any of its items until after it terminates (c or a) so this transactⁿ is still in expanding phase.

Finite Automata and Regular expressions

1) Find

Conversion of FA to RE there are 3 methods

(i) McNaughton-Yamada theorem (6 hours)
(Reg^k method)

(ii) Arden's method (2 or 3 mins)

(iii) state elimination method (1 min)

Arden's method

(1) This method is not applicable for E-NFA.

(ii) If P, Q are two RE's and P does not contain ϵ , then the unique solution

$$R = Q + RP \text{ can be written as } R = QP^*$$

(iii) If P contains ϵ , then we will get infinite solutions

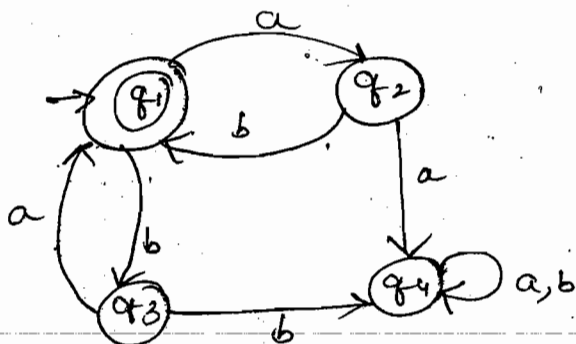
→ Find

→ $\begin{pmatrix} q \\ a \end{pmatrix}$

q
q
q

Questions

Give the RE for following FA, by using Arden's method

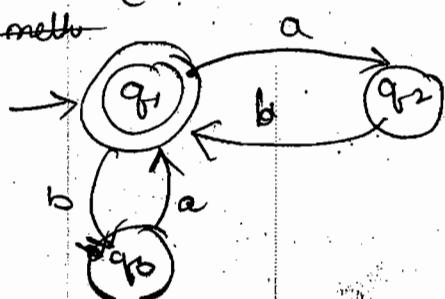


⇓
simplification

1) ~~set~~ Remove non reachable states.

2) q_4 does not reach final state so remove it.
Remove those states which does not leads to final state.

Arden's method



i) Find characteristic equation for every state

$$q_1 = \epsilon + q_2 b + q_3 a \rightarrow (1)$$

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

$$q_3 = q_1 b \rightarrow (3)$$

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

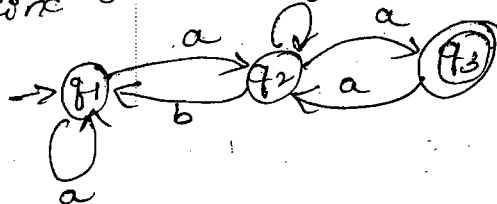
$$= \epsilon + q_1 (a b + a b a)$$

$$= \epsilon + q_1 a b$$

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

$$2) q_1 \in (a b + b a)^*$$

→ Find the RE for following FA by using arden's method.



$$q_1 = \epsilon + a q_1 + b q_2$$

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

$$q_3 = q_2 a$$

$$q_2 = q_2 b + q_1 a + q_2 a a$$

$$q_2 = q_1 a + q_2 (b + a a) \Rightarrow q_2 = (q_1 a) (b + a a)^*$$

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

$$q_1 = \epsilon + q_1 a + q_1 a a + q_2$$

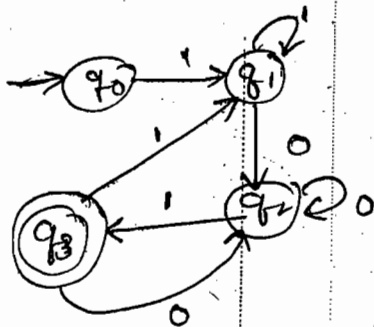
$$q_1 = \epsilon + q_1 (a + a (b + a a)^* b)$$

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

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

$$q_2 = q_1 (a + a (b + a a)^* b)^* a + \epsilon$$

→ Find RE for following FA using arden's method.



$$\begin{aligned} q_0 &= \epsilon \\ q_1 &= q_0 1 + q_1 1 + q_3 1 \\ q_2 &= q_1 0 + q_2 0 + q_3 0 \\ q_3 &= q_2 0 \end{aligned}$$

$$q_2 = q_1 0 + q_2 0 + q_3 0$$

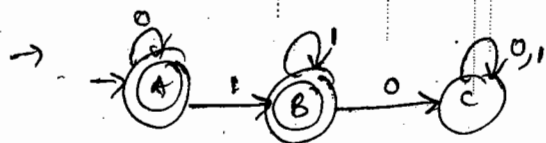
$$q_2 = q_1 0 + (0 + 10)^* 1$$

$$q_1 = q_1 1 + q_1 1 + q_1 0 (0 + 10)^* 1 1$$

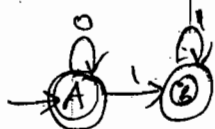
$$q_1 = 1 (1 + 0 (0 + 10)^* 11)^*$$

$$q_2 = 1 (1 + 0 (0 + 10)^* 11)^* + q_2 0 + q_2 1 0$$

$$q_3 = 1 (1 + 0 (0 + 10)^* 11)^* (0 + 10)^* 1$$



As 'C' can be removed.



$$A = \epsilon + A0$$

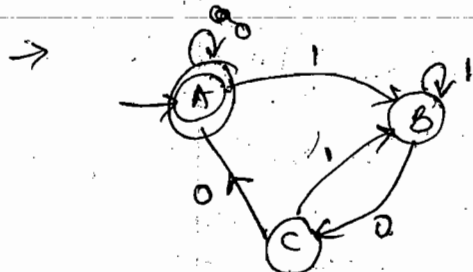
$$\Rightarrow A = 0^*$$

$$B = A1 + B1$$

$$\Rightarrow B = 0^* 1 1^*$$

$$A + B = 0^* + 0^* 1 1^*$$

$$= 0^* (\epsilon + 11^*) = 0^* 1^*$$



$$A = \epsilon + A0 + C0$$

$$B = B1 + A1 + C1$$

$$C = B0$$

$$B = B1 + A1 + B01$$

$$B = A1(1 + 01)^*$$

$$A = \epsilon + A0 + A1(1 + 01)^* 00$$

$$A = \epsilon (0 + 1(1 + 01)^* 00)^*$$

→ Give RE
sub state
Contrary!
not con

we
nom

q
q q

q = -

A
r = →