

Advance DBMS

①

Field $\rightarrow$ Marks = [10, 20, 30, 40]

Record $\rightarrow$ Collection of fields

Admn	Name	Marks
1	x	140
2	y	150

Data files $\rightarrow$ Collection of Records

Database $\rightarrow$ Collection of Data files

Data $\rightarrow$ fields $\rightarrow$ Record $\rightarrow$ Data files $\rightarrow$ Database

DBMS (Database Management System)

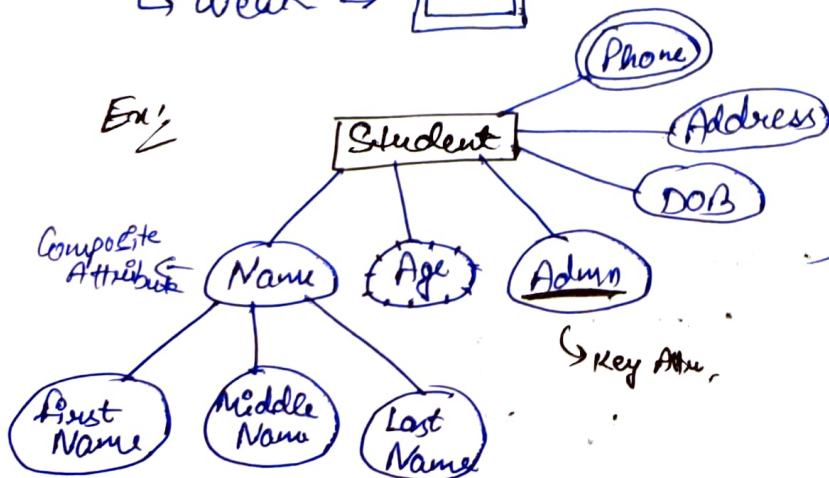
- $\hookrightarrow$ Efficient retrieval, insertion of data.
- $\hookrightarrow$ Main thing is time.

Entity Relationship Diagram

$\hookrightarrow$ Strong $\rightarrow$ 

$\hookrightarrow$ Weak $\rightarrow$ 

Ex:



Attributes

- $\hookrightarrow$ Single Valued 
- $\hookrightarrow$ Multivalued 
- $\hookrightarrow$ Composite 
- $\hookrightarrow$ Derived 
- $\hookrightarrow$ Descriptive Attribute 

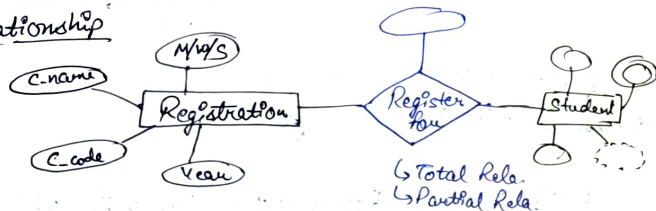
Keys:-

- Primary Key
- Partial Key
- Candidate Key
- Super Key
- Foreign Key

ER-Model

↓
Relational Model

Relationship



Conversion into Relations/Tables

Student

Ad.no	DOB	Name First Name	Name LN	Age
=	=	=	=	=

Student-Phone

Ad.no	Phone
=	=

Anomalies

- Insert
- Delete
- Update

Branch

B-Name	Ad.
=	=

Account

A-no.	B-No
=	=

If the Branch name is there in Branch Relation, then only Account can be created.

Branch (BN, BC, Assets)

Customer (CN, CS, CC)

Account (AN, BN, BAL)

Loan (LN, BN, AMT)

Depositor (CN, AN)

Borrower (CN, LN)

Basic Operations:-

- Select
- Project
- Union
- Set Difference
- Cartesian Product
- Rename

Additional Operations

- Set Intersection $(R \cap S) = \{t \mid t \in R \text{ and } t \in S\}$
- Natural Join $[R \bowtie S]$
- Division $[R \div S]$
- Assignment $[x \leftarrow \pi_{BN}(r_{cc}(\text{Branch}))]$

Q. Find names of the customer who have both loan & Account

Sol. $\pi_{CN}(\text{Depositor}) \cap \pi_{CN}(\text{Borrower})$

Natural Join:-

Let x, s be the relations on schema $x(R)$ & $s(S)$ respectively. Then, $x \bowtie s$ is a relation on schema $R \cup S$ obtained as follows:-

- Consider each pair of tuples t_x from x and t_s from s .
- If t_x & t_s have the same value on each of the attributes in $x \cap s$, add a tuple t to the result where,
 - t has the same value as t_x on x and t_s
 - t has the same value as t_s on s .

	A	B	C	D
x	α	1	α	α
	β	2	γ	α
	γ	4	β	β
	α	1	γ	α
	β	2	β	β

	B	D	E
s	1	α	α
	3	α	β
	1	α	γ
	2	β	β
	3	β	β

$R \bowtie S$

A	B	C	D	E
x	1	x	a	x
x	1	x	a	y
x	1	y	a	x
x	1	y	a	y
s	2	b	b	s

$$R \bowtie S = \pi_{R.A, R.B, R.C, R.D, R.E} \left(\sigma_{R.B=S.B \wedge R.D=S.D} (R \times S) \right)$$

Q. And the names of all branches with Customers who have an account at the bank and who live in Dhanbad City.

Solⁿ Depositor (CN, AN) → for accounts

Customers (CN, CS, CO) → for City

Account (AN, BN, BAL) → for Branch Name

[Customer $\bowtie$ Depositor] [CN, CS, CC, AN]

$\bowtie$ Account [CN, CS, CC, AN, BN, BAL]

$$\pi_{BN} \left(\sigma_{CC="Dhanbad"} \left((Customer \bowtie Depositor) \bowtie Account \right) \right)$$

Division

A	B
x	1
x	2
x	3
y	1
s	1
s	3
s	4
e	6
e	1
b	2

B
1
2

A
x
b

⇒ for all values of B which A values we have tuple is the answer

$$R \div S = \pi_{R-S} (R) - \pi_{R-S} \left(\left(\pi_{R-S}(R) \times S \right) - \pi_{R-S,S}(R) \right)$$

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

B
b ₁
b ₂

→

A
a ₁
a ₅

B
b ₁

→

A
a ₁
a ₂
a ₃
a ₅

B
b ₁
b ₂
b ₃

→

A
-
-
-

B
-
-

→

A
a ₁
a ₂
a ₃
a ₅

Q. And names of all Customers who have an account at all branches at Dhanbad City.

Solⁿ Branch, Account, Depositor

$$\Rightarrow \pi_{BN} \left(\sigma_{BC="Dhanbad"} (Branch) \right)$$

(List of branches located in Dhanbad)

$$\Rightarrow \pi_{CN, BN}(\text{Depositor} \bowtie \text{Account})$$

$$\pi_{BN}(\tau_{BC="Dhankad"}(\text{branch}))$$

CN	BN
—	—
—	—

Extended Relational Algebra operation

- Generalized Projection $\rightarrow \pi_{f_1, f_2, \dots, f_n}(exp)$
- Outer Join
- Aggregate function

$$\pi_{A, B, C, A+B+C, 2 \times A, C/B}(x)$$

A	B	C
1	2	3
4	5	6

A+B+C	2xA	C/B
—	—	—
—	—	—

Outer Join

A	B	C
1	2	3
4	5	6
7	8	9

C	D
3	10
9	12
18	11

$x \bowtie x$ (Inner Join or Natural Join)

A	B	C	D
1	2	3	10
7	8	9	12

$x \bowtie x$ (Left Outer Join)

A	B	C	D
1	2	3	10
7	8	9	12
4	5	6	Null

$x \bowtie x$ (Right Outer Join)

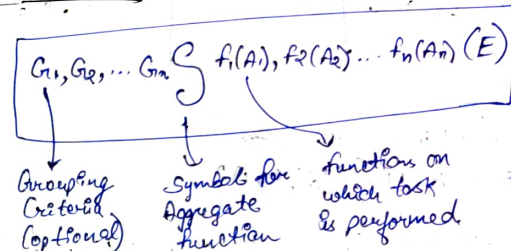
A	B	C	D
1	2	3	10
7	8	9	12
Null	Null	18	11

$x \bowtie x$ (Full Outer Join)

A	B	C	D
1	2	3	10
7	8	9	12
4	5	6	Null
Null	Null	18	11

Aggregate Function

- avg()
- max()
- min()
- count()
- sum()



$$\sum_{sum(C)}(x) \Rightarrow 27$$

A	B	C
α	α	7
α	β	7
β	β	3
β	β	10

$\sum_{sum(BAL)}(\text{account})$

BN	Sum(BAL)
Dhan Ranchi	1300
Dhan Ranchi	1500
Sum	700

Account

BN	AN	BAL
Dhan	A1	400
Dhan	A2	900
Ranchi	A3	750
Ranchi	A4	750
Plum	A5	700

Modification of the Database

- Deletion (-)
- Updation (+) (-, U)
- Insertion (+ U)

Deletion

Want to remove this tuple

A	B
1	2
3	4
5	6

$\sigma_{A=3}(R)$

A	B
3	4

$R \leftarrow R - \sigma_{A=3}(R)$

A	B
1	2
5	6

Insertion

To add (8,9) in the above Table

$$R \leftarrow R \cup \{('8', '9')\}$$

Q. Delete all account records of Dhanbad Branch.

Ans. Account (AN, BN, BAL)

account $\leftarrow$ account - $\sigma_{BN="Dhanbad"}$ (account)

Q. Delete all loan records of range (0 to 50).

Ans. Loan (LN, BN, AMT)

loan $\leftarrow$ loan - $\sigma_{AMT \geq 0 \wedge AMT \leq 50}$ (loan)

Q. Delete all account at branches located in Dhanbad city.

Updation

1 Case - If you want to select some tuple from R & to update then you can use the following

$$R \leftarrow \pi_{f_1, f_2, \dots, f_n}(\sigma_p(R)) \cup (R - \sigma_p(R))$$

Conditions (Predicate)

②

Integrity Constraints

→ Domain Constraints on Columns (check)

Student

Admno.	Name	Dob	Branches
			CSE ECE DA

check if value is in Branch table or not

Referencing Relation

→ Referential Integrity Constraint

Branch

BN	dep id	Prop id	Est
CSE			
ECE			
EC			

If value is present here

Referenced Relation

Branch id is foreign Key

→ Functional Dependencies

$$X \rightarrow Y$$

∴ Y is functionally dependent on X

[X → Y does not mean that, Y → X will also be true.]

R

X	...	Y	...
2		5	
3		5	
1		8	

Using the FD notation, we say that, 'A' is Superkey of 'R' if $K \rightarrow R$, that is K is a Superkey if whenever $t_1[K] = t_2[K]$ it is also the case that $t_1[R] = t_2[R]$, it means $t_1 = t_2$.

$R(R)$

$$R = \{A, B, C, D, E\}$$

K

$$BC \rightarrow A, B, C, D, E$$

If none of A or C can uniquely define A, B, C, D, E then BC is candidate key (minimal Super key)

Closure of FDs

A set of all FDs logically implied by F is the closure of F .

$F = \text{set of FDs}$

$F^+ \Rightarrow \text{Closure} \Rightarrow \{F\}$

$\Rightarrow$ If α is set of attributes & $\beta \subseteq \alpha$, then β is FD dependent on α (Reflexivity)

$\alpha = \{A, B, C, D\}$, $\beta = \{B, C\}$
 $\alpha \rightarrow \beta$

$\Rightarrow$ If $\alpha \rightarrow \beta$, then $\gamma \alpha \rightarrow \gamma \beta$ (Augmentation)

$\Rightarrow$ If $\alpha \rightarrow \beta$, $\beta \rightarrow \gamma$, then $\alpha \rightarrow \gamma$ (Transitivity)

$\Rightarrow$ If $\alpha \rightarrow \beta$ and $\alpha \rightarrow \gamma$, then $\alpha \rightarrow \beta \gamma$ (Union)

$\Rightarrow$ If $\alpha \rightarrow \beta \gamma$, then $\alpha \rightarrow \beta$ & $\alpha \rightarrow \gamma$ (Decomposition)

$\Rightarrow$ If $\alpha \rightarrow \underline{\beta}$ and $\gamma \underline{\beta} \rightarrow \delta$, then $\alpha \gamma \rightarrow \delta$ (Pseudotransitivity)

Example:

$R = \{A, B, C, G, H, I\}$

$F = \{A \rightarrow B, A \rightarrow C, CG \rightarrow H, CG \rightarrow I, B \rightarrow H\}$

$F^+ = \{?\}$

1) $A \rightarrow H$ [$A \rightarrow B, B \rightarrow H$]

2) $AG \rightarrow H$ [$A \rightarrow C, CG \rightarrow H$]

3) $AG \rightarrow I$ [$A \rightarrow C, CG \rightarrow I$]

4) $CG \rightarrow HI$ [$CG \rightarrow H, CG \rightarrow I$]

$AG^+ = \{A, B, C, G, H, I\} = R$ [Candidate Key]

$A^+ = \{A, B, C, H\} \neq R$

$G^+ = \{G\} \neq R$

Closure of Attribute set

$\alpha \rightarrow \alpha^+ = ?$

result = α

while (change to result) do

for each FD $\beta \rightarrow \gamma$ in F do

begin

if $\beta \subseteq \text{result}$ then $\text{result} = \text{result} \cup \gamma$

end

Usage of Attr. Closure

1) Testing of Super Key

2) Testing FDs.

$A \rightarrow (H)$

$A^+ = \{A, B, C, H\}$

3) Computing closure of F

Canonical Cover

$\rightarrow$ Remove the Redundant Dependencies from F .

$F = \{A, B, C\}$

$FD = \{A \rightarrow B,$

$B \rightarrow C,$

$A \rightarrow C\}$

$\rightarrow$ Redundant

$\rightarrow$ Remove Extra Attributes

$AG \rightarrow CH$

$\hookrightarrow$ may be redundant

$F_c = F$

$F_c = \text{Canonical}$

Repeat

- use the union rule to replace any dependencies in F_c of the form

$\alpha_1 \rightarrow \beta_1$ and $\alpha_2 \rightarrow \beta_2$ with $\alpha_1 \rightarrow \beta_1 \beta_2$

- find a FD $\alpha \rightarrow \beta$ in F_c with an extraneous Attr either in α or β , then delete it from $\alpha \rightarrow \beta$

until F_c does not change.

Entireness Attribute

1) If $A \in \beta$, to check if A is entireness, consider the set,
 $F' = (f - \{A \rightarrow B\}) \cup \{A \rightarrow (B-A)\}$
 and check if $A \rightarrow A$

$$F = \left\{ \begin{array}{l} A \rightarrow B \\ A \rightarrow H \\ C \rightarrow D \end{array} \right\}, F' = \left\{ \begin{array}{l} A \rightarrow B \\ A \rightarrow H \\ C \rightarrow D \end{array} \right\}$$

$$A^+ = \{A, B, H\} \text{ not } C$$

$\Rightarrow$ To check, compute A^+ under F'
 $\Rightarrow$ If A^+ includes A , then A is entireness in β

2) If $A \in \alpha$, to check if A is entireness,
 $\gamma = \alpha - \{A\}$ and check $\alpha \rightarrow \beta$ can be inferred from F
 γ^+ under F .

Finding no. of Candidate Keys

$$R = \{A, B, C, D, E, F, G, H\}$$

$$FD = \left\{ \begin{array}{l} AB \rightarrow C \\ A \rightarrow DE \\ B \rightarrow F \\ F \rightarrow GH \end{array} \right\}$$

Dependency/Edge Diagram



Now, Attribute which do not have incoming edge :-
 $\Rightarrow A$ and B

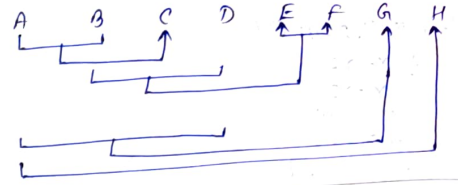
$$(AB)^+ = \{A, B, C, D, E, F, G, H\} = R \text{ (Candidate Key)}$$

$$A^+ = \{A, D, E\} \neq R$$

$$B^+ = \{B, F, G, H\} \neq R$$

(10) $R = \{A, B, C, D, E, F, G, H\}$

$$F = \left\{ \begin{array}{l} AB \rightarrow C \\ BD \rightarrow EF \\ AD \rightarrow G \\ A \rightarrow H \end{array} \right\}$$



$$(ABD)^+ = \{A, B, C, D, E, F, G, H\} = R$$

Candidate Key.

$$R = \{A, B, C, D, E\}$$

$$F = \left\{ \begin{array}{l} BC \rightarrow ADE \\ D \rightarrow B \end{array} \right\}$$



$$C^+ = \{C\} \neq R \text{ (But an essential key and must be a part of CK)}$$

$$AC^+ = \{A, C\} \neq R$$

$$BC^+ = \{A, B, C, D, E\} = R$$

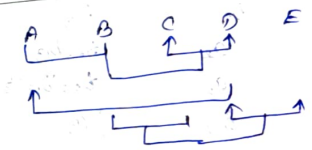
$$DC^+ = \{A, B, C, D, E\} = R$$

$$EC^+ = \{C, E\} \neq R$$

{CR}

(11) $R = \{A, B, C, D, E\}$

$$F = \left\{ \begin{array}{l} AB \rightarrow CD \\ D \rightarrow A \\ BC \rightarrow DE \end{array} \right\}$$



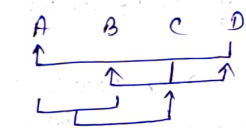
$$B^+ = \{B\} \neq R$$

$$AB^+ = \{A, B, C, D, E\} = R$$

$$BC^+ = \{A, B, C, D, E\} = R$$

$$BD^+ = \{A, B, D, A, C, E\} = R$$

(12) $R = \{A, B, C, D, E\}$, $F = \left\{ \begin{array}{l} D \rightarrow A \\ C \rightarrow BD \\ AB \rightarrow C \end{array} \right\}$



$$A^+ \neq R$$

$$B^+ \neq R$$

$$C^+ = \{A, B, C, D\} = R$$

$$D^+ \neq R$$

$$AB^+ = \{A, B, C, D\} = R$$

$$BD^+ = \{A, B, D, A, C\} = R$$

$$AD^+ = \{A, D\}$$

$\Rightarrow$ No. essential Atth.

Lossless Join Algorithm

$R = \{X, Y, Z, W, Q\}$

$F = \{$
 $X \rightarrow Z$
 $Y \rightarrow Z$
 $Z \rightarrow W$
 $WQ \rightarrow Z$
 $ZQ \rightarrow X$
 $\}$

Decomposition Set = $\{R_1(X, W), R_2(X, Y), R_3(Y, Q), R_4(Z, W, Q), R_5(X, Q)\}$

$\therefore$ whether this Decomp is lossless or not?

Step 1: Create a Table, as many ~~columns~~ ^{Attributes (5)} as many Columns.
 Rows = no. of Decomposed Relations

	X (A ₁)	Y (A ₂)	Z (A ₃)	W (A ₄)	Q (A ₅)
R ₁	a ₁	b ₁₂	b ₁₃	a ₄	b ₁₅
R ₂	a ₁	a ₂	b ₂₃ b ₁₃	b ₂₄ a ₄	b ₂₅
R ₃	b ₃₁ a ₁	a ₂	b ₃₃ b ₁₃ a ₃	b ₃₄ a ₄	a ₅
R ₄	b ₄₁ a ₁	b ₄₂	a ₃	a ₄	a ₅
R ₅	a ₁	b ₅₂	b ₅₃ b ₁₃ a ₃	b ₅₄ a ₄	a ₅

for $X \rightarrow Z$, for same values of X we should have same Z .
 so for all a₁ in Row 1 we should put same value of Z in R₃ as b₁₃.

for $Y \rightarrow Z$, Row 2 = Row 3 for same values of $a_2 \rightarrow b_{13}$

for $Z \rightarrow W$, for b_{13} in Row 3 put same a_4 in Row 4 (W).

for $WQ \rightarrow Z$,

R₃, a₄ a₅ $\rightarrow$ b₁₃ a₃
 R₄, a₄ a₅ $\rightarrow$ a₃
 R₅, a₄ a₅ $\rightarrow$ b₁₃ a₃

(14)

for, $ZQ \rightarrow X$

R₃, a₃ a₅ $\rightarrow$ b₃₁ a₁
 R₄, a₃ a₅ $\rightarrow$ b₄₁ a₁
 R₅, a₃ a₅ $\rightarrow$ a₁

when we got a's & b's then replace b's with a's

(15)

After operations, if any Row has all a's value then Relation/Decomposition is lossless.

Normalization

Indicators,

- 1) Redundant data
- 2) Insertion Anomalies
- 3) Updation Anomalies
- 4) Deletion Anomalies
- 5) NULL Handling

- 1) 1NF $\rightarrow$ Domains are Atomic
- 2) 2NF $\rightarrow$ 1NF + No Partial dependency
- 3) 3NF $\rightarrow$ 2NF + No Transitivity depen.
- 4) BCNF $\rightarrow$ 3NF + Determinant (LHS) should be candidate key
- 5) 4NF - BCNF + not more than 1 MVD
- 6) 5NF - 4NF + No Join Dependence (PJNF)

(Project Join NF)

2NF

$\rightarrow R$ is in 2NF $\Leftrightarrow [(R \text{ should be in 1NF}) \wedge (\text{Not Contain any Partial dependency})]$

$\Rightarrow X \rightarrow Y$, is fully dependent, iff X is not a proper subset of Candidate Key or Y is a prime Attribute.

$$[(X \not\subseteq CK) \vee (Y \text{ is PA})]$$

$\Rightarrow X \rightarrow Y$, is partially dependent $\Leftrightarrow [(Y \subseteq CK) \wedge (X \text{ is Not PA})]$

Q. $R = \{A, B, C, D\}$, $f = \{AB \rightarrow D, B \rightarrow C\}$ $\rightarrow CK$
 Solⁿ $(AB)^+ = ABCD = R = CK$
 $PA = \{A, B\}$, $NPA = \{C, D\}$

$AB \rightarrow D$, Full dep.
 $B \rightarrow C$, Partial dep.
 $\hookrightarrow PA$

3NF R is in 3NF, iff $\forall FD_{NT}$ (for all FD which are non-trivial)
 $C \rightarrow Y, [(C \text{ is Superkey}) \vee (Y \text{ is PA})]$

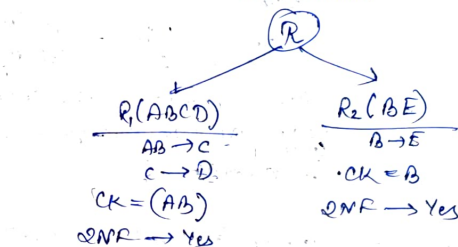
$\Rightarrow (X \rightarrow Y \text{ is Transitive Dep}) \Leftrightarrow (('X' \text{ is not a SK}) \wedge ('Y' \text{ is NPA}))$

Example:-

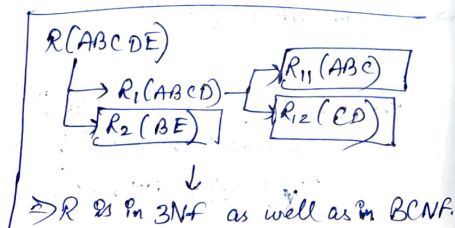
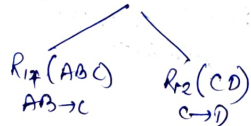
$R = \{A, B, C, D, E\}, F = \{ \begin{matrix} AB \rightarrow C \\ C \rightarrow D \\ B \rightarrow E \end{matrix} \}$

Solⁿ: $AB \rightarrow \{A, B, C, D, E\}$ $\left\{ \begin{matrix} B \rightarrow E \\ \text{CK} \end{matrix} \right\}$ in 3NF. $\rightarrow PA$ and D is not PA, Not in 2NF

$\therefore$ So relation is not in 2NF.



Because of $(C \rightarrow D)$,
 $\Rightarrow R_2$ is not in 3NF



4NF $\rightarrow$ Remove all MV Dependencies

Multivalued Dependency

$X \twoheadrightarrow Y$ (for same value of X we have multiple values of Y)

Three Necessary Conditions for MVD:-

- 1) $A \twoheadrightarrow B$, for a single value of A , more than 1 value of B exists.
- 2) Table/Relation should have atleast 3 columns
- 3) for a table with column A, B, C and there is MVD b/w A & B , then B and C should be independent of each other.

A	B	C
A ₁	B ₁	C ₁
A ₁	B ₂	C ₂

$A \twoheadrightarrow B$
 $A \twoheadrightarrow C$

Ex:-

Adm no.	Course	hobby
1	ADDBMS	chess
1	ADSA	Tennis
2	-	-

$\rightarrow$

Adm	Course	hobby
1	ADDBMS	Tennis
1	ADSA	chess

$\Rightarrow$ It will give two more rows in Table.

$\Rightarrow$ Its a bad design, because there is no relation b/w Course and hobby.

Decomposition $\rightarrow R_1(\text{Adm, Course})$
 $R_2(\text{Adm, hobby})$

$\Rightarrow$ MVD occurs if two or more independent relations are kept in a single relation.

$\bar{A} \twoheadrightarrow \bar{B}$
 $\forall t, u \in R: t[\bar{A}] = u[\bar{A}] \text{ then}$
 $\exists v \in R: v[\bar{A}] = t[\bar{A}] \text{ and } v[\bar{B}] = t[\bar{B}] \text{ and } v[\text{rest}] = u[\text{rest}]$

	A	B	rest
t	a	b ₁	x ₁
u	a	b ₂	x ₂
v	a	b ₁	x ₂
w	a	b ₂	x ₁

MVD are also known as
 Tuple generating Dependencies.

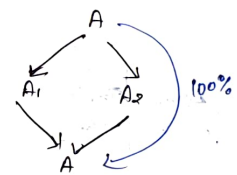
Admno	Phone	email
e3	1	c1
e3	1	c2
e3	1	c3
	2	c4
	2	c5

Adm	Phone
x	p1
x	p2
x	p3
x	p4
x	p5

Adm	email
x	c1
x	c2
x	c3
x	c4
x	c5

$A \twoheadrightarrow B$, $A \twoheadrightarrow C$

MVD



Join
 Dependency

Ex:- R

A	B	C
1	2	1
2	2	2
3	3	2

So PTNP/5NF says choose common attribute as key or part of key like A in this example.

$R_1 (A, B)$

A	B
1	2
2	2
3	3

$R_2 (B, C)$

B	C
2	1
2	2
3	2

Join

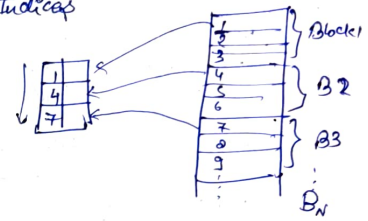
A	B	C
1	2	1
2	2	2
3	3	2

→ Common Attr. should be candidate or super key of either relations (R, R₁, R₂), then Join will be lossless.

Indexing

- Ordered Indices
- Hashed Indices

Ordered:-



- 1) Access Time:- specific value, Range of values (2 to 200)
- 2) Access Time:- Individual access time taken.
- 3) Insertion Time:- Time to insert in DB (indexes). Making Index Table up to the mark.
- 4) Deletion Time:- After deleting record, updation of indexes.
- 5) Space Overhead:- Space/Memory to store indexes

Ordered Indices

- Dense Index
- Sparse Index

All keys are present

Key	Pointer
B	A1
D	A2
M	A4
P	A5
R	A9
U	A10

Dense

Key	Pointer
B	A1
M	A4
R	A8

Sparse

AN	BN	Col
A1	B	-
A2	D	-
A3	D	-
A4	M	-
A5	P	-
A6	P	-
A7	P	-
A8	R	-
A9	R	-
A10	U	-

∴ Pointer will point to the 1st Occurrence of key.

Trade-Off B/w Both results into dividing the Relation into Blocks and store first value of each block in the Sparse Index table.

Types of Indices

- Primary
- Clustered
- Secondary

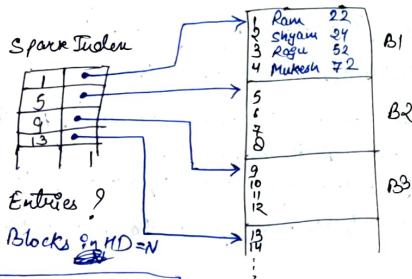
Ordered Type (Sorted)
Unordered Type

Primary Index	Clustered Index
1, 2, 3, 4, 5, 6...	1, 1, 2, 3, 3, 3, 4...
Secondary Index	Secondary Index
4, 0, 1, 2, 6, 9, 3...	3, 1, 3, 2, 1, 6, 1, 2...

Key (only Unique values)

Non-Key

Primary Index



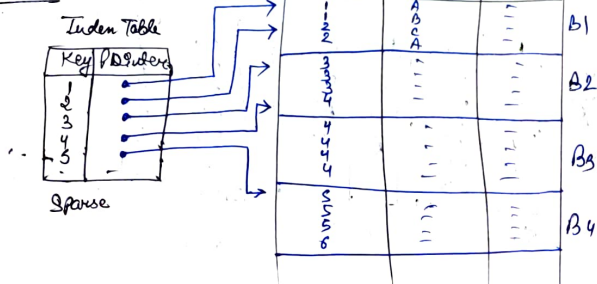
Q. How many Entries?

Ans. Number of Blocks in HD = N

$$\text{Search Time} = \log_2 N + 1$$

where N = no. of Blocks in Index Table.
 $\Rightarrow O(\log_2 N)$

Clustered Index



When searching in reg. for 4 value then, we require to access two Blocks (B2 & B3) of HD and this is called Block Hopper.

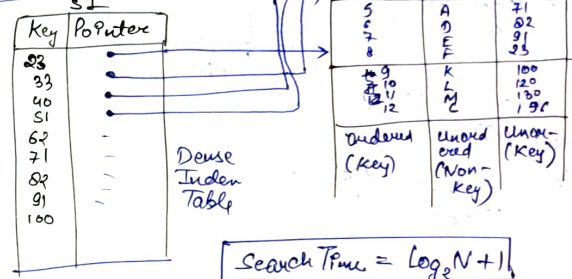
$$\text{Search Time} = \log_2 N + 1 + 1$$

For Index Table (Index Tree)

For another Block (may search more than 1 Block due to Block Hopper)

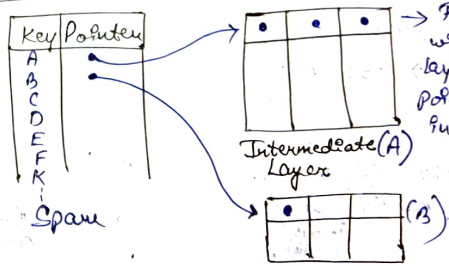
Secondary Index

For Unordered + Non-Key (PAN)



$$\text{Search Time} = \log_2 N + 1$$

For Unordered + Non-Key (Name)



For multiple value, we will create a Intermediate layer and each column will point to all locations of value in the HD.

$$\text{Overall Index} = \text{Dense}$$

$$\text{Search Time} \Rightarrow \log_2 N + 1 + 1$$

7th Access

Intermediate Layer

Main IT

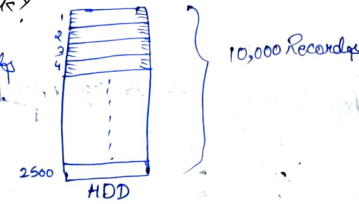
Ex: Find out the avg Time Comp. to search a record from the HD. Given the following:-

Total Records = 10000
 Record Size = 250 Bytes/Record
 Block Size = 1000 Bytes/Block
 Data entered in HDD without any order (Unordered)

1) What is the no. of Records/Block?

Ans $\frac{\text{Block Size}}{\text{Record Size}} = \frac{1000}{250} = 4 \text{ records}$

No. of Blocks = $\frac{10000}{4} \Rightarrow 2500$



2) Sorted Data

Sol Binary search (Indexing)

Best Case = $O(1)$

Worst Case = $O(\log_2 N) \Rightarrow \log_2(2500)$
 $\Rightarrow 12 \text{ Blocks}$

3) Unsorted Data

Sol Linear Search

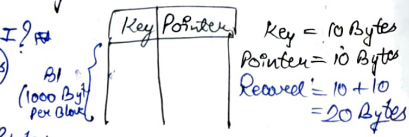
Best Case = $O(1)$

Worst Case = $O(N)$
 $= 2500 \text{ Blocks}$

Average Case = $O(N/2) \Rightarrow 1250 \text{ Blocks}$

4) Sorted Data + Sparse Indexing.

Q: How many no. of Entries in IT?
 $\Rightarrow 2500$ (Same as no. of Blocks)



One Block of Index Table $\Rightarrow 1000 \text{ Bytes}$

One entry Size of Index T. $\Rightarrow 20 \text{ Bytes}$

No. of entries in 1 Block of IT $\Rightarrow \frac{1000}{20} \Rightarrow 50$

No. of Blocks in IT $\Rightarrow \frac{2500}{50} \Rightarrow 50 \text{ Blocks}$

Total Search Time $\neq \log_2 n + 1$
 $\Rightarrow \lceil \log_2(50) \rceil + 1$
 $\Rightarrow 7$

5) Unsorted + Dense Indexing

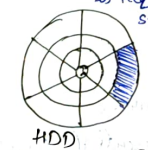
No. of Blocks in IT $\Rightarrow \frac{10000}{50} \Rightarrow 200 \text{ Blocks}$

TST $\Rightarrow \log_2(200) + 1$
 $\Rightarrow 9$

Roll List (10) (30) (10) (8) (50) $\Rightarrow 120 \text{ Bytes}$

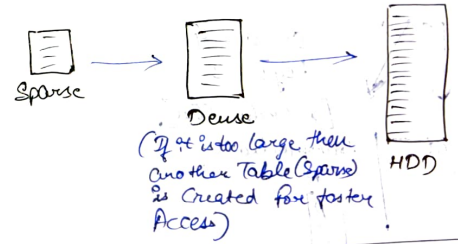
Admno.	name	Prog.	Branch	Dept
...	...	...	...	...

100 Records $\rightarrow$ 25 Blocks of HDD
 20 MB to store



HDD Block Size = 512 Bytes

1 Record Size = 128 Bytes



Index Updation

Insertion

- $\Rightarrow$ Insert a Record to HD
- $\Rightarrow$ If data is already present then go to the pointer & to the last replication and insert the record.
- $\Rightarrow$ If not then insert and add entry to IT.

⇒ Sparse

- ⇒ Entry is present then add to HDD
- ⇒ Not present, then Add to HDD & then in Index Table
- ⇒ If HDD full, then Create Block.

⇒ Deletion

- ⇒ First go to the pointer, then delete all repetitions & then } Dense
- Delete entry from Index Table.

For Sparse

- ⇒ Search in the HDD & delete the repetitions.
- ⇒ If the entry is present in Sparse IT then, after deleting increment the Sparse entry to the next Element.

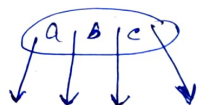
Multilevel Searching

B Tree / B⁺ Tree

- Balanced
- Multiway Search
- Random

B⁺ Tree

- Balanced
- Multiway Search
- Random & Sequential Search



Order = 4
Keys = 3

P-Order BTree

Children	Root	Intermediate
max	P	P
min	2	$\lceil \frac{P}{2} \rceil$

(24)

Order = 5

1, 7, 6, 2, 11, 4, 8, 13, 10, 5, 19, 9, 18, 24, 3, 12, 14, 20, 21, 16

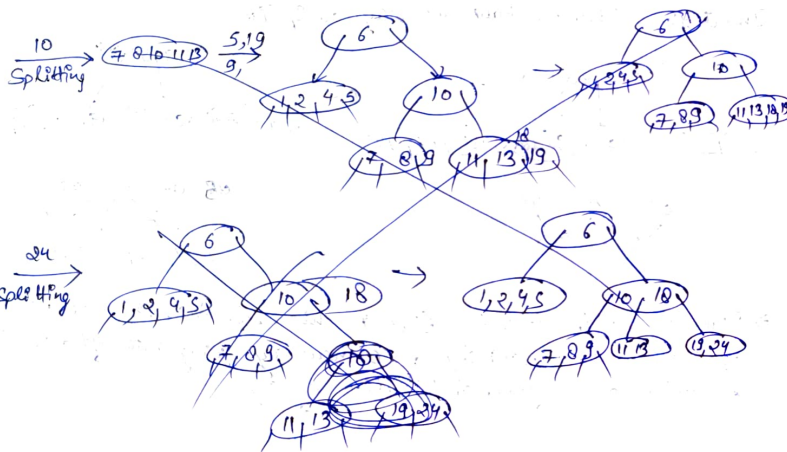
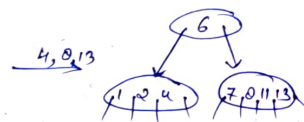
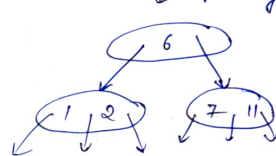
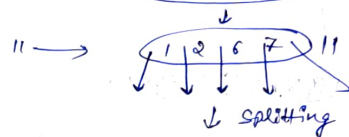
max children = 5

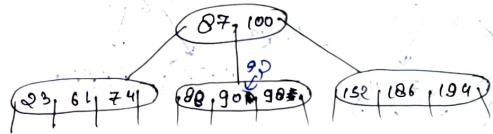
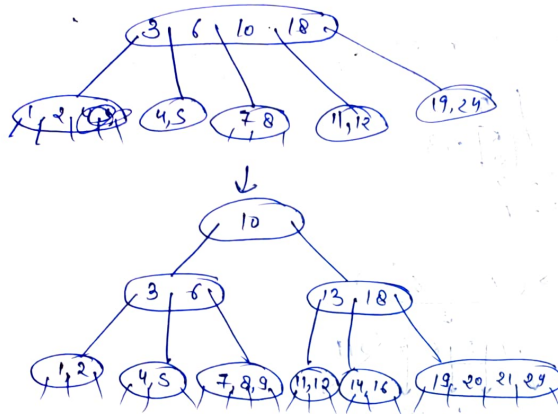
min children = 2 $\lceil \frac{5}{2} \rceil = 3$

max keys = 5 - 1 = 4

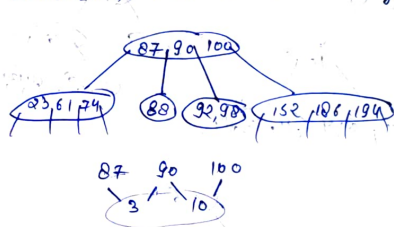
min keys = 2 - 1 = 1

1, 7, 6, 2 → (1 7 6 2)

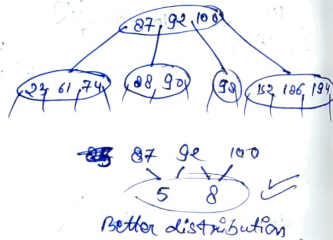




Insert 92 → Then we have two ways to split



or



⇒ Thus we would select left bias in this case. Since, it more nearly equalizes the probability of a new key going into left & right. Assuming a uniform distribution of keys.

86

B⁺ Tree :-

→ Repetition of some keys.

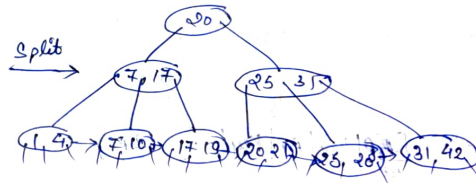
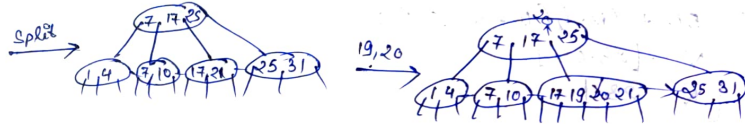
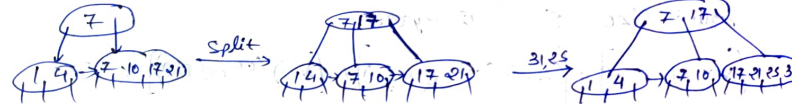
→ All queries are resolved after going to the leaf and all leaf are connected. (Sequential Search).

Insertion in B⁺ Tree

1, 4, 7, 10, 17, 21, 31, 25, 19, 20, 28, 42 (Order=4)

Max	4	4	3
Min	2	2	1

Key



Q Consider a B Tree of Key Size = 10 Bytes, Block Size = 512 Bytes.
Data Pointer Size = 8 bytes, Block Pointer Size = 5 bytes.
Find the Order of B Tree.

Sol:

BP	Key Value	DP	BP	Key Value	DP
----	-----------	----	----	-----------	----

 ...

BP

Let no. of Block Pointer = N

Total size of Block Pointers = $5N \rightarrow ①$

Number of Keys $\Rightarrow N-1$

Total size of Keys $\Rightarrow 10(N-1) \rightarrow ②$

Number of DP $\Rightarrow N-1$

Total size of DP $\Rightarrow 8(N-1) \rightarrow ③$

Block Size ≤ 512 Bytes

$$[N \times \text{Size of BP}] + [(N-1)(\text{Size of Keys})] + [(N-1)(\text{Size of DP})] \leq 512$$

$$\Rightarrow 5N + (N-1)(10+8) \leq 512$$

$$\Rightarrow N = \lfloor 23.04 \rfloor$$

$$N = 23$$

Q Consider B⁺ Tree,

Key Size = 10 Bytes, Block Size = 512 B, Data P size = 8 bytes,
BP size = 5 Bytes.

What is the Order of Leaf & non-leaf nodes?

Sol for non-leaf Node $\Rightarrow$

BP	Key	BP	Key	...	BP
----	-----	----	-----	-----	----

$$5N + (N-1)10 \leq 512$$

$$15N \leq 522$$

$$N = 34.8$$

$$N = 34$$

for leaf $\Rightarrow$

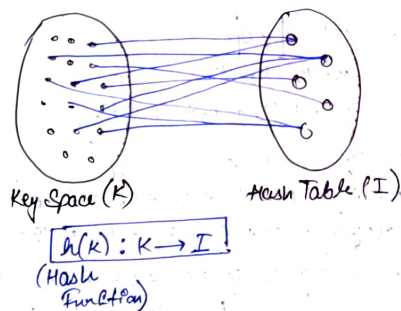
BP	Key	BP	Key	DP	...	BP
----	-----	----	-----	----	-----	----

$$5 + (N-1)(10+8) \leq 512$$

$$N = 28$$

Hashing

$\hookrightarrow$ Hashed Indices



Hash function

10, 19, 35, 43, 62, 59, 31, 49, 77, 33

$h(K) = \text{sum of digits}$

$$10 \rightarrow (1+0) \% 10 = 1$$

$$19 \rightarrow (1+9) \% 10 = 10$$

$$h(K) = \left(\sum_{i=0}^n d_i \right) \% 10$$

0	19
1	10
2	
3	49
4	59, 31, 77
5	
6	33
7	43
8	35, 62
9	

Collisions $\rightarrow$

Most Used Hash function

1) Division Method :-

$$H(K) = K \bmod h$$

$$\text{Index} \begin{matrix} 0 \\ 1 \end{matrix} \quad \begin{matrix} H(K) = K \% h \\ H(K) = K \div h + 1 \end{matrix}$$

$h = \text{no. of spaces in hash table}$

$\Rightarrow$ Mid Square Method :-

$$K = 1234$$

$$K^2 = 1522756$$

$\uparrow \uparrow \uparrow$
Store at (525)

3. Folding Method

$$K = 15 | 22 | 75 | 6$$

$$\rightarrow K_1 + K_2 + K_3 + K_4$$

$$\Rightarrow K_1 + K_2 + K_3 + K_4 \Rightarrow R_{sum}$$

$$\text{Chopping} \Rightarrow 01 \ 52 \ 27 \ 56$$

$$\text{Pure folding} \Rightarrow 01 + 52 + 27 + 56 \Rightarrow 136$$

$$\text{Fold shifting} \Rightarrow 10 + 52 + 72 + 56 \Rightarrow 190$$

$$\text{Fold Boundary} \Rightarrow 10 + 52 + 27 + 55 \Rightarrow 154$$

(First & last no. folded)

136
152 27 56

Collision Resolution Technique

Open Hashing
(Chaining)

Closed Hashing
(Open Addressing)

42	19	10	12
----	----	----	----

$$h(K) = K \% 5$$

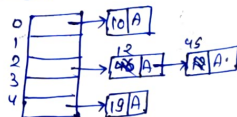
$$42 \% 5 = 2$$

$$19 \% 5 = 4$$

$$10 \% 5 = 0$$

$$12 \% 5 = 2$$

HT



$\Rightarrow$ Insert at front of Linked list
 $\Rightarrow$ Disadvantages, if all keys mapped to same location (Extra space)

Disadvantage $\Rightarrow$ Primary Clustering
 $\Rightarrow$ Secondary Clustering

$\Rightarrow$ Linear time ($O(n)$), if most spaces are filled
 $\Rightarrow$ Deletion Difficulty

$\Rightarrow$ No Extra Space

Linear Probing

$$h(K) = K \bmod n$$

$$h'(K, i) = (h(K) + i) \bmod n$$

(in case of Collision)

$i = \text{collision/probe number}$

43	135	72	23	99	19	82
----	-----	----	----	----	----	----

$$n = 10$$

$$h(K) = K \bmod 10$$

$$h'(K, i) = (h(K) + i) \bmod 10$$

for, 23

$$h(K, i) = (h(23) + i) \% 10$$

$$\Rightarrow 4$$

for 19,

$$h(19, i) = (h(19) + i) \% 10$$

$$\Rightarrow 20 \% 10$$

$$\Rightarrow 0$$

Same for 82,

0	19
1	
2	72
3	43
4	23*
5	135
6	82
7	
8	
9	99
10	

Disadvantages

$\Rightarrow$ if two or more Key are following same probing sequence, thats called Secondary Clustering.

$\Rightarrow$ Primary Clustering

$\Rightarrow$ In worst Case searching time is $O(n)$

$\Rightarrow$ Deletion is difficult.

Advantages

$\Rightarrow$ No Extra Space.

Quadratic Probing

$$h(K) = K \bmod n$$

$$h(K, i) = (h(K) + i^2) \bmod n$$

42	16	91	33	18	27	36	62
----	----	----	----	----	----	----	----

for 26,

$$1) (h(26) + 1^2) \% 10 \Rightarrow 7, \text{ Collision}$$

$$2) (h(26) + 2^2) \% 10 = 0$$

for 62,

$$1) (h(62) + 1^2) \% 10 \Rightarrow 3, \text{ Collision}$$

$$2) (h(62) + 2^2) \% 10 \Rightarrow 6, \text{ Collision}$$

$$3) (h(62) + 4^2) \% 10 \Rightarrow 2, \text{ Collision}$$

So on

Not getting space for some keys, even if space are there, This is biggest disadvantage.

Disadvantages

$\Rightarrow$ Search $O(n)$

$\Rightarrow$ Secondary Clustering

$\Rightarrow$ No guaranteed of space in HT

Advantages

1) No extra space

2) Primary clustering is not here

0	36
1	91
2	42
3	33
4	
5	
6	16
7	27
8	18
9	

Double Hashing (Rehashing)

$$h_1(K) = K \bmod 11$$

$$h_2(K) = 8 - (K \bmod 8)$$

$$[h_1(K) + i h_2(K)] \bmod 11$$

20	34	45	70	56
----	----	----	----	----

for 45, (collision)

$$h_1(K) = 1$$

$$h_2(K) = 8 - (45 \% 8) \Rightarrow 3$$

$$\Rightarrow (h_1(K) + i h_2(K)) \% 11$$

$$\Rightarrow (1 + (1)(3)) \% 11$$

$$\Rightarrow 4$$

0	
1	34
2	
3	56
4	45
5	
6	70
7	
8	
9	20
10	

for 70, Collision

$$h_1(K) = 4$$

$$h_2(K) = 8 - (70 \% 8) \Rightarrow 2$$

$$(4 + 1(2)) \% 11 \Rightarrow 6$$

for 56, Collision

$$h_1(K) = 1$$

$$h_2(K) = 8 - (56 \% 8) \Rightarrow 8$$

$$(1 + 1(8)) \% 11 \Rightarrow 9, \text{ Collision}$$

$$(1 + 2(8)) \% 11 \Rightarrow 6, \text{ Collision}$$

$$(1 + 3(8)) \% 11 \Rightarrow 3,$$

Advantages

→ No Primary & Secondary Clustering.

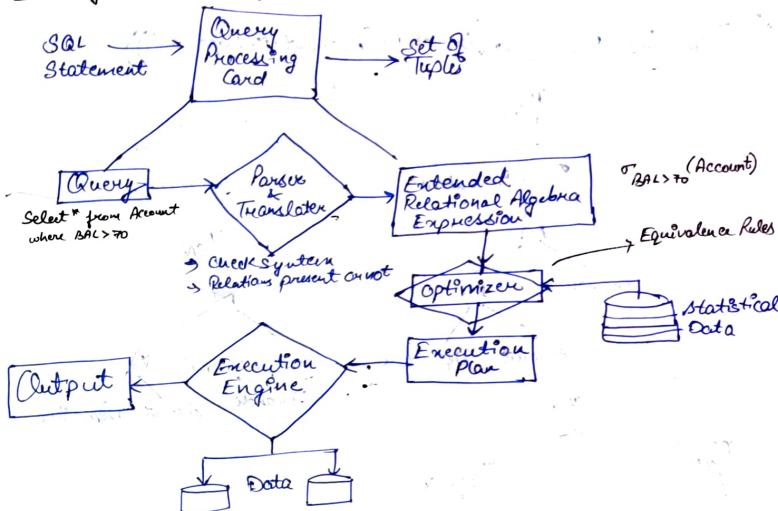
→ Always get space, if present

→ No extra space

Disadvantages

→ No Linear time $O(n)$ Searching.

Query Processing & Optimization



Account

AN	BN	BAL
A1	A	2000
A2	B	2000
A3	A	6000
A4	B	7000
A5	C	1000

π_{BAL}

$\sigma_{BAL > 2000}$

Account

AN	BN	BAL
A2	B	3000
A3	A	6000
A4	B	7000

$\pi_{BAL > 2000}$

π_{BAL}

Account

BAL
3000
6000
7000

More Optimized

Account

AN	BN	BAL

Bytes/Record

→ How many Records

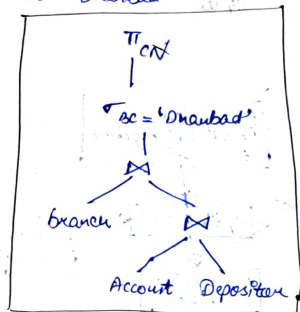
→ How many Blocks?

→ Records/Block

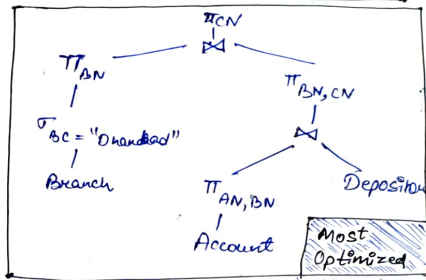
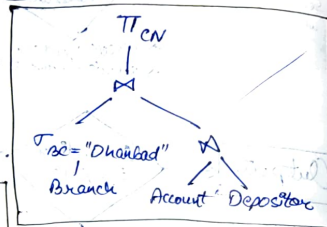
→ (AN, Account) = Change Value Count

Q. Find the names of all Customers who have an account at any branch located in Dhanbad.

Sol: $\pi_{CN} (\sigma_{BC = 'Dhanbad'} (branch \bowtie (Account \bowtie Depositor)))$



Branch (BN, BC, Ass ests)
Account (AN, BN, BAL)
Depositor (CN, AN)



Selection Operation (σ)

→ File Scanning [$\sigma_p(12)$]

Basic Algorithms

- Linear Search (A1)
- Binary Search (A2)

Selection Using Indices

- Primary Index, Equality on Key
- Primary Index, Equality on Non-Key
- Secondary Index, Equality on Key
- Secondary Index, Equality on Non-Key.

(2)

Selection Involving Comparison

- Primary Index Comparison.
- Secondary Index Comparison

More Complex Selection Predicate.

- Conjunction (Series of AND ($\wedge$))
- Disjunction (Series of OR ($\vee$))
- Negation (NOT)

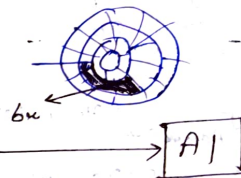
Cost Involvement

- Disk Scanning Time (DSTS) → Seek Time
- Block Transfer Time (BT) →

Let, n blocks to be transferred,

$$Cost_{LS} = TS + b_n \cdot TB \quad (\text{Linear Search})$$

$$Cost_{Avg} = TS + \frac{b_n \cdot TB}{2} \quad (\text{Equality on Key})$$



→ If the file is ordered on a attribute and selection condition is Equality Comparison on the Attribute.

$$Cost_{BS} = \lceil \log(b_n) \rceil \cdot (TS + TB) \quad (\text{Binary Search}) \rightarrow A2$$

Selection Operation

A1: Linear Search [$Cost = TS + b_n \cdot TB$]

A2: Binary Search [$Cost = \lceil \log_2(b_n) \rceil \cdot (TS + TB)$]

A3: Primary Index with Equality on Key

~~Student record~~ (B+ Tree)

$$Cost_{LS} = [TS + b_n \cdot TB + TS + TB]$$

$$Cost_{BS} = [\lceil \log_2(b_n) \rceil + TS + TB]$$

$$[Cost = (h_i + 1) \cdot (TS + TB)]$$

A4: Primary Index (Non-Key) [$Cost = h_0(TS + TB) + TS + n \cdot TB$]

A5: Secondary I (Key) [$Cost = \lceil \log_2(b_n) \rceil \cdot (TS + TB) + TS + TB$] (Same as A3)

A6: SI with Non-Key [$Cost = h_0(TS + TB) + n(TS + TB)$] [Key is on n multiple log]

Index Table

Key	Address
1	a ₁
2	a ₂
3	a ₃
4	a ₄
5	a ₅
6	a ₆
7	a ₇
8	a ₈
9	a ₉
10	a ₁₀

Dense Index

$T_{Adm} = 7$ (Student Record)

$$Cost_B = h_2(T_S + T_B)$$

$$+ (T_S + T_B)$$

$$Cost_B = (A+1)(T_S + T_B)$$

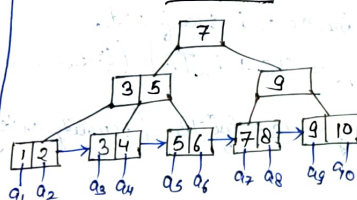
$$Cost_{B_s} = \log_2(b_4)(T_S + T_B) + (T_S + T_B)$$

$$Cost_{T_S} = (T_S + b_4 T_B) + (T_S + T_B)$$

Student Record

Address	Admno.	Name	Marks
a ₁	1	Ankit	20
a ₂	2	Shyam	40
a ₃	3	Roni	31
a ₄	4	Robin	69
a ₅	5	Risabh	82
a ₆	6	Anita	73
a ₇	7	Prakash	43
a ₈	8	Meera	52
a ₉	9	Ramjan	39
a ₁₀	10	Mamaj	28

B+ Tree



Selection Using Comparison

Algo ~~1~~ (Ordered)

I. $A > V$ or $A \geq V$

$\sigma_{Admno. > 5}$ (Student Record)

$\sigma_{Admno. \geq 5}$ (SR)

II. $A < V$ or $A \leq V$

$$Cost = T_S + n T_B$$

The SI provides pointers to the Records but to get actual Records we have to fetch the records by using pointers.

$$Cost = h_2(T_S + T_B) + (T_S + T_B)$$

No. of Blocks to be moved

for Secondary Index, B+ Tree, Comparison

$$Cost = h_2(T_S + T_B) + n(T_S + T_B)$$

The lowest level index are scanned either from the smallest value up to V or from V up to max. value

A₆ and A₇

More Complex Selection Predicate

A₈: Conjunctive Selection using one Index

$$\sigma_{Q_1 \wedge Q_2 \wedge Q_3 \wedge \dots \wedge Q_n(x)}$$

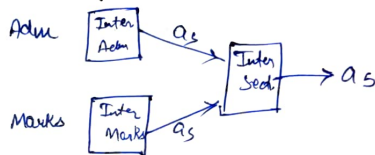
Select * from student record where Adm > 5 and Adm < 8

⇒ After Applying $Adm > 5$, then second Cond. will only be applied on the resulting data (Attribute must be same).

A₉: Conjunctive Selection using Composite Index

Select * from SR where Adm = 5 and marks = 82

A₁₀: CS by Intersection of Identifier



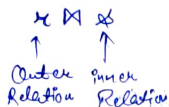
Marks	Address
20	a ₁₀
28	a ₃
31	a ₉
39	a ₂
40	a ₇
43	a ₈
52	a ₆
62	a ₄
73	a ₅
82	a ₁

A₁₁: ~~CS~~ Disjunctive Selection by union of Identifier

$$\sigma_{Q_1 \vee Q_2 \vee Q_3 \vee \dots \vee Q_n(x)}$$

A₁₂: Negation (¬)

Join Operation



$R \bowtie_{R.A=S.B} S \rightarrow$ Equi Join

Depositor $\bowtie$ Customer
(AN, CN) (CN, CS, CC)

Number of Records in Depositor = 5000 & Customer = 10,000
" " Blocks " " = 100 & " = 400

Nested Loop Join

For each tuple t_x in 'R' do begin

for each tuple t_s in 'S' do begin

- test $pair(t_x, t_s)$ to see if they satisfy the Join Condition
- If they do, add (t_x, t_s) to the result

end

end



Comparisons = $n_R \times n_S$

Worst Condition

- $\Rightarrow$ 1 Block of each is allowed at a time
- $\Rightarrow$ Total no. of Block Transfer = $b_R + n_R b_S$
- $\Rightarrow$ Total seeks = $b_R + n_R \times 100$

Best Case

- $\Rightarrow$ Block Transfer = $b_R + b_S$ [Bring all Blocks of R & S into memory]
- $\Rightarrow$ Seek Time = $1 + 1 \Rightarrow 2$

Block Nested Loop Join

for each Block B_R of 'R' do Begin

for each Block B_S of 'S' do Begin

for each tuple t_x in 'R'

for each tuple t_s in 'S'

- test pair (t_x, t_s) to see if they satisfy the Join Condition
- If they do, add (t_x, t_s) to the result

end

end

(38)

Worst Case

$$\text{Total BT} = b_R + b_R + b_S = 40100$$

$$\text{Total ST} = b_R + b_R = 200$$

Improvements

- # If the Join Attribute form Key on the inner relation, then for each outer relation tuple, the inner loop can terminate as soon as the first match is found.
- # Instead of using disk blocks as the blocking unit for outer relation we can use the biggest size that can fit in memory while leaving enough space for the buffers of the inner relation & outer.
- # Memory has m blocks.
Read $(m-2)$ blocks of outer relation at a time.
Each block of inner relation will join with $(m-2)$ blocks of outer relation.
Reduces the no. of scan of inner relation from $b_R \rightarrow \left\lceil \frac{b_R}{m-2} \right\rceil$

$$\text{Total Block Transfer} = b_R + \left\lceil \frac{b_R}{m-2} \right\rceil \times b_S$$

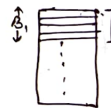
$$\text{Total Seek Time} = 2 \left\lceil \frac{b_R}{m-2} \right\rceil$$

- # Inner loop can be scanned alternatively forward & Backward.
This scanning method orders the request for disk blocks so that the data remaining in the buffer from the previous scan can be reused. Thus, the reducing the no. of disk accessing.

Indexed Nested-Loop Joins

$$\text{Total Cost} = b_R \left(\frac{1}{b_S} + c \right) + n_R \times c$$

where, c = The Cost of single
Selection on S using join condition



Merge Join

a	3	A
b	1	G
d	8	N
d	13	N
m	5	B

$R_1 \rightarrow$

x	a_1	a_2
	a	3
	b	1
	d	8
	d	13
	f	7
	m	5
	l	6

$R_2 \leftarrow$

a_1	a_2
a	A
b	G
c	L
d	N
m	B

- 1) It associates 1 pointer with each relation.
- 2) Pointer points to the initially tuple of resp. relation.
- 3) Pointer value will be incremented only if there is an unequal occurrence.

Cost = $b_n + b_b$ Seek time = 2

4) If either of the input relation is unsorted on a Join Attribute, they can be sorted first before applying the merge join Algorithm.

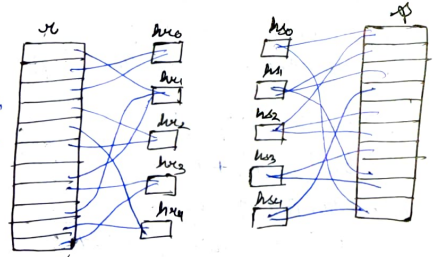
Seek time (worst case) = $\left\lceil \frac{b_n}{b_b} \right\rceil + \left\lceil \frac{b_s}{b_b} \right\rceil$ (Total Block Cop = b_b)

If $b_n = b_s = b_b$,

ST = 2

Hash Join

- 1) Partition the record.
- 2) One by one take partition.
- 3) Build the Hashed Index & build input.
- 4) Probe the tuples from other relation.



Cost = $(b_n + b_s) + (b_n + b_s) + (b_n + b_s)$

[loading all blocks of x & y in m/m for Hashing/Partitioning] [Back to HD after Partitioning] [for build and probe]

Seek time = $2(b_n + b_s) + 2n_h$ [n_h = no. of Hash partition of the relation]

Hash Join

Eid	Name	Salary	Did
10	P	1000	1
11	Q	2000	1
12	R	3000	1
13	S	4000	2
14	T	5000	2
15	U	6000	3
16	V	7000	3
17	W	8000	3
18	X	9000	4

Did	Name
1	A
2	B
3	C
4	D
5	E

We will apply Hash Join,

$H(x) = (Did \% 3)$

Eid	Name	Sal	Did
15	U	6000	3
16	V	7000	3
17	W	8000	3
10	P	1000	1
11	Q	2000	1
12	R	3000	1
18	X	9000	4
13	S	4000	2
14	T	5000	2

h_{x0}

Did	Name
3	C

Bucket 0

Bucket 1

h_{x1}

Did	Name
1	A
4	D

Bucket 0

Bucket 1

h_{x2}

Did	Name
2	B
5	E

Bucket 0

Bucket 1

Probe Relation

Build Relation (Keep in memory)

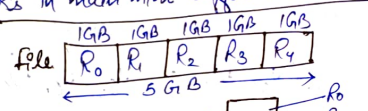
Sorting Operation

- Internal Sorting: If entire relation can be accommodate into m/m and sorting is applied. (Heap, Quick, merge, etc.)
- External Sorting: If the relation is large and can't fit into m/m.

External Sort Merge

Let M denote the no. of blocks in main m/m Buffer available for sorting,

$M = 1 \text{ GB}$
File size = 5GB



Read 150MB of data from each file (R_0, R_1, R_2, R_3, R_4)
Total = 750MB (150x5)
Left M/m $\Rightarrow$ 250MB (for output)

Ex

S.id	Name
1	Ravi
2	Amit
3	Pawan
4	Magav
5	Neha
6	Rahul
7	Abhishek
8	Rajesh
9	Santosh
10	Jaya
11	Monu

M=3

R ₀	R ₁	R ₂	R ₃
2	6	7	10
3	4	8	11
1	5	9	

Back to disk

2	Amit
6	Rahul
4	Magav
5	Neha
3	Pawan
1	Ravi

7	Abhishek
10	Jaya
11	Monu
8	Rajesh
9	Santosh

$$\text{Total Cost} = b_r (\log_{m-1} (b_r/m) + 1)$$

where, m = no. of runs

$$\text{Total Seek Time} = 2 \lceil b_r/m \rceil + \lceil b_r/b_o \rceil$$

$$(\log_{m-1} (b_r/m) - 1)$$

7	Abhishek
2	Amit
10	Jaya
11	Monu
6	Rahul
8	Rajesh
4	Magav
5	Neha
3	Pawan
1	Ravi
9	Santosh

Query Processing & Optimization

SQL $\rightarrow$ Ext. Rel. Algebra $\rightarrow$ [Various Execution Plan] $\downarrow$ Best Execution Plans

Catalog Information

Equivalence Rule

1) Conjunctive selection Opⁿ:- can be deconstructive into a sequence of individual selection.

$$\sigma_{O_1 \wedge O_2}(E) \cong \sigma_{O_1}(\sigma_{O_2}(E))$$

2) Selection Opⁿ are Commutative:-

$$\sigma_{O_1}(\sigma_{O_2}(E)) = \sigma_{O_2}(\sigma_{O_1}(E))$$

3) $\pi_{L_1}(\pi_{L_2}(\pi_{L_3} \dots \pi_{L_n}(E))) = \pi_{L_1}(E)$

4) Cartesian Combination:- $\sigma_{\theta}(E_1 \times E_2) = E_1 \bowtie_{\theta} E_2$

$$\sigma_{O_1}(E_1 \bowtie_{O_2} E_2) = E_1 \bowtie_{O_1 \wedge O_2} E_2$$

5) Theta Join Opⁿ are Commutative:-

$$E_1 \bowtie_{\theta} E_2 = E_2 \bowtie_{\theta} E_1$$

6) Associative (3 Join):-

$$(E_1 \bowtie_{\theta} E_2) \bowtie_{\theta} E_3 = E_1 \bowtie_{\theta} (E_2 \bowtie_{\theta} E_3) \rightarrow \text{Natural Join}$$

$$(E_1 \bowtie_{O_1} E_2) \bowtie_{O_2 \wedge O_3} E_3 = E_1 \bowtie_{O_1 \wedge O_3} (E_2 \bowtie_{O_2} E_3) \rightarrow \text{Theta Join}$$

(Possible iff O_2 involves attr. only from E_2 & E_3)

7) Selection Opⁿ distributes over 3 Join Opⁿ under two conditions:

a) when all Attr. in O_1 involve only the Attr. of one of the Exp(E_i) being Join.

$$\sigma_{O_1}(E_1 \bowtie_{\theta} E_2) = \sigma_{O_1}(E_1) \bowtie_{\theta} E_2$$

b) when O_1 involve only Attr. of E_1 & O_2 of E_2 and if you have this,

$$\sigma_{O_1 \wedge O_2}(E_1 \bowtie_{\theta} E_2) = \sigma_{O_1}(E_1) \bowtie_{\theta} \sigma_{O_2}(E_2)$$

8) Projection Opⁿ distributes over theta Join:-

a) If θ involves only L_1 & L_2 Attr. (θ from L_1, L_2)

$$\pi_{L_1, L_2}(E_1 \bowtie_{\theta} E_2) = \pi_{L_1}(E_1) \bowtie_{\theta} \pi_{L_2}(E_2)$$

b) $E_1 \bowtie_{\theta} E_2$

L_1 = Set of Attr. of E_1

L_2 = " " " " E_2

L_3 = " " " " of E_3 that involves in θ but are not in L_1, L_2 .

$E_1 = (A, B, C, D, E)$, $L_1 = (A, B, C, D)$

$E_2 = (B, C, D, F, Q)$, $L_2 = (B, C, D)$

$L_3 = (E, F)$

L_4 = " " " " E_1 that involves in θ but not in L_1, L_2

$$\pi_{L_1, L_2}(E_1 \bowtie_{\theta} E_2) = \pi_{L_1, L_2}(\pi_{L_1, L_3}(E_1) \bowtie_{\theta} \pi_{L_2, L_4}(E_2))$$

9- Union & Intersection are Commutative:-

$$E_1 \cup E_2 = E_2 \cup E_1$$

$$E_1 \cap E_2 = E_2 \cap E_1$$

10) Union & Intersection are Associative:-

$$(E_1 \cup E_2) \cup E_3 = E_1 \cup (E_2 \cup E_3)$$

$$(E_1 \cap E_2) \cap E_3 = E_1 \cap (E_2 \cap E_3)$$

11) Selection opⁿ distributes over Union, Intersection & Set Difference

$$\sigma \cup \cap \Delta -$$

$$\sigma_{\theta}(E_1 - E_2) = \sigma_{\theta}(E_1) - \sigma_{\theta}(E_2)$$

↑
U & ∩

$$\sigma_{\theta}(E_1 \cap E_2) = \sigma_{\theta}(E_1) \cap \sigma_{\theta}(E_2)$$

↑
X & Δ ∩

12) Project opⁿ distributes over Union:-

$$\pi_{L_1}(E_1 \cup E_2) = \pi_{L_1}(E_1) \cup \pi_{L_1}(E_2)$$

13) Join Ordering:-

$$(R_1 \bowtie R_2) \bowtie R_3 = R_1 \bowtie (R_2 \bowtie R_3)$$

which order to prefer?

If $R_2 \bowtie R_3$ is quite large & $R_1 \bowtie R_2$ is small, then LHS to choose.

Statistical Information for Cost Estimation

Catalog Information ↘ ↘ ↘

I- for each relation

- n_r — number of tuples in 'r'
- b_r — # of blocks containing tuples of 'r'
- l_r — size of a tuple of 'r'
- f_r — blocking factor — No. of tuples of 'r' that fit into 1 block.

$$f_r = \frac{n_r}{b_r}$$

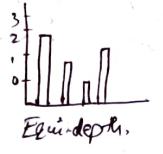
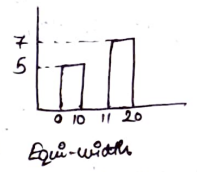
$$b_r = \frac{n_r}{f_r}$$

II. for each Attrn of a relation:-

$V(A, r) = \#$ of distinct values that appear in relation r for A (Attrn).
($\pi_A(r)$)

$\min(A, r)$
 $\max(A, r)$
Histogram

↙ Equi-width
↘ Equi-depth



- ⇒ Equi-width divides the range of values into equal size range
- ⇒ Equi-depth adjust the boundaries of the range i.e. each range has same # of value

III. Index (B⁺ Tree):-

- ↳ Height
- ↳ # of Leaf Tree Nodes
- ↳ Order of Tree

Statistical Information for Cost Estimation

- n_r — no. of tuples in 'r'
- b_r — no. of blocks containing tuples of 'r'
- l_r — size of a tuple of 'r'
- f_r — blocking factor
- $V(A, r)$ — No. of distinct values that appears in 'r' for Attrn. A
- $b_r = \left\lceil \frac{n_r}{f_r} \right\rceil$

- $\min(A, r)$ — min. value of A in 'r'
- $\max(A, r)$ — max. value of A in 'r'

Selection Size Estimation

1) $\sigma_{A=V}(X)$, $\frac{n_k}{V(A,X)}$
 $\Rightarrow$ If A is Key, then result will be 1/0 tuple.
 $\Rightarrow$ Assume uniform distribution of values, $\frac{n_k}{V(A,X)}$, $V(A,X)$ = values of occurrence of A in n

2) $\sigma_{A \leq V}(X)$

a) If A is less than equal to v as 0, if $V < \min(A, X)$

b) If A is $\geq \max(A, X)$ then n_k tuples (All tuples).

c) Otherwise, (b/w 22 to 90),

$$\frac{n_k * (V - \max(A, X))}{\max(A, X) - \min(A, X)}$$

AdmNo	Marks
1	90
2	80
3	40
4	35
5	22



d) $\sigma_{A < V}(X)$, $\sigma_{A \leq V}(X) - \sigma_{A=V}(X)$

e) $\sigma_{A > V}(X)$, $n_k - \sigma_{A \leq V}(X)$

f) $\sigma_{A \neq V}(X)$, $n_k - \sigma_{A=V}(X)$

g) Conjunction ($\sigma_{O_1 \wedge O_2 \wedge \dots \wedge O_n}(X)$)

$\Rightarrow$ Selectivity of Condition O_i is the probability that a tuple in relation n satisfies O_i . = $\frac{S_i}{n_k}$

$$\Rightarrow \left[\frac{S_1}{n_k} \times \frac{S_2}{n_k} \times \dots \times \frac{S_n}{n_k} \right] n_k$$

$$\Rightarrow n_k * \frac{S_1 * S_2 * S_3 * \dots * S_n}{(n_k)^n}$$

Adm	Prog	Gender
1	BT	M
2	BT	F
3	MT	M
4	MT	M
5	BT	F
6	DD	M

$$S[BT] = 3/6 \Rightarrow 1/2$$

$$S[MT] = 2/6 = 1/3$$

$$S[DD] = 1/6$$

$$S[M] = 4/6 = 2/3$$

$$S[F] = 2/6 = 1/3$$

Q. Find $\sigma_{Prog=BT \wedge Gender=M}(X) = ?$

$$\Rightarrow 6 * \frac{3}{6} * \frac{4}{6} \Rightarrow 2$$

$$\sigma_{Prog=MT \wedge Gender=F}(X) = 6 * \frac{2}{6} * \frac{2}{6} \Rightarrow \frac{2}{3}$$

h. Disjunction: $\sigma_{O_1 \vee O_2 \vee O_3 \dots \vee O_n}(X) :-$

$$n_k - n_k \left[\left(1 - \frac{S_1}{n_k}\right) \left(1 - \frac{S_2}{n_k}\right) \dots \left(1 - \frac{S_n}{n_k}\right) \right]$$

$$\text{Negation: } \sigma_{\neg O}(X) = n_k - \sigma_O(X)$$

Size of Joins

$$R \times S = n_k \times n_s$$

1) IF $(RNS) = \emptyset$, then $R \bowtie S$

2) IF (RNS) is a key for 'R'

$$\# \text{ of tuples in } R \bowtie S < n_s$$

3) RNS is a foreign key in S referencing R

$$\# \text{ of tuples in } R \bowtie S = n_s$$

4) $RNS = \{A\}$ is not a key for R or S

$\Rightarrow$ If we assume that every tuple T in R produces tuples in $R \bowtie S$, then # of tuples in $R \bowtie S =$

$$\frac{n_k * n_s}{V(A, S)}$$

$\Rightarrow$ If we assume every tuple T in S produces in $R \bowtie S$, then

$$\# \text{ of tuples} = \frac{n_k * n_s}{V(A, S)}$$

Adm	Name
1	A
2	B
3	C
4	D

RNS = Adm (key of R)

X	S
1	1
2	2
3	3
4	4

result will never be more than n_s as RNS is key of R.

Adm	Sid
1	S1
2	S2
3	S3
4	S4

$$n_k = 4$$

$$V(Sid, n) = 4$$

$$R \bowtie S = 4 * \frac{3}{4} \Rightarrow 4$$

Tid	Sid
T1	S1
T2	S2
T3	S3

$$n_s = 3$$

$$V(Sid, n) = 3$$

Estimation of Size of Projection

$$\pi_A(x) = v(A, x)$$

$$1) \text{ } \sigma_{F}(x) = v(A, x)$$

2) Set operation

$$\sigma_{O_1}(x) \cup \sigma_{O_2}(x) = \sigma_{O_1 \vee O_2}(x) \text{ [Disjunction]}$$

$$x \cup y = [n_x + n_y]$$

$$x \cap y = [\min(n_x, n_y)]$$

$$x - y = [n_x]$$

$$x \bowtie y = [n_x] + \text{Size of } x \bowtie y$$

$$x \bowtie y = \text{Size of } x \bowtie y + [n_y]$$

$$x \bowtie y = \text{Size of } x \bowtie y + [n_x + n_y]$$

Choice of Evaluation Plan

→ Cost based Approach

→ Heuristic

$$x, y, z, \dots, n$$

$$\text{ways} = \frac{[2(n-1)]!}{(n-1)!}$$

→ Perform Selection early (not always true)

$$\text{Ex: } \sigma_{\theta}(x \bowtie y)$$

and s is very large compared to x .

$$\sigma(x \bowtie \sigma_{\theta}(s))$$

$$\Rightarrow x < s$$

→ Index on the join attr. of s exists. No index on the attribute by θ of 's'

→ So join will take place by non key (B, C) & takes time

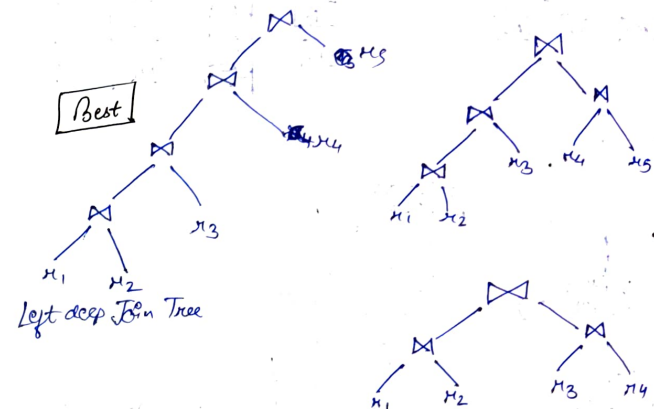
$$s(A, B, C, D)$$

↓
Index

(40)

→ Perform Projection early

→ Perform most restrictive selection & join Op^n (i.e. with smallest result size) before other similar Op^n .

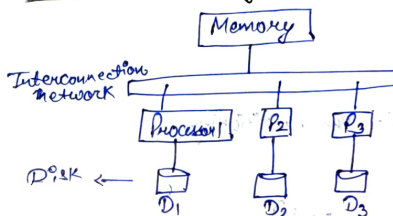


Parallel Databases

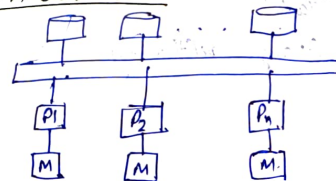
→ Objective = To provide Parallelism

→ Query can be processed from multiple places.

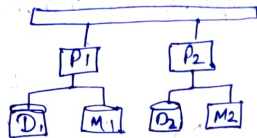
Shared Memory Architecture



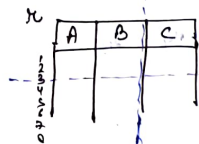
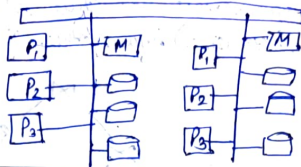
3) Disk Architecture



3) Shared Nothing Architecture



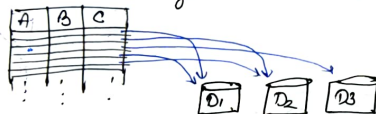
4) Hierarchical Architecture



Horizontal Partitioning

Vertical Partitioning

- Round Robin
- Hash Partitioning
- Range Partitioning



1) Round Robin

Disk $\% \text{ mod } n$
→ one by one into Disks.

2) Hash Partitioning

$h(i)$ → it will return disk no.

3) Range Partitioning

Partition Vector = [5, 11]

→ Blocks 1 to 5 → D_1

→ Blocks 6 to 11 → D_2

→ Blocks 12 to ∞ → D_3

→ Better the Partitioning Vector better the Partitioning.

Types of Query

1) Scanning the entire relation (Round Robin)

2) Point Query $B=5$ (Hash)

3) Range Query $5 < B < 11$ (Range Partitioning)

Advantage of Round Robin

→ Best suited for sequential scan of entire relation.

→ Complex for range queries.

Advantage of Hash Partitioning

→ Sequential scan

→ If Point Query & search key is from same Attr. on which Hash Join has been performed.

→ In worst case of Range Query, scanning the whole disks.

Advantage of Range Partitioning

→ Very good for sequential scan & Range Query.

→ for Point Query also it is easy to search.

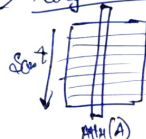
Handling Skew

→ Attribute-value skew

→ Partition skew.

→ Some values appear in the Partitioning Attr. of many tuples. So all tuples will fall into the same Partitioning.

1) Range Partitioning



Histogram

Sort relation based on skewed Attr. Then it will divide m blocks into n partitions (m/n) to avoid skewness

Virtual Processes Partitioning

→ n real processes

→ $m \gg n$ (VP)

2) Hash Partitioning

→ To avoid skewness, make good hash function.

Advantages of Partitioning

1) Different queries can run in parallel on data at different sites.

2) Time required for retrieval of relation (data) from disk is reduced.

3) Increase in throughput of transaction processing.

Parallelism

Interquery Parallelism

Multiple Query running

Intraquery Parallelism

Single Query involved

Intrasoperation Parallelism

Multiple operations (o, n)

Interoperation Parallelism

Parallel Sort

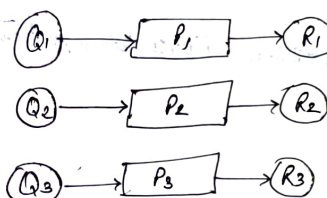
Range Partitioning Sort
Parallel External Merge Sort

Parallel Join

Partition Join
Fragment & Replicate Join
Partition Parallel Hash Join
Parallel Nested Join

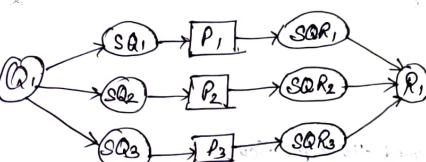
Pipeline Parallelism

Independent Parallelism

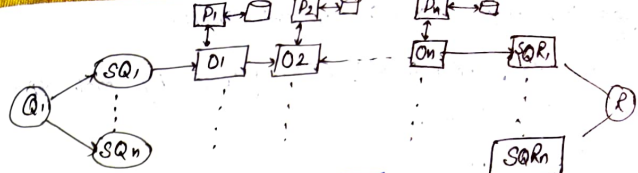


where, Q = Query
P = Processor
R = Result

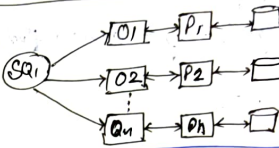
Interquery Parallelism



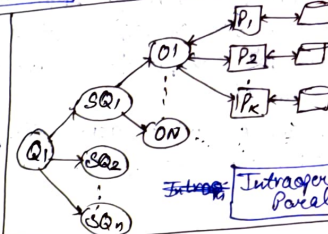
Intraquery Parallelism



Pipeline Parallelism



Independent Parallelism



Intraoperation Parallelism

Interquery Parallelism (Parallelism among Query)

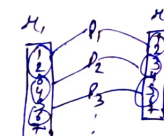
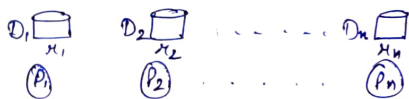
- Easiest form of 11sm to support particularly in shared m/m db.
- More complicated to implement on shared disk or sharing Nothing Arch.
- Some of the examples of cache coherence protocols of shared disk systems are Locking of data (Shared/Exclusive).
- Cache Coherence protocols for shared Nothing System uses Home processor (All independent M/C)

Intraquery Parallelism (Parallelism within Query)

Parallel Sort

- Sort a relation that resides on n disks.
- $R \rightarrow D_1, D_2, \dots, D_n$

Range Partitioning Sort



- Identify Range Partitioning vector
- Redistribute the relation using RP vector.
 - All tuples that lie in the i^{th} range are sent to the processors P_i
 - P_i stores the result temporarily if necessary on D_i .

① ~~Disk~~ Steps require I/O & Communication Overheads
There

We get the sorted data in all Disks followed by each disk

④ Final merge opⁿ is final.

② Parallel External Merge Sort

① Assume relation is already partitioned.

② Each processor will locally sort the data.

③ Sorted runs for each processors are merged to get sorted final data.

A Sorted Partition

⇒ P_i are range partitioned across the processor.

⇒ Each P_i has own sorted part.

Distributed Databases

→ DS - loosely Coupled System

→ OS - NOS/DOS

↓
Networks → Distributed OS

→ System Image

→ Autonomy

→ Fault Tolerance Capability

↳ Low for NOS (Data is present on one node only)

↳ High for DOS (Because of replicas of data)

→ Each process is having its own m/m & disk.

NOS - no. of systems are connected together

DOS - Responsibility of system to hide the identity of multiple m/c.

① Homogeneous DB → Same type of DBMS are present on every node

→ all sites have identical softwares.

→ Each site surrenders parts of its autonomy in terms of right to change Schemas/softwares

→ Query processing is easy.

② Heterogeneous DB :- Different types of DBMS.

→ Different Schemas have different softwares

→ Query Processing is typical so that Transaction management

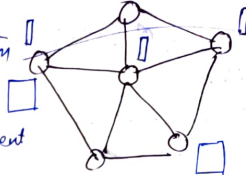
Distributed Data Storage (R)

① Replication:- Multiple copies of Relations are made. and distributed to every nodes. Challenges are write opⁿ to ensure consistency.

② Fragmentation:- Store fragments of relations in different nodes.

③ Combination of Replication & fragmentation:-

Vertical fragmentation



④ Transparency:- User is provided experience of single coherent system. (Everything is present on the user's system).

Fully Redundant DBs

→ Every site contains a copy of entire database.

Advantages of Replication

① Availability ↑

② Parallelism ↑

③ Reduced data Transfer

Disadvantages of Replication

① Write/update (In all places of copies)

Horizontal Fragmentation



account (AN, BN, BAL)

T_{BN} = 'A' (account)
'B'

Advantages

① Parallel Processing on fragments of relations.

Vertical Fragmentation



π_{AN, BN} (account) → fragment 1

π_{AN, BAL} (account) → fragment 2

Common Attr as key to join them again

Advantages:-

① Parallel Processing

Fragmentation Transparency → User is unaware of location of the fragmented data storage.

Replication Transparency = Replication data location

Location Transparency -)

Naming of Data Items

7) Criteria

- 1) Every data item must have system-wide unique name.
- 2) It should be possible to find the location of data items efficiently.
- 3) It should be possible to change the location of Data item transparency.
- 4) Each site should be able to create new data item autonomously.

Centralized Scheme
(Name Server)

- It's a name server responsible for assigning naming.
- Each site maintains a record of local items of data.

Advantages:-

- 1) System wide unique name.
- 2) Efficient to find location of Data Items.
- 3) Possible to change the location of data items.

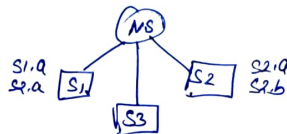
Disadvantages

- 1) Bottleneck
- 2) Single point of failure.
- 3) Each site can't create new data item autonomously.

→ Solution

Use of aliasing like S1.a, S2.a

But locat. Transp. is failed, so we can't use aliasing this way.



58

Solution

Create a set of aliases for data items; Store the mapping of aliases to the real names at each site.

57

Distributed Transactions → Local Trans.
→ Global Trans.

- Atomicity ✓
- Consistency ✓
- Isolation ✓
- Durability ✓

→ Each site has local Transaction Manager & Transaction Coordinator.



→ Responsibilities of TM

- 1) Ensuring ACID Properties to local Transactions. (maintains log file)

→ Responsibilities of TC:-

- 1) Starting the execution of Transaction that originates at the site.
- 2) Distributing sub Transactions at appropriate sites for execution.
- 3) Coordinating the Termination of each Transaction that originates at the site which may result in the Trans. being committed at all sites or aborted at all sites.

System Structure

Commit Protocols

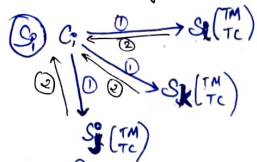
- 2PC (2 Phase CP)
- 3PC (3 Phase CP)



Fail Stop Model

→ When any sort of failure occurs the site immediately stops executing.

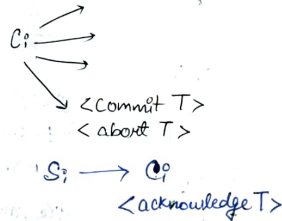
Phase 1 (Obtaining Decision)



1) $\langle \text{Prepare } T \rangle$
 [TM at sites determines if it can commit the Trans.]
 2) $\begin{cases} \text{If No, } \rightarrow \langle \text{no } T \rangle \rightarrow \langle \text{Abort } T \rangle \text{ to } C_i \\ \text{If Yes } \rightarrow \langle \text{ready } T \rangle \rightarrow \langle \text{ready } T \rangle \text{ to } C_i \end{cases}$

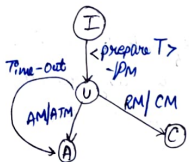
Phase 2 (Recording the Decision)

$\Rightarrow T$ can be committed?
 $\langle \text{ready } T \rangle$ from all sites
 $\Rightarrow T$ should be aborted?
 Any one of site has aborted or any site has not responded in stipulated time.



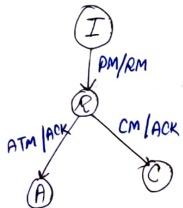
State Diagram

Coordinator



I = Initial state
 U = undecided state
 RM = Ready message
 CM = Commit message
 PM = Prepare T message
 AM = Abort message

Participant



R = Ready state
 ACK = acknowledgement

Handling of failure

1) Site Failure

a) If the Coordinator C_i detects that site has failed
 Actions - 1) Site is at initial state (It means site has a failure) $\Rightarrow$ Abort
 Abort $\Rightarrow$
 2) Site failure at (R) , after ready message has been sent,
 $\hookrightarrow$ failure & IGNORED

On Site S_i Recovery

$\Rightarrow$ Check LOGS
 $\hookrightarrow$ log contains $\langle \text{commit } T \rangle$ record : redo(T)
 $\hookrightarrow$ log contains $\langle \text{abort } T \rangle$ record : undo(T)
 $\hookrightarrow$ log contains $\langle \text{ready } T \rangle$ record : Site must consults C_i to determine whether it's committed or aborted.

UP

$\rightarrow$ Check status, then send commit(redo T) or abort(undo T)

DOWN

$\rightarrow$ If C_i is down, S_i must try to find out the fate of T from other sites.
 $\langle \text{query status } T \rangle \rightarrow$ to all sites of the system.

Yes $\rightarrow$ If any site has the information about T

No $\rightarrow$ In periodic time, it again send message

2) Coordinate Failure

IPC: A new Coordinator to be elected, Reexecute Transaction Protocol.

If some acting participating sites does contain $\langle \text{Ready } T \rangle$ Block in its log, then the failed Coordinator C_i can not have decided to

Commit. $\Rightarrow$ **ABORT**
 $\langle \text{commit } T \rangle \checkmark$
 $\langle \text{abort } T \rangle \checkmark$
 $\langle \text{ready } T \rangle \times$

If none of the above cases hold then all active sites must have ⁽⁶⁰⁾
 $\langle \text{ready } T \rangle$ records in their log but no additional control records
 such as $\langle \text{abort } T \rangle$ or $\langle \text{commit } T \rangle$

Blocking Problem:- when all sites have send $\langle \text{ready } T \rangle$ msg.
 to coordinator, then coordinator fails & all sites have to wait
 to get current state of Transaction (abort/commit) after it recovers.
 And all sites will wait and block the resources they are using.

Scenario 1: One of the live participant received the decision

→ Non-Blocking Problem.

Scenario 2: None of the live participant know the decision of
 Coordinator.

Case 1 → All participants are live (No Failure)

Case 2 → Participants also failed along with Coordinator.
 (Have to wait)

3) Network Partition

- 1) If coordinator & its participants are in one partition then no failure
- 2) If they belong to several partitions, they may be blocking problem
- 3) Sites that are not in the partition where coordinator is present, think that coord. has failed and executes the protocol to deal with failure of coord.
- 4) Coord. & sites are in same partition think that the sites in other partition have failed.

Recovery & Concurrency Control ⁽⁶¹⁾

- In-doubt Transaction is that we don't know its Abort or Committed.
- The Recovering site must determine the Commit/Abort status of such Transactions by contacting other sites. This can slow & potentially block recovery.
- Recovery Algo. can look ~~know~~ note lock info. in the log.
 $\langle \text{ready } T, L \rangle$, where L = List of all write locks held by T
 where the log is written.
- For every in-doubt Transaction T , all the locks noted in the
 $\langle \text{ready } T, L \rangle$ log record are reacquired.

3 Phase Commit Protocol

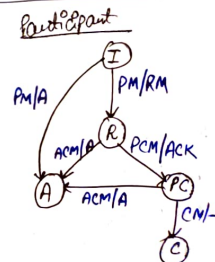
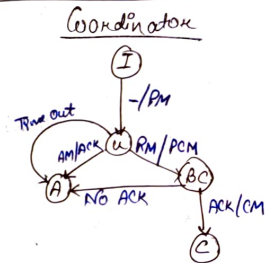
→ To avoid the Blocking Problem of 2PC.

→ Assume, no New Partitioning.

, At any point at least one site is up.

, At most K sites (Participants as well as Coordinator)
 can fail.

where,
 PCM = Prepare to
 Commit msg
 BC = Before Commit
 PM = Prep. msg
 RM = Ready msg
 AM = Abort msg
 ACM = Abort commit
 msg



Coordinate failure

1 PC $\rightarrow$ Same as 2PC

2 PC $\rightarrow$ Also have a Participant failure
Participant will be in PC (Prepare to commit State) &
Coordinator will be in BC (Before Commit State). Then
we can Rest of the participants elect new coordinator.

3 PC $\rightarrow$ Coord. passed to 3PC & failed. (Commit State)
Taken decision to commit & convey to participants.

Concurrency Control

1) Locking Protocol

a $\rightarrow$ Single Lock Manager Approach: One Single is SL manager & other
willing to sites to take lock will ask from manager.

If opⁿ is Read, then data can be
read from any Replica.

If opⁿ is write, then data/change
should be written back to all replicas.



Advantages of SLMA:-

- 1) Implementation is easy.
- 2) If there is deadlock, the manager can easily resolve.

Disadvantages of SLMA:-

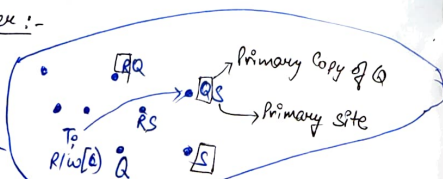
- 1) If manager is down, then system will shut (Single Point of failure)
- 2) There will be bottleneck (Traffic).

b) Distributed Lock Manager:-

1) Primary Copy:-

Primary copy of a single
data item work as single
Lock Manager.

Once a write opⁿ is started then the Primary Copy will
send msg to invalidate all copy of data (Replicas).



Drawback of Primary Copy

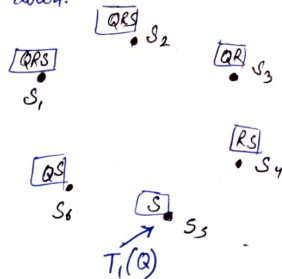
1) If Primary Copy site is down.

2) Majority Protocol

a) Q is Replicated at
 $n=4$ sites, then a
lock request msg
must be sent to
more than half of
the n sites on which
 Q is stored.

$$\text{msg Exchange} = 2\left(\frac{n}{2} + 1\right)$$

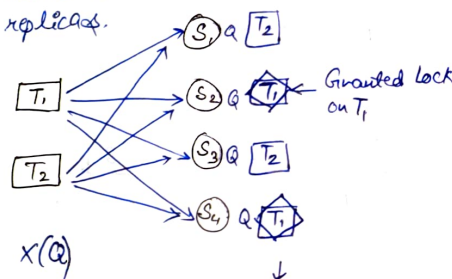
$$\text{lock Released} = \left(\frac{n}{2} + 1\right)$$



b) The Transaction doesn't operate on Q until it has obtained
a lock on majority of Replicas of Q .

c) When writing the data items, Transactions performs
writes on all replicas.

3) Deadlock
occurs.



Deadlock

Both Read & write opⁿ involves some level of complexity.

4) Implementation is complex.

3) Biased Protocol:-

$\rightarrow$ Biased towards Read/shared Locks.

Advantage

1) Read opⁿ can be handled faster.

Disadvantage:- Possibility of Deadlock, Complexity

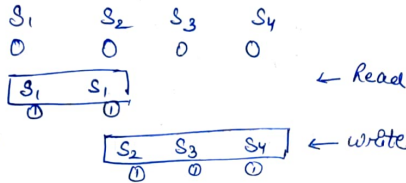
(See Disu of Majority Protocol)

4. Quorum Consensus Protocol

→ The Protocol assigns each site that have replica with a weight
 → for any data item, the Protocol assign read Quorum (Q_r) & write Quorum (Q_w) → integers, (Sum of weights of some sites)
 and these 2 integers are chosen according to the following condition put together:

$$\begin{aligned} Q_r + Q_w &> S \\ 2Q_w &> S \end{aligned}$$

Q to All (S_1 to S_4)
 weight = 1
 $S = 4$

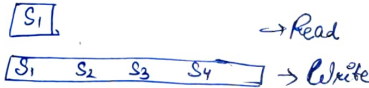


Ex: $Q_r = 2$
 $Q_w = 3$
 $S = 4$

$$\begin{aligned} 2 + 3 &> 4 \\ 2 \times 3 &> 4 \end{aligned}$$

Ex: $Q_r = 1$
 $Q_w = 4$
 $S = 4$

$$\begin{aligned} 1 + 4 &> 4 \\ 2 \times 4 &> 4 \end{aligned}$$



2. Timestamping

T_i issue read(Q)

a) If $\text{Timestamp}(T_i) < W\text{-TS}(Q)$
 then T_i needs to read value of Q that is already written
 ⇒ read(Q) is Rejected
 ⇒ T_i is rollback.

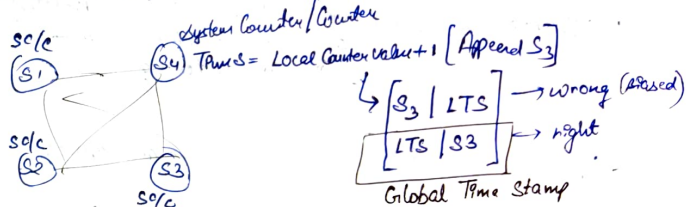
b) If $\text{TS}(T_i) \geq W\text{-TS}(Q)$
 then, read(Q) is executed
 $R\text{-TS}(Q) = \max(R\text{-TS}(Q), \text{TS}(T_i))$

T_i issued write(Q)

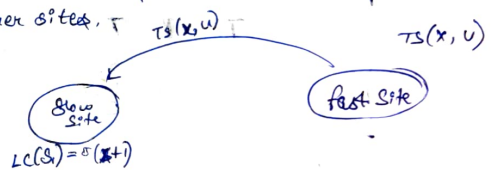
a) If $\text{TS}(T_i) < R\text{-TS}(Q)$, then value of Q that T_i is producing was needed previously & system assume that value would never be produced.
 write(Q) is REJECTED, T_i is Rollback.

b) If $\text{TS}(T_i) < W\text{-TS}(Q)$
 write(Q) is REJECTED, T_i is Rollback

c) Rest all conditions, write(Q) is EXECUTED.
 $W\text{-TS}(Q) = \text{TS}(T_i)$



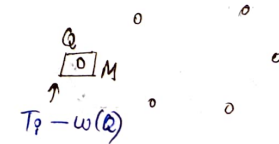
Note:- A site with slow clock will assign smaller Timestamp
 The fast sites logical Counter will be larger than that of other sites.



3. Replication with weak Consistency

↳ DB is updated slowly

Master-Slave Replication



Replica should see transaction consistent Snapshot

Multi Master replication: In this updates are permitted at any replica & are automatically propagated to all the replicas.

Lazy Propagation: Many systems supports Lazy Propagation where updates are transmitted after transaction commits but at a cost of consistency.

It also allows updates to occur even if some sites are disconnected from the net but at the cost of consistency.

Update at Replica [Go to primary site immediately and then propagated lazily to all Replicas. To ensure ordering]

Update at any Replica → Propagate to all replica

4. Deadlock Handling :-

⇒ Easy in case of centralized Systems.

- Centralized
- Hierarchical
- Distributed
 - ↳ wait for graph
 - ↳ probe based

1) Transaction T_1 requested for R_1

2) T_2 requests R_2

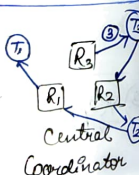
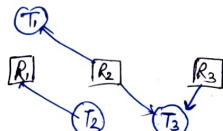
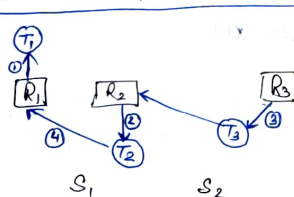
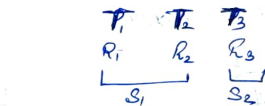
3) T_3 requests R_3

4) $T_2 \leftarrow R_1$ (waiting)

5) $T_3 \leftarrow R_2$ (waiting)

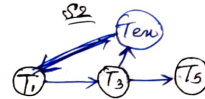
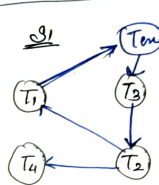
6) T_1 releases R_1 , then R_1 will go to T_2

7) $T_1 \leftarrow R_3$ and is waiting



⇒ Resolve: By appending a unique Global Time Stamp with each msg

Wait For Graph (WFG)



[Such situation is known as false Cycle.]

⇒ An extra node T_m is added to local graph wait for graph of each site & this node is connected to wfg of the corresponding site in the following manner.

a) an edge T_i to T_m is added if a T_i is waiting for Resource in another site being held by any Transaction.

b) An edge T_m to T_j is added if T_j is a Transaction of another site that is waiting for a resource currently being held by a Transaction of this site.

At sites

a) An edge T_1 to T_m is added because a Trans (T_1) is waiting for resource at another site 2. that is held by T_3 . and an edge T_m to T_3 is added because a Transaction T_3 is a process of site 2 that is waiting to acquire a res. currently held by a Transaction T_2 of site 1.

An edge T_3 to T_m because a Transaction T_3 is waiting for resource inside S_1 that is held by T_2 and an edge T_x to T_1 because a Trans. T_1 of Site S_1 that is waiting to acquire a resource currently held by a Transaction T_3 of Site 2.

To confirm a distributed deadlock a DistDDetection Algo. is invoked by a site whose WFG contains a cycle involving T_m .

$$(T_m, T_1, T_2, \dots, T_k, T_m) \rightarrow S_j$$

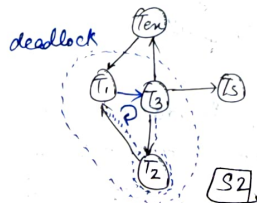
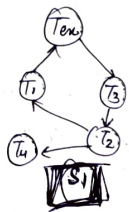
S_j = initiating the Algo (S_j has detected formation of cycle)

- a) S_j will send deadlock detection Message to S_i .
- b) The DDM doesn't contain complete WFG of Site S_i but only the part of WFG that forms the cycle.

Ex:- In above Graphs msg send will be:-

at site 1 $\Rightarrow (T_m, T_3, T_2, T_1, T_m)$ send to S_2

- c) On receiving the msg, Site S_j updates its local WFG by adding those edges of the cycle that don't involve T_m to its WFG (wait for Graph).

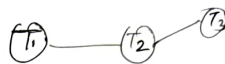


Newly Constructed Graph

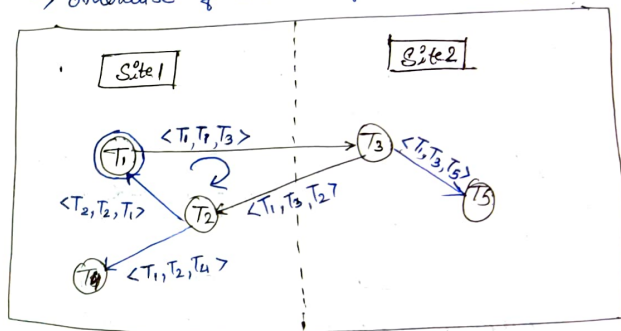
$\Rightarrow$ If a Site S_j contains a cycle that doesn't involve T_m then deadlock exists and appropriate recovery process must be initiated.

Probe Based

Probe message = ~~Sender~~
 $\langle \text{originator, sender, receiver} \rangle$
 $\langle T_1, T_1, T_2 \rangle$



- $\hookrightarrow$ if T_1 grants a resource which it wants then it will proceed
- $\hookrightarrow$ otherwise if will check for deadlock by sending Probe msg.



- 1) T_1 requests for T_3 , then wait for sometime & then send probe msg = $\langle T_1, T_1, T_3 \rangle$
- 2) Then T_3 is also waiting for some resource at T_5 , then it will send $\langle T_1, T_3, T_5 \rangle$
- 3) Then T_5 is doing nothing & will ignore the msg.
- 4) After that T_1 will finally received the msg it has send at the starting that means there is deadlock.

Distributed Query Processing

- In centralized based on Disk access we are choosing Optimized Query.
- But in Distributed we need to calculate Communication Cost to transfer b/w different sites.
- It allows Parallelism, due to replicas.

Q. Find all records of accounts relation.

Ans. Select * from account (in centralized system)

But in Distributed → Replication $\hookrightarrow R$
 → fragmentation $\hookrightarrow R, F$
 → R & F. So on

⇒ So Basically portion of account relation needed to find that where account relation is situated. ~~and~~ we need to construct back account.

- If only Replication is used, then choose any replica of account.
- If replica is fragmented, then we need to use many joins and unions to reconstruct the relation.
- If no replicas are fragmented, then whichever replica is closer to site.
- Relation is Fragmented

Simple Join Strategy

account $\bowtie$ depositor $\bowtie$ branch
 $\downarrow$ $\downarrow$ $\downarrow$
 S_1 S_2 S_3

Query for Join $\rightarrow S_1$

Different ways to handle Joins:-

- Shift Copies of all 3 Relations to S_1 .
- Shift a Copy of Account Relation to S_2

Then at S_2 , find [temp1 = account $\bowtie$ depositor]

After that, Shift temp1 to S_3 ,

Then, temp2 = temp1 $\bowtie$ branch [At S_3]

Shift temp2 to S_1 (Result).

[Shift ✓
Shift X]

Must Considered following factors to choose ① or ②:

- Amount of data being Shipped.
- Cost of transmitting data blocks b/w Sites.
- Relative processing speeds at each sites

Semi-Join Strategy

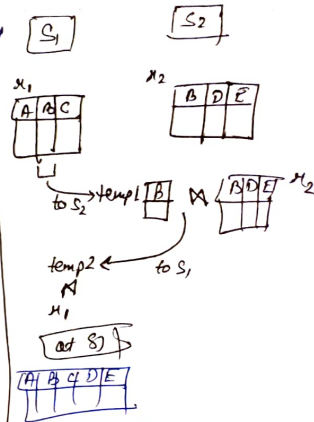
$r_1(R_1) \rightarrow \text{Site 1}$

$r_2(R_2) \rightarrow \text{Site 2}$

Q. $r_1 \bowtie r_2$ and obtain the result at S_1 .

Steps:-

- Compute temp1 $\leftarrow \pi_{R_1, R_2}(x_1)$ at S_1 .
- Ship temp1 from S_1 to S_2 .
- Compute temp2 $\leftarrow r_2 \bowtie \text{temp1}$ at S_2 .
- Ship temp2 from S_2 to S_1 .
- Compute $x_1 \bowtie \text{temp2}$ at S_1 .



3rd Step: $x_2 \bowtie \pi_{R, OR_2}(x_1)$

5th Step: $x_1 \bowtie x_2 \bowtie \pi_{R, OR_2}(x_1)$

$\Rightarrow x_1 \bowtie \pi_{R, OR_2}(x_1) \bowtie x_2$

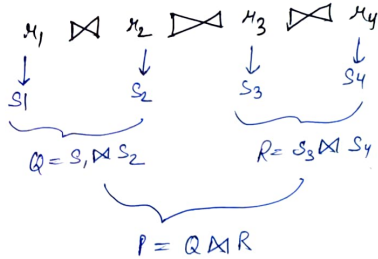
$x_1 \bowtie x_2$

Semijoin at x_1 with $x_2 \Rightarrow x_1 \ltimes x_2$

$\rightarrow$ selects those tuple at x_1 that contributes to $x_1 \bowtie x_2$.

$\pi_{R_1}(x_1 \ltimes x_2)$

Using Parallelism



Distributed Heterogeneous Databases (Multi-Databases System)

Wrapper: A Software which performs schema translation.

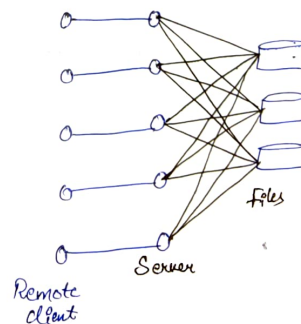
Advanced Transaction Processing

Transaction Processing Manager: Monitor:

- $\rightarrow$ It provides infrastructure for building & administering complex transaction processing system with large no of clients & multiple servers.
- $\rightarrow$ It provides functionalities such as :- Managing, Deploying and Developing transactional Distributed Information Systems.
- $\rightarrow$ TP monitor acts as middleware, its main task is to support and handle interaction b/w applications on a variety of computer platforms.
- $\rightarrow$ It controls programs that monitor or manager, a transaction of data as it passes from one stage in a process to another in an ~~one~~ organised or transaction oriented manner.
- $\rightarrow$ ~~It~~ Transaction Processing is critical in Multi-tier Architecture with processes running on different platforms, a given Transaction may be forwarded to any one of the several servers.
- $\rightarrow$ TP monitor handles all load balancer.

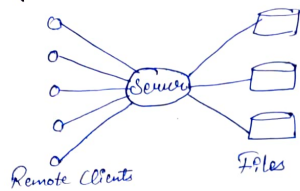
TP Monitor Architecture

a) Process Per-client Model :-



- $\rightarrow$ Instead of individual logging session per terminal, server process communicates with the terminal.
- ~~Problems~~ $\rightarrow$ High CPU overhead for context switching
- $\rightarrow$ Memory leak

b) Single Server Model:-

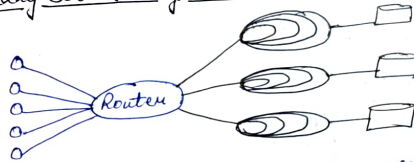


- Multithreading Concept is used
- Server will create multiple threads on port to multiple clients

Problem:-

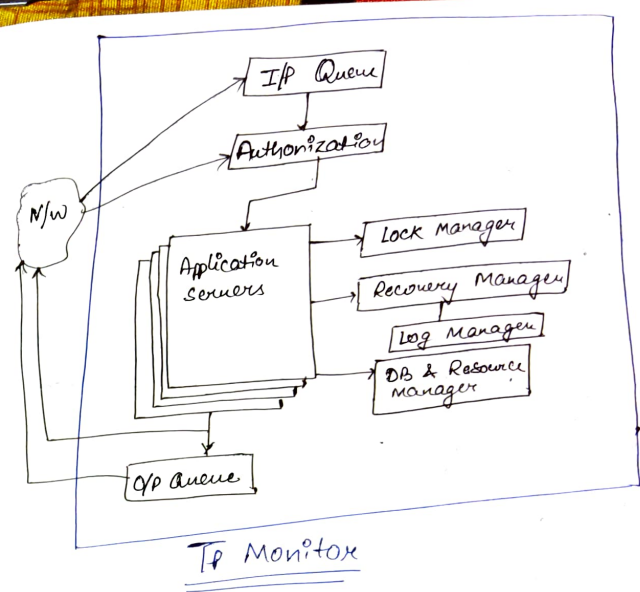
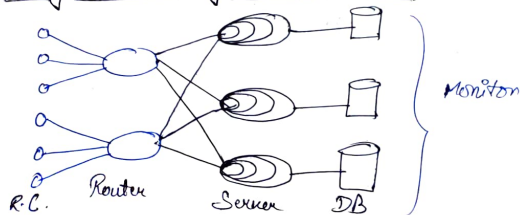
- 1) Because of sharing Address space each thread will share address space. Hence concern of protection among application.
- 2) Moving System to another, then there will be change in address space.

c) Many Servers, Single Server Model



- Multiple applications server process access a common database
- Clients communicates with the application through a single communication process that routes request.
- It supports independent server processor for multiple applications
- Each app. can have a pool of server processors.
- Multithread server processors are allowed.
- Communication process handles communication among process.

d) Many Server, Many Router Model:-



Persistent & Durable Messaging Queue

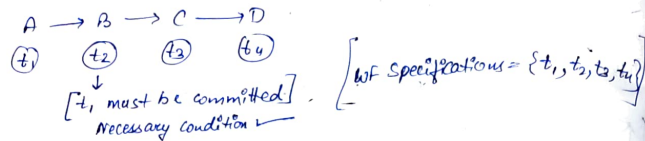
- I/P Queue
- O/P Queue
- Queues are safe even after system will fail

Remote Procedure Call (RPC)

- Widely used in Distributed Systems.

Workflow System

- Workflow Specifications - Predefined things
- WF Execution - Whether in a particular fashion



Cases:-

- 1) Previous activity must have committed.
- 2) Previous activity has committed plus the output must be ~~is acc~~ that is required.
- 3) Based on time any activity is triggered.

→ Workflow Specifications

- ↳ static spec.
- ↳ Dynamic spec.

Static Specifications

- Tasks & Dependencies among them are defined before the execution of workflow starts.

Pre-Execution

- 1) Execution state of other tasks.

$$t_2 \rightarrow t_3$$

- 2) Output value of other tasks

$$t_2 \rightarrow t_3$$

[$x > 25$]

- 3) External variables that are modified by some external parameters

Ex:- Clock

$$t_2 \xrightarrow{\text{clock}} t_3$$

Dynamic Specifications:-

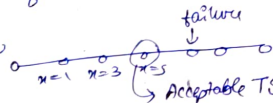
- Not decided by the specifier in prior.
- Ex:- Entering a URL in Browser, Sending mail (Different path from the Router)

Failure Atomicity Requirements

- 1) Acceptable Termination States:- Every execution of workflow will terminate in a state that satisfies the failure Atomicity requirements defined by the design.

→ Committed:- Objectives of a workflow have been achieved

→ Absorted:- valid termination state in which a wf has failed to achieve its objective.



- 2) A workflow must reach an ATState even in the presence of system failure.

- 3) Compensating Transaction:- To undo the result of already committed Transaction.

Execution of Workflow

Workflow management Systems:- It includes

- 1) Scheduler → Everything is planned
- 2) Task Agents
- 3) Query Mechanism

Job of Scheduler:-

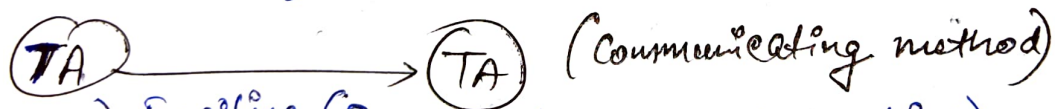
Program that process workflows by submitting various tasks for execution, monitoring various events & evaluating conditions related to intertask dependencies.

Task Agents:- Controls the execution of a task by processing entities.

Query Mechanism:- Different workflow running.

There are Architectural Approaches to the development of Workflow Management Systems.

- 1) Centralized (Single Scheduler)
- 2) Partially Distributed (for each ~~workflow~~ ^{workflow} there is one ^{instantiated} Scheduler)
- 3) Fully Distributed (No Scheduler, but task agents start communicating with ~~itself~~ ^{each other} & coordinate the execution to satisfy task dependencies & other workflow execution requirements.



- 1) Enabling (Dropping of messages very high)
- 2) Persistent Messaging (Guarantee of message delivery)