

Final (2019-20):

1(a): What is database (DBMS)? List and explain four reasons why DBMS is used instead of file processing system.

• DBMS: DBMS is a software system that allows users to define, create, maintain, and control access to a database. It manages data efficiently and allows safe, consistent access by multiple users.

• Four reasons:

- Data redundancy control: In file system the same data might be stored in multiple files. DBMS eliminates duplication by centralizing data, ensuring data is stored only once.

- Data integrity and consistency: DBMS enforces rules to ensure data is accurate and consistent. File systems can not ensure that all files have up-to-date or valid data.

- Data security: DBMS provides user-based access control - different users can have different permissions. File system don't offer this level of control.

- Easy data access and querying: DBMS uses query languages like SQL to retrieve and manipulate data efficiently. In file systems, data retrieval is manual and slower.

1(b): $R = (M, N, P, S, T)$

1) Super Key: multiple ($\underline{M}N, MP, MS, MT, NP, NS, NT, PS, PT, ST, MNP, MPS, MNS, MNT, STP, SPN, MNPS, MNPT, MNPST$)

2) Candidate Key: 2 (MN, NS)

3) primary Key: 1 (can be either 'MN' or NS depends on user).

1(c) How many attributes you can use in a table?
Is there any limitations? Why you need to split the attributes in multiple tables?

The number of attributes in a table depends on the DBMS.

Example limits:

- MySQL: up to 4096 columns per table

- PostgreSQL: up to 1600 columns per table.

- SQL server: up to 2147483647 columns

- Oracle: up to 1000 columns per table.

So, yes there is a limitation, but it depends on the database system and row size.

• split attributes into multiple tables:

1. Remove Redundancy: prevents duplicate data.

2. Ensure data consistency: When data is update in one place, changes reflect everywhere.

3. Improve data Integrity: Foreign key constraints ensure valid and related data.

4. Better query performance: Smaller tables mean faster queries and less memory usage.

5. Scalability and maintainability: Easier to update and maintain smaller tables than one huge table with hundreds of columns.

1 (d) Keyword queries in web search are quite different from database queries. List the difference between the two, in terms of way the queries are specified and in terms of what is the result of a query.

Way queries are specified:

Feature	Web search	Database Queries
Language	Natural language / Keywords	Structured query language
Syntax	Informal, no strict syntax	Formal syntax
User Knowledge	User often doesn't know data structure.	User must know schema
Query Example	best scientific movie 2025	SELECT * FROM MOVIES WHERE Genre = 'Sci-Fi' AND Year = 2025;

Result of a Query:

Feature	Web search	Database Query
Type of result	Ranked list of documents / Web pages.	Exact rows from a structured database
Nature of output	Unstructured or semi-structured (text, link)	Structured data.
Ranking	Results ranked by relevance	No ranking, exact match
Uncertainty	Results may be approximate or fuzzy	Results are deterministic and precise.

Web search is designed for human understanding, uses keywords, and gives relevant documents.

Database queries are for data retrieval, use structured syntax, and return exact data.

2(a) What is database trigger? Discuss its Strength and Weakness of the trigger mechanism.

A trigger is a special kind of stored procedure in a database that automatically executes in response to certain events on a table or view.

Strength of trigger mechanism:

- Automation: Automatically enforces rules or actions without manual input.
- Consistency: Ensures data integrity across related tables.
- Centralized logic: Business rules are stored in one place, not repeated in multiple applications.
- Immediate response: Executes instantly after the specified event.
- Security: Can enforce complex security constraints beyond basic privileges.

Weakness of trigger mechanism:

- Hidden logic: Triggers run in the background; developers may forget they exist, causing confusion.
- Debugging difficulty: Hard to trace and debug, especially when multiple triggers are chained.
- Performance overhead: Extra processing can slow down large INSERT, UPDATE or DELETE operations.
- Complexity: Overuse can make the system hard to maintain and understand.
- Limited portability: Trigger syntax and behaviour can differ across DBMS.

2(b) We can convert any weak entity set to a strong entity set by simply adding appropriate attributes. Then, why do we have weak entity sets.

We can convert any weak entity set to a strong entity set to a strong entity set by simply adding appropriate attributes. (Like a primary key). Weak entity sets are still used in database design for the following important reasons:

- Real world dependency: Weak entities represent objects that depend on another entity for their identification. For example, a dependent of an employee cannot be uniquely identified without the employee.
- Avoiding artificial key: Creating unique attributes may be unrealistic or unnecessary. Weak entities allow us to use a partial key combined with the primary key of the owner entity.
- Semantic clarity: Using weak entity helps to clearly express the existence dependency in the ER model. It shows that the entity cannot exist independently.
- Maintaining referential integrity: Weak entities ensure that no such entity exists without its related strong entity. This simplifies database constraints and helps maintain data consistency.
- Better ER design and normalization: Weak entity helps to design a more normalized and structured database that avoids data redundancy.

2(c) Explain the distinct between database schema and instance.

Database Schema	Database instance
- The overall structure or design of a database -	- The actual data stored in the database at a moment
- It is permanent	- It is Temporary
- It's like a blueprint of a building.	- It's like furniture and people currently in the building
- Defines tables, attributes types, relationships	- The current rows in the table.

2(d)

A file processing system has several disadvantages when used for a large, complex platform like a video site. Here's how these affect both actual video data and metadata.

1. Data Redundancy and Inconsistency:

- Actual video data: If the same video is copied in different folders for different purposes, storage is wasted.

- Metadata: User information or tags might be stored in multiples. If updated in one place but not others, it causes inconsistency.

2. Lack of data sharing and Isolation;

- video data: Multiple user may not be able to access or stream the same video efficiently.

- Metadata: Different applications may find it difficult to access metadata uniformly.

3. Difficulties in Data Access:

- video Data: Accessing specific video parts becomes difficult without structure.

• Metadata: Searching videos by title, tags, or uploader becomes slow and complex using plain files.

4. Integrity Problems:

• Video Data: File corruption might not be detected automatically.

• Metadata: Ensuring values is difficult without constraints.

5. Security issues:

• Video data: Unauthorized users might access, modify or delete videos.

• metadata: private user data can be leaked without proper access control.

3(a) ER-diagram:

• Entities and Attributes:

1. Book (ISBN, title, year, price, category)

2. Customer (customer-id, name, email, phone, address)

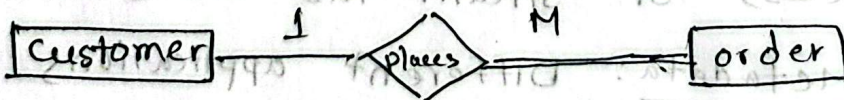
3. Order (order-id, Date, Quantity, total-amount)

4. Payment (payment-id, payment-method, payment-status)

5. Interest (preferred-category)

• Relationships and Cardinality and participation:

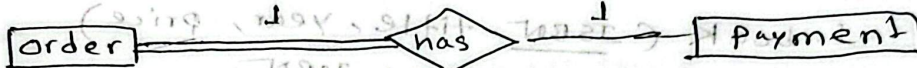
1. customer places order.



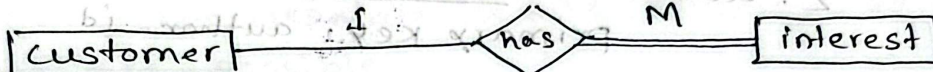
2. Order contains book



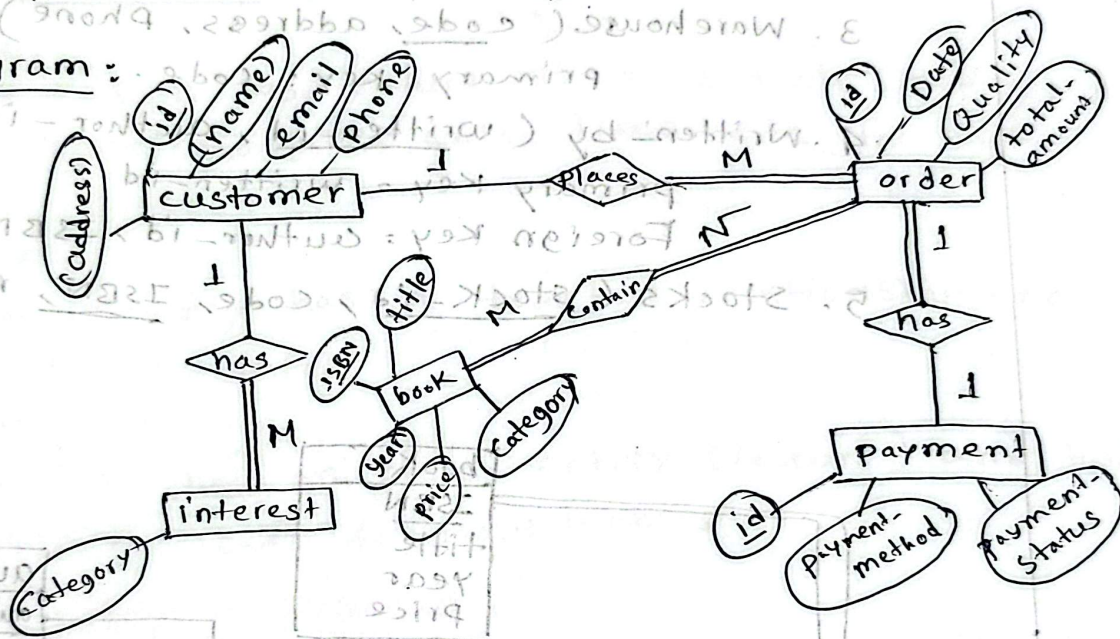
3. order has payment.



4. customer has interest.



ER-diagram:



3(b): Meaning of % and _ in string operators:

In SQL LIKE Operators:

- % (Percentage): Matches any sequence of characters (including empty).
- Example: 'A%' matches 'Any', 'Apple', 'A' etc.
- _ (Undercore): Matches exactly one character.
- Example: 'A_' matches 'An', 'Ax' but not 'Apple'.

3(c) Schema Diagram:

1. book (ISBN, title, year, price)

primary key: ISBN

2. author (author-id, name, address, url)

primary key: author-id

3. Warehouse (code, address, phone)

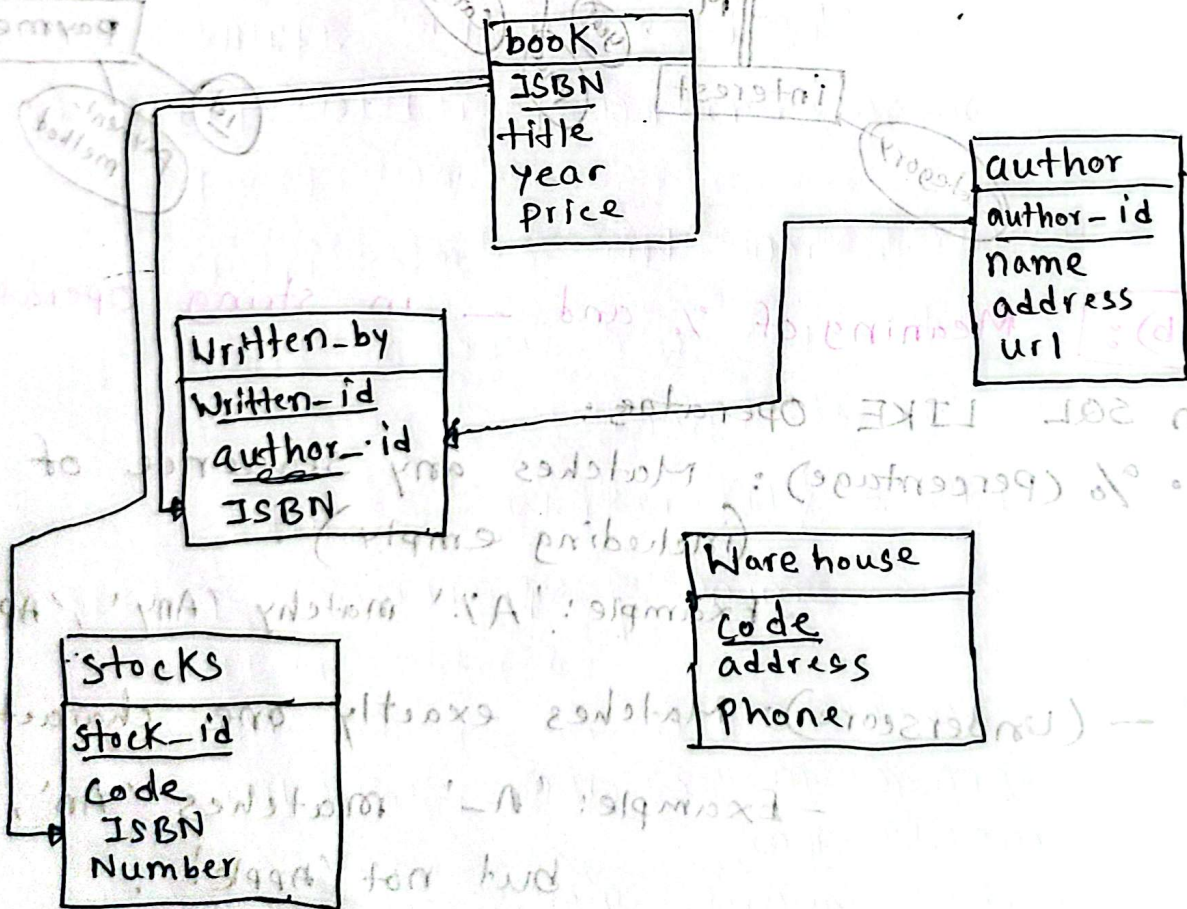
primary key: code

4. Written-by (written-id, author-id, ISBN)

primary key: written-id

Foreign key: author-id, ISBN

5. Stocks (stock-id, code, ISBN, number)



4(a) : Compares ER Model A and B.

- (i) False
 - Model A only shows a lecturer teaches a course → it does not show multiple lectures.
 - Model B allows separate lecture entities, so multiple lectures are possible.

(ii) True

- Since Model B has lecture as a separate entity, we can store info like date, time etc for each lecture.

(iii) True.

- Model A only has a teaches relationship - no specific lecture info.

(iv) True.

- Model B has an extra entity (lecture) and two relationships → leads to more tables.

4(b) Entities and Relationships:

1. Lot (LotNumber, createDate, ControlMaterial, LeftByUser).
2. productionUnits (serial#, exactWeight, productType, QualityTest, productDate).
3. RawMaterials (MaterialID, Name, Type, Unitcost)
4. Includes (many to many between Lot and productionUnits)
5. createdForm (Many to many between productionUnits and Raw materials)

4(c) composite attributes:

In RawMaterials, the attribute 'Name' could be composite if it stores full material description like: 'MaterialName - Brand - code'.

• Explanation: If we break name into Material Name / Brand, code, it becomes a composite attribute.

4(d) Identify Data Redundancy:

In productionUnits, if productType, Quality Test, and productionDate are repeated for each unit even if they belong to the same product batch — that's redundant.

• Reason: These attributes could be moved to a separate product entity, reducing repetition.

5(a)

(i) SELECT Wname
FROM Works

WHERE orgname = 'google';

(ii) SELECT DISTINCT Worker.Wname
FROM Worker

JOIN Works ON Worker.Wname = Works.Wname

JOIN organization ON Works.orgname =
organization.orgname

WHERE Worker.city = organization.city;

(iii) SELECT DISTINCT e.Wname
FROM Worker e

JOIN manages m ON e.Wname = m.Wname

JOIN Worker mgr ON m."manager-name" =
mgr.Wname

WHERE e.city = mgr.city AND e.street = mgr.street;

(iv) UPDATE Works
 SET salary = salary * 1.075
 WHERE Wname IN (SELECT DISTINCT
 "manager-name" FROM manages);

(v) SELECT Wname
 FROM Worker
 WHERE city = (
 SELECT city
 FROM Worker
 WHERE Wname = 'Rahim');

(vi) SELECT orgname
 FROM Works
 GROUP BY orgname
 ORDER BY COUNT(DISTINCT Wname) DESC
 LIMIT 1;

(vii) CREATE VIEW HighEarnings AS
 SELECT Wname, salary
 FROM Works
 WHERE salary > (
 SELECT AVG(salary)
 FROM Works);

(viii) SELECT Wname
 FROM Works
 WHERE orgname = 'Facebook'
 AND jdate <= CURRENT-DATE -INTERVAL
 '5 years';

5(b)

(i). The table has:

FOREIGN KEY (manager-name) REFERENCES
manager (employee-name) ON DELETE CASCADE;

that means:

- If a manager (a row where employee-name = some manager's name) is deleted, then all the employees who had that person as their manager (rows where manager-name = that name) will also be deleted automatically.
- This is called cascading delete. It ensures no employee is left pointing to a non-existent manager.

(ii). If the statement was:

FOREIGN KEY (manager-name) REFERENCES
manager (employee name) ON DELETE RESTRICT;

then:

- If we try to delete a manager's record,
- The deletion will fail if any employee still references that manager in their manager-name.
- No deletion is allowed unless all dependent employees are first removed or ~~see~~ reassigned.

- (6) (a) given table:
- books (accessionno, isbn, title, author, publisher).
 - users (userid, name, deptid, deptname);

given FDS:

accessionno $\rightarrow$ isbn

isbn $\rightarrow$ title

isbn $\rightarrow$ author

isbn $\rightarrow$ publisher.

userid $\rightarrow$ name

userid $\rightarrow$ deptid

userid $\rightarrow$ deptname

already 4NF, 2NF done.

For 3NF:

Create following tables:

Books (accessionno, isbn)

ISBN-info (isbn, title, author, publisher)

Users (userid, name, deptid).

Departments (deptid, deptname):

BCNF

(b) (i) $AB \rightarrow C$

Row 1 and Row 3: $A=1, B=2$

Row 1: $C=3$, Row 3: $C=4 \rightarrow X$ Not equal $\rightarrow$ Doesn't hold.

(ii), $B \rightarrow D$

Row 1 and 3: $B=2 \rightarrow D=4$ (correct)

Row 2: $B=4 \rightarrow D=4$ (correct)

$B \rightarrow D$ holds.

(iii), $DE \rightarrow A$

Row 1 and 2: $D=1, E=5 \rightarrow A=1$ (correct)

$DE \rightarrow A$ holds.

Q1

FDS: $\{AB \rightarrow C, BC \rightarrow D\}$

Step 1:

$AB \rightarrow C$

$AB \rightarrow D$

$BC \rightarrow D$

Step 2:

check $AB \rightarrow C$

- can we remove A from AB?

- $B \rightarrow C$? No (not given)

- can we remove B?

$A \rightarrow C$? NO

→ so $AB \rightarrow C$ is minimal.

check $AB \rightarrow D$

- Already $BC \rightarrow D$ exists. Try removing A.

- Is $B \rightarrow D$ or $BC \rightarrow D$ enough to get $AB \rightarrow D$?

∴ minimal cover: $\{AB \rightarrow C, BC \rightarrow D\}$.

7

Q1 What is database transaction? Discuss

Acid properties of database transaction

⇒ A database transaction is a logical unit of work that contains one or more database operations which are executed as a single unit. It must be either fully completed or fully failed.

ACID properties:

• A (Atomicity): All operations in a transaction are completed; if one fails, the whole

Transaction fails.

- C (consistency): The database remains in a consistent state before and after the transaction.
- I (Isolation): Transactions occur independently without interference.
- D (Durability): Once a transaction is committed, the changes are permanent even in the case of a crash.

7b

T_1 :	LOCK (A)	T_2 :	LOCK (A)
	READ (A)		WRITE (A)
	LOCK (B)		LOCK (B)
	READ (B)		READ (B)
	$B = A + B$		UNLOCK (A)
	WRITE (B)		UNLOCK (B)
	UNLOCK (A)		
	UNLOCK (B)		

Is it deadlock free?

No, this may lead to a deadlock. If T_1 locks A and T_2 locks B at the same time, both will wait for each other to release the lock leading to a circular wait. To prevent deadlock, transactions should lock resources in a consistent order.

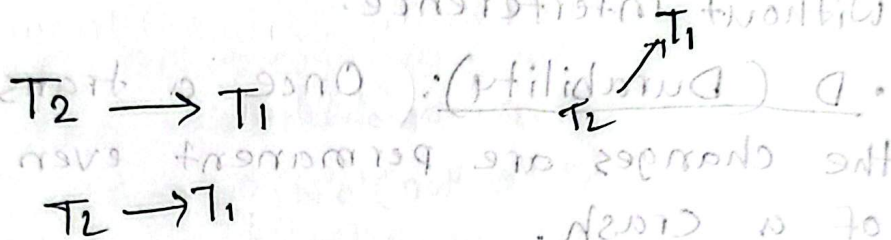
7c

S: $R_1(A); R_2(A); R_3(B); W_1(A); R_2(C); R_2(B); W_2(B); W_1(C)$

check conflicting operations:

• $R_1(A)$ and $W_1(A)$ - same transaction, no conflict.

- $R_2(A)$ and $W_1(A)$ - conflict
- $R_3(B), R_2(B), W_2(B)$ - $R_2(B)$ and $W_2(B)$ is a conflict
- $W_1(C)$ and $R_2(C)$ - conflict



No cycle in graph.



• Dirty read: When one transaction reads data that has been written by another uncommitted transaction.

Example:

T_1 : UPDATE balance SET value = 500

T_2 : SELECT value (reads 500)

T_1 : ROLLBACK

→ T_2 read a dirty value.

• Lost update: When two transactions read the same value and update it, one update is lost.

Example:

T_1 = READ balance = 100 → balance = balance + 50 → write (150)

T_2 = READ balance = 100 → balance = balance + 30 → write (130).

→ final value is 130, T_1 's update is lost.

[8] [a] Differentiate ~~to~~ among Primary, Secondary and clustering Indexing.

- ⇒
- primary index: Created on the primary Key, entries are sorted according to Key order.
 - Secondary index: created on non-primary ~~in~~ attributes, may have duplicates.
 - Clustering index: Based on a clustering field, physically arranges data in file order.

[8] [b] When is it preferable to use a dense index rather than a sparse index?

- ⇒ Dense index is preferable When:
- Records are frequently accessed randomly.
 - Records are small and Φ uniformly distributed.
 - For search heavy applications, as every record has an index.

[8] [c] Is it possible in general to have two primary indices on the same relation for different search keys?

- ⇒ No, ~~in one relat~~ in the same relation there cannot have two primary indices.
- A primary index is based on the primary Key, which must be unique.

— only one attribute can act as the primary key per relation.

[8d]

.course (course-name, room, instructor)

.enrollment (course-name, student-name, grade)

cluster file structure (Example):

course: DBMS, Room: 101, Instructor: Dr. Karim

→ Students: [Tamin, A; Arif, B; Nuyeen, A]

Course: OS, Room: 202, Instructor: Dr. Rahman

→ Students: [Mehedi, A+; Arif, A; Tamin, B]

Explanation:

Here the file is clustered on course-name. The related enrollment records are physically stored with the course information to improve I/O performance on queries like:

SELECT * FROM course JOIN enrollment USING (course_name).