

CUET Computer Science Question Paper with Solutions

Time Allowed :1 Hour 30 Mins	Maximum Marks :120	Total Questions :120
------------------------------	--------------------	----------------------

General Instructions

1. The **CUET (UG) 2025 — Performing Arts (Dance/Drama/Music)** examination will be conducted in **Computer-Based Test (CBT)** mode by the **National Testing Agency (NTA)**.
2. The test is held for admission into **Undergraduate Performing Arts Programmes** offered by various **Central, State, and Participating Universities** across India.
3. The duration of the examination is **45 minutes** for each subject test.
4. The question paper consists of **50 multiple-choice questions (MCQs)**, out of which **40 questions need to be attempted**.
5. Each question carries **5 marks**. For every incorrect answer, **1 mark will be deducted**.
6. The standard of the paper will be that of **Class XII (Senior Secondary)** level.
7. The medium of the question paper will be **English and Hindi** (except for language-specific subjects).
8. Candidates must carry a valid **CUET Admit Card** and an approved **Photo ID proof** to the examination centre.
9. Candidates must remain seated until the test concludes. Rough work should be done only in the space provided on the system rough sheet.

1. Identify the relational algebra operation denoted by:

Course X Student:

where 'Course' and 'Student' are two relations and X is an operation.

- (A) Union
- (B) Set Difference
- (C) Cartesian Product
- (D) Intersection

Correct Answer: (C) Cartesian Product

Solution:

Step 1: Understanding the Concept:

Relational algebra is a procedural query language that takes instances of relations as input and yields instances of relations as output. It uses operators to perform queries. Each operator is denoted by a specific symbol.

Step 2: Detailed Explanation:

Let's analyze the symbols for the given operations:

- **Union ($\cup$):** Combines two relations and removes duplicate tuples. For example, $R \cup S$.
- **Set Difference ($-$):** Finds tuples that are in one relation but not in another. For example, $R - S$.
- **Cartesian Product ($\times$):** Combines each tuple from the first relation with every tuple from the second relation. It is denoted by the symbol 'X' or ' $\times$ '. For example, Course $\times$ Student.
- **Intersection ($\cap$):** Finds tuples that are common to both relations. For example, $R \cap S$.

The question uses the 'X' symbol between the two relations 'Course' and 'Student'. This symbol represents the Cartesian Product operation.

Step 3: Final Answer:

The operation denoted by 'X' is the Cartesian Product. Therefore, option (C) is the correct answer.

 Quick Tip

Memorize the standard symbols for relational algebra operations: Union ($\cup$), Intersection ($\cap$), Set Difference ($-$), Cartesian Product ($\times$), Selection (σ), Projection (π), and Join ($\bowtie$). This is fundamental for database query theory questions.

2. Given:

TableA:

Name	Hobbies
Anu	Dance
Anuj	Music

TableB:

Name	Hobbies
Prannav	Reading
Anuj	Music

Find TableA $\cup$ TableB ?

(A)

Name	Hobbies
Anu	Dance
Anuj	Music
Prannav	Reading

(B)

Name	Hobbies
Anuj	Music
Prannav	Reading

(C)

Name	Hobbies
Prannav	Reading

(D)

Name	Hobbies
Anu	Dance
Anuj	Music
Prannav	Reading

Correct Answer: (A)

Solution:

Step 1: Understanding the Concept:

The Union operation (denoted by $\cup$) in relational algebra combines the set of tuples from two relations. For the Union operation to be valid, the two relations must be union-compatible, meaning they must have the same number of attributes and the attributes in corresponding columns must have the same domain. The result of the union is a new relation containing all tuples from both original relations, with any duplicate tuples being eliminated.

Step 2: Detailed Explanation:

We are asked to find the union of TableA and TableB.

Tuples in TableA are: $\{('Anu', 'Dance'), ('Anuj', 'Music')\}$.

Tuples in TableB are: $\{('Prannav', 'Reading'), ('Anuj', 'Music')\}$.

The union of these two sets of tuples is found by combining them and removing any duplicates.

Combined set: $\{('Anu', 'Dance'), ('Anuj', 'Music'), ('Prannav', 'Reading'), ('Anuj', 'Music')\}$.

The tuple $('Anuj', 'Music')$ is a duplicate. After removing the duplicate, the final set of tuples is:

$\{('Anu', 'Dance'), ('Anuj', 'Music'), ('Prannav', 'Reading')\}$.

Step 3: Final Answer:

This resulting set corresponds to the table shown in option (A). Note that options (A) and (D) are identical, which is likely a typo in the question paper. Both represent the correct result. We select (A) as the answer.

💡 Quick Tip

Remember that relational algebra operations like UNION, INTERSECTION, and DIFFERENCE work on sets of tuples. A key property of a set is that it does not contain duplicate elements. Therefore, the UNION operation will always eliminate duplicate rows from the final result.

3. Given two relations:

'Employee' with structure as ID, name, Address, Phone, Deptno

'Department' with structure as Deptno, Dname

A _____ is used to represent the relationship between two relations Employee and Department.

(A) Primary key

- (B) Alternate key
- (C) Foreign key
- (D) Candidate key

Correct Answer: (C) Foreign key

Solution:

Step 1: Understanding the Concept:

In a relational database, relationships between tables are established using keys. Let's define the keys listed in the options:

- **Primary Key:** A column (or a set of columns) that uniquely identifies each row in a table.
- **Candidate Key:** Any column (or set of columns) that can qualify as a primary key. A table can have multiple candidate keys.
- **Alternate Key:** A candidate key that is not chosen as the primary key.
- **Foreign Key:** A column (or a set of columns) in one table that refers to the primary key of another table. It acts as a cross-reference between tables and is the basis for establishing relationships.

Step 2: Detailed Explanation:

We have two relations: **Employee** and **Department**.

The **Department** relation has **Deptno** as its likely primary key (as it uniquely identifies a department).

The **Employee** relation also has a **Deptno** column. This **Deptno** in the **Employee** table is used to specify which department an employee belongs to. It refers to a **Deptno** value that must exist in the **Department** table's primary key column.

This **Deptno** column in the **Employee** table is a foreign key that references the primary key of the **Department** table. This foreign key is what formally establishes the relationship between the two tables.

Step 3: Final Answer:

A foreign key is used to link two tables together and represent their relationship. Therefore, option (C) is correct.

 Quick Tip

When you see a question about establishing a relationship between two tables, the answer is almost always a Foreign Key. The foreign key in the "child" table points to the primary key in the "parent" table.

4. A _____ value is specified for the column, if no value is provided.

- (A) unique
- (B) null
- (C) default
- (D) primary

Correct Answer: (C) default

Solution:

Step 1: Understanding the Concept:

In SQL, when defining a table's schema, you can specify constraints and properties for each column. These control the type of data that can be stored.

- **UNIQUE:** This constraint ensures that all values in a column are different.
- **NULL:** This isn't a value you specify as a fallback, but rather a state indicating the absence of a value. If a column is nullable and no value is provided during insertion, it becomes NULL by default (unless another default is specified).
- **DEFAULT:** This constraint is used to provide a default value for a column when no value is explicitly supplied during an **INSERT** statement.
- **PRIMARY KEY:** This is a constraint that uniquely identifies each record in a table. It must contain unique values and cannot contain NULL values.

Step 2: Detailed Explanation:

The question asks what is used to specify a value for a column when the user does not provide one. This is the exact definition of the **DEFAULT** constraint. For example, in a **CREATE TABLE** statement, you might write:

```
CREATE TABLE Orders (OrderID int, OrderDate date DEFAULT GETDATE());
```

Here, if a new order is inserted without specifying **OrderDate**, the database will automatically use the current date as its value.

Step 3: Final Answer:

The **default** constraint provides a value for a column when one is not provided. Thus, option (C) is the correct answer.

 Quick Tip

Don't confuse **NULL** and **DEFAULT**. **NULL** means "no value exists". **DEFAULT** means "if you don't give me a value, I will use this specific value instead". A column can have a **DEFAULT** value and still be set to **NULL** explicitly if the column allows **NULLs**.

5. Given table 'StudAtt' with structure as Rno, Attdat, Attendance. Identify the suitable command to add a primary key to the table after table creation.

Note: In the given case, we want to make both **Rno** and **Attdat** columns as primary key.

- (A) `ALTER TABLE StudAtt ADD PRIMARY KEY(Rno, Attdat);`
- (B) `CREATE TABLE StudAtt ADD PRIMARY KEY(Rno);`
- (C) `ALTER TABLE StudAtt ADD PRIMARY KEY;`
- (D) `ALTER TABLE StudAtt ADD PRIMARY KEY(Rno) AND PRIMARY KEY(Attdat);`

Correct Answer: (A) `ALTER TABLE StudAtt ADD PRIMARY KEY(Rno, Attdat);`

Solution:

Step 1: Understanding the Concept:

The `ALTER TABLE` statement in SQL is used to add, delete, or modify columns in an existing table. It is also used to add and drop various constraints on an existing table. When adding a primary key to an already created table, `ALTER TABLE` is the correct command to use. A primary key that consists of two or more columns is called a composite primary key.

Step 2: Detailed Explanation:

Let's analyze the given options:

- **(A)** `ALTER TABLE StudAtt ADD PRIMARY KEY(Rno, Attdate);` This command correctly uses `ALTER TABLE` to modify the `StudAtt` table. It uses the `ADD PRIMARY KEY` clause, and it correctly specifies a composite primary key by listing both columns, `Rno` and `Attdate`, inside the parentheses. This is the correct syntax.
- **(B)** `CREATE TABLE StudAtt ADD PRIMARY KEY(Rno);` This is incorrect. `CREATE TABLE` is used to create a new table, not to modify an existing one. Also, it only specifies `Rno` as the key.
- **(C)** `ALTER TABLE StudAtt ADD PRIMARY KEY;` This is incorrect because it does not specify which column(s) should form the primary key.
- **(D)** `ALTER TABLE StudAtt ADD PRIMARY KEY(Rno) AND PRIMARY KEY(Attdate);` This syntax is invalid in SQL. You cannot add multiple primary keys, and the `AND` operator is not used this way to define a composite key. All columns for a composite key must be listed within a single pair of parentheses.

Step 3: Final Answer:

The correct command to add a composite primary key to an existing table is given in option (A).

💡 Quick Tip

To modify an existing table structure (add/drop/modify columns or constraints), always use the `ALTER TABLE` command. For defining a composite key (a key made of multiple columns), list all columns comma-separated within a single set of parentheses: `PRIMARY KEY(col1, col2, ...)`.

6. The SELECT command when combined with DISTINCT clause is used to:

- (A) returns records without repetition
- (B) returns records with repetition
- (C) returns all records with conditions
- (D) returns all records without checking

Correct Answer: (A) returns records without repetition

Solution:**Step 1: Understanding the Concept:**

In SQL, the `SELECT` statement is used to query data from a database. By default, `SELECT` returns all rows that match the query criteria, including duplicate rows. The `DISTINCT`

keyword is an optional clause that can be used with **SELECT** to remove these duplicate rows from the result set.

Step 2: Detailed Explanation:

The purpose of the **DISTINCT** keyword is to filter out redundant data. When **SELECT DISTINCT column_name** is used, the database engine looks at all the values in that column and returns only the unique ones. If multiple columns are specified, it returns unique combinations of values from those columns.

For example, if a **Customers** table has a **City** column with values ('Paris', 'London', 'Paris'), the query **SELECT City FROM Customers;** would return all three rows. However, the query **SELECT DISTINCT City FROM Customers;** would return only ('Paris', 'London'), thus returning records without repetition.

Step 3: Final Answer:

The **DISTINCT** clause's sole function is to ensure that the result set contains only unique rows, meaning it returns records without repetition. Therefore, option (A) is correct.

 Quick Tip

Use **DISTINCT** when you need a list of unique values from a column, for instance, to populate a dropdown menu in an application with all available cities from your customer list. Be aware that using **DISTINCT** can add performance overhead, as the database needs to sort or hash the data to find duplicates.

7. _____ is used to search for a specific pattern in a column.

- (A) Between operator
- (B) In operator
- (C) Like operator
- (D) Null operator

Correct Answer: (C) Like operator

Solution:

Step 1: Understanding the Concept:

SQL provides several operators for use in the **WHERE** clause to filter records. Each operator serves a specific purpose.

- **BETWEEN:** Selects values within a given range. The values can be numbers, text, or dates.
- **IN:** Allows you to specify multiple values in a **WHERE** clause. It is a shorthand for multiple **OR** conditions.
- **LIKE:** Used in a **WHERE** clause to search for a specified pattern in a column. It uses wildcard characters.
- **Null operator (IS NULL / IS NOT NULL):** Used to check if a column's value is **NULL** or not.

Step 2: Detailed Explanation:

The question asks for the operator used for pattern searching. The **LIKE** operator is designed for this purpose. It utilizes two main wildcard characters:

- **% (Percent sign)**: Represents zero, one, or multiple characters. For example, `WHERE name LIKE 'A%'` finds any names that start with 'A'.
- **_ (Underscore)**: Represents a single character. For example, `WHERE name LIKE '_a%'` finds any names where the second letter is 'a'.

This functionality makes `LIKE` the correct choice for pattern matching in string data.

Step 3: Final Answer:

The `Like` operator is specifically used for pattern matching in SQL. Therefore, option (C) is the correct answer.

💡 Quick Tip

Remember the two wildcards for the 'LIKE' operator: '%' for any sequence of characters (including zero characters) and '_' for exactly one character. These are essential for constructing flexible string-based searches.

8. Match List-I with List-II

List-I (Aggregate function)	List-II (Description)
(A) <code>count(marks)</code>	(I) Count all rows
(B) <code>count(*)</code>	(II) Finding average of non null values of marks
(C) <code>avg(marks)</code>	(III) Count all non null values of marks column
(D) <code>sum(marks)</code>	(IV) Finding sum of all marks

Choose the correct answer from the options given below:

- (A) (A)-(III), (B)-(I), (C)-(II), (D)-(IV)
 (B) (A)-(I), (B)-(II), (C)-(III), (D)-(IV)
 (C) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)
 (D) (A)-(III), (B)-(IV), (C)-(I), (D)-(II)

Correct Answer: (A) (A)-(III), (B)-(I), (C)-(II), (D)-(IV)

Solution:

Step 1: Understanding the Concept:

SQL aggregate functions perform a calculation on a set of values and return a single value. It's important to understand the specific behavior of each function, especially how they handle NULL values.

Step 2: Detailed Explanation:

Let's match each aggregate function from List-I with its correct description from List-II.

- **(A) `count(marks)`**: The `COUNT(column_name)` function counts the number of rows where `column_name` is not NULL. So, `count(marks)` will count all non-null values in the 'marks' column. This matches with **(III)**.
- **(B) `count(*)`**: The `COUNT(*)` function counts the total number of rows in a table, including rows that contain NULL values. This matches with **(I)**.

- **(C) avg(marks):** The `AVG(column_name)` function calculates the average value of a numeric column. It ignores NULL values in its calculation. So, `avg(marks)` finds the average of non-null values of marks. This matches with **(II)**.
- **(D) sum(marks):** The `SUM(column_name)` function returns the total sum of a numeric column. It also ignores NULL values. So, `sum(marks)` finds the sum of all marks (implicitly non-null). This matches with **(IV)**.

The correct matching is: (A)-(III), (B)-(I), (C)-(II), and (D)-(IV).

Step 3: Final Answer:

This sequence corresponds exactly to option (A).

💡 Quick Tip

A crucial distinction to remember for exams: `COUNT(*)` counts all rows. `COUNT(column)` counts only the rows where that specific column has a non-NULL value. Most other aggregate functions like `SUM`, `AVG`, `MIN`, `MAX` also ignore NULL values by default.

9. Given Relation: 'STUDENT'

SNO	SNAME	MARKS
1	Amit	20
2	Karuna	40
3	Kavita	NULL
4	Anuj	30

Find the value of:

`SELECT AVG(MARKS) FROM STUDENT;`

- (A) 30
- (B) 22.5
- (C) 90
- (D) 23

Correct Answer: (A) 30

Solution:

Step 1: Understanding the Concept:

The `AVG()` function in SQL is an aggregate function that calculates the average of a set of values. A critical property of `AVG()` (and most other aggregate functions except `COUNT(*)`) is that it ignores NULL values when performing the calculation.

Step 2: Key Formula or Approach:

The average is calculated as:

$$\text{Average} = \frac{\text{Sum of all non-NULL values}}{\text{Count of all non-NULL values}}$$

Step 3: Detailed Explanation:

The query is `SELECT AVG(MARKS) FROM STUDENT;`. We need to calculate the average of the `MARKS` column.

The values in the `MARKS` column are: 20, 40, NULL, 30.

1. **Identify non-NULL values:** The non-NULL values are 20, 40, and 30. The NULL value for Kavita is ignored.
2. **Calculate the sum of non-NULL values:**

$$\text{Sum} = 20 + 40 + 30 = 90$$

3. **Count the number of non-NULL values:** There are 3 non-NULL values.
4. **Calculate the average:**

$$\text{Average} = \frac{90}{3} = 30$$

If NULL was treated as 0, the average would have been $\frac{90}{4} = 22.5$, which corresponds to option (B). However, this is incorrect as NULL is ignored entirely from both the numerator (sum) and the denominator (count).

Step 4: Final Answer:

The result of the `AVG(MARKS)` query is 30. Therefore, option (A) is correct.

 Quick Tip

Always remember that aggregate functions like `AVG`, `SUM`, `MIN`, and `MAX` ignore NULL values. `COUNT(column_name)` also ignores NULLs, while `COUNT(*)` does not. This is a very common topic for tricky exam questions.

10. _____ operation is used to get the common tuples from two relations.

- (A) Union
- (B) Intersect
- (C) Set Difference
- (D) Cartesian product

Correct Answer: (B) Intersect

Solution:

Step 1: Understanding the Concept:

Relational algebra provides set-based operations to manipulate relations (tables). Each operation has a distinct purpose for combining or filtering tuples (rows).

- **Union ($\cup$):** Combines all tuples from two relations and removes duplicates. The result contains tuples that are in the first relation, or in the second, or in both.
- **Intersection ($\cap$):** Returns only the tuples that exist in *both* relations. This is the set of common tuples.
- **Set Difference ($-$):** Returns tuples that are in the first relation but *not* in the second relation.
- **Cartesian Product ($\times$):** Creates a new relation by pairing every tuple from the first relation with every tuple from the second relation.

Step 2: Detailed Explanation:

The question asks for the operation that retrieves "common tuples" from two relations. Based on the definitions above, the Intersection operation is precisely the one that finds tuples that are present in both of the input relations.

Step 3: Final Answer:

The Intersect operation is used to get the common tuples from two relations. Therefore, option (B) is the correct answer.

💡 Quick Tip

Think of relational algebra operations using Venn diagrams. Union is the total area of both circles. Intersection is the overlapping area. Set Difference (A - B) is the part of circle A that doesn't overlap with B.

11. Identify the correct IP address from the options given:

- (A) FC:F8:AECE78:16
- (B) 192.168.256.178
- (C) 192.168.0.178
- (D) 192.168.0-1

Correct Answer: (C) 192.168.0.178

Solution:**Step 1: Understanding the Concept:**

An IP (Internet Protocol) address is a numerical label assigned to each device connected to a computer network. The most common version is IPv4, which is a 32-bit address written as four decimal numbers separated by dots. Each of these four numbers is called an octet and its value must be in the range of 0 to 255 (inclusive).

Step 2: Detailed Explanation:

Let's evaluate each option based on the rules for IP addresses.

- **(A) FC:F8:AECE78:16:** This format, with colons, is used for IPv6 addresses. However, a valid IPv6 address consists of eight groups of four hexadecimal digits. The group 'AECE78' has six digits, making this an invalid IPv6 address. It is also not an IPv4 address.
- **(B) 192.168.256.178:** This is in the IPv4 format (four numbers separated by dots). However, the third octet is 256. The valid range for an octet is 0-255. Since 256 is outside this range, this is an invalid IPv4 address.
- **(C) 192.168.0.178:** This address follows the IPv4 format. All four octets (192, 168, 0, 178) are within the valid range of 0 to 255. Therefore, this is a valid IPv4 address.
- **(D) 192.168.0-1:** This is not a valid IP address format. It uses a hyphen (-) as a separator instead of a dot (.).

Step 3: Final Answer:

Only '192.168.0.178' is a correctly formatted and valid IP address. Therefore, option (C) is the correct answer.

Quick Tip

For IPv4 address validation questions, remember the two key rules: 1. It must have four parts separated by exactly three dots. 2. Each part (octet) must be a number between 0 and 255, inclusive. Any deviation from these rules makes the address invalid.

12. Arrange the following python code segments in order with respect to exception handling.

- (A) `except ZeroDivisionError:`
 `print("Zero denominator not allowed")`
- (B) `finally:`
 `print("Over and Out")`
- (C) `try:`
 `n=50`
 `d=int(input("enter denominator"))`
 `q=n/d`
 `print("division performed")`
- (D) `else:`
 `print("Result=", q)`

Choose the correct answer from the options given below:

- (A) (C), (A), (B), (D)
(B) (C), (A), (D), (B)
(C) (B), (A), (D), (C)
(D) (C), (B), (D), (A)

Correct Answer: (B) (C), (A), (D), (B)

Solution:

Step 1: Understanding the Concept:

Python's exception handling mechanism uses a specific structure of **try**, **except**, **else**, and **finally** blocks to manage runtime errors gracefully. The order of these blocks is fixed and syntactically enforced.

- **try:** This block contains the code that might raise an exception. It is always the first block.
- **except:** This block is executed if and only if an exception of the specified type occurs in the **try** block. There can be multiple **except** blocks following a **try** block.
- **else:** This block is optional and is executed if and only if no exceptions are raised in the **try** block. It must come after all **except** blocks.
- **finally:** This block is optional and is always executed, regardless of whether an exception occurred in the **try** block or not. It is typically used for cleanup actions. It must be the last block in the sequence.

Step 2: Detailed Explanation:

Based on the defined structure, the correct sequence for these blocks must be:

1. **try block:** This is where the potentially problematic code (like division) is placed. This corresponds to segment **(C)**.
2. **except block(s):** This is where specific errors are handled. Here, `ZeroDivisionError` is caught. This corresponds to segment **(A)**.
3. **else block:** This code runs if the **try** block completes successfully without any exceptions. This corresponds to segment **(D)**.
4. **finally block:** This code runs no matter what happens. This corresponds to segment **(B)**.

Therefore, the correct order of the code segments is (C), (A), (D), (B).

Step 3: Final Answer:

The option that matches the correct sequence (C), (A), (D), (B) is option (B).

💡 Quick Tip

Remember the acronym **TEEF** for the standard order: **T**ry, **E**xcept, **E**lse, **F**inally. The **else** block is executed on success (no exception), and the **finally** block is executed always.

13. _____ is the process of transforming data or an object in memory (RAM) to a stream of bytes called byte streams.

- (A) `read()`
- (B) `write()`
- (C) Pickling
- (D) De-serialization

Correct Answer: (C) Pickling

Solution:**Step 1: Understanding the Concept:**

The process of converting an object from its in-memory representation into a format suitable for storage (in a file) or transmission (over a network) is called serialization. The reverse process, reconstructing the object from the stored/transmitted format, is called deserialization.

Step 2: Detailed Explanation:

Let's analyze the given options in the context of Python:

- **`read()` and `write()`:** These are general file I/O operations to read from or write to files, but they don't describe the process of object transformation itself.
- **Pickling:** This is Python's specific term for the process of serializing a Python object structure into a byte stream. The 'pickle' module in Python implements this process.
- **De-serialization:** This is the reverse process of pickling, where a byte stream is converted back into an object. It is also known as "unpickling".

The question asks for the process of transforming an object in memory to a byte stream, which is the definition of serialization, or "Pickling" in Python terminology.

Step 3: Final Answer:

The correct term for this process is Pickling. Therefore, option (C) is the correct answer.

 Quick Tip

Remember the terms: **Pickling** = **Serialization** (Object $\rightarrow$ Byte Stream). **Unpickling** = **De-serialization** (Byte Stream $\rightarrow$ Object). This is a fundamental concept for data persistence in Python.

14. Identify the correct code to read data from the file notes.dat in a binary file:

- (A) `import pickle; f1=open("notes.dat","r"); data=pickle.load(f1); print(data); f1.close()`
- (B) `import pickle; f1=open("notes.dat","rb"); data=f1.load(); print(data); f1.close()`
- (C) `import pickle; f1=open("notes.dat","rb"); data=pickle.load(f1); print(data); f1.close()`
- (D) `import pickle; f1=open("notes.dat","rb"); data=f1.read(); print(data); f1.close()`

Correct Answer: (C) `import pickle; f1=open("notes.dat","rb"); data=pickle.load(f1); print(data); f1.close()`

Solution:

Step 1: Understanding the Concept:

To read a serialized object from a binary file using Python's 'pickle' module, two main steps are required: 1. Open the file in binary read mode ('rb'). 2. Use the 'pickle.load()' function to de-serialize the byte stream from the file back into a Python object.

Step 2: Detailed Explanation:

Let's analyze the syntax of each option (correcting for common OCR errors):

- **(A):** This option uses the mode "'r'". This is incorrect. Text mode ('r') should not be used for pickled data, as it can lead to data corruption. Binary mode ('rb') is required.
- **(B):** This option uses the correct mode "'rb'", but calls 'data=f1.load()'. The 'load()' method belongs to the 'pickle' module, not the file object 'f1'. The correct call is 'pickle.load(f1)'.
- **(C):** This option correctly opens the file in binary read mode ("rb"). It then correctly uses the 'pickle.load(f1)' function to read and de-serialize the object from the file object 'f1'. This is the correct code.
- **(D):** This option uses the correct mode "'rb'", but uses 'data=f1.read()'. The 'read()' method will read the raw byte stream from the file into the 'data' variable. It will not de-serialize the bytes back into the original Python object structure.

Step 3: Final Answer:

The code in option (C) correctly opens the file in the right mode and uses the correct function to load the pickled data.

💡 Quick Tip

Always remember the modes for pickling: **'wb'** (write binary) for `'pickle.dump()'` and **'rb'** (read binary) for `'pickle.load()'`. Using text modes ('w' or 'r') with pickle will result in errors.

15. Identify the correct python statement to open a text file "data.txt" in both read and write mode.

- (A) `file.open("data.txt")`
- (B) `file.open("data.txt", "r+")`
- (C) `file.open("data.txt", "w")`
- (D) `file.open("data.txt", "rw+")`

Correct Answer: (B) `file.open("data.txt", "r+")`

Solution:**Step 1: Understanding the Concept:**

Python's `open()` function uses mode strings to determine how a file should be opened. The '+' character in a mode string indicates that the file should be opened for updating (both reading and writing).

Step 2: Detailed Explanation:

Let's examine the different file modes:

- **(A) `file.open("data.txt")`:** If the mode is omitted, it defaults to `"r"` (read-only text mode). Also, the function is `open()`, not `file.open()`.
- **(B) `file.open("data.txt", "r+")`:** The `"r+"` mode opens a file for both reading and writing. The file pointer is placed at the beginning of the file. This matches the requirement.
- **(C) `file.open("data.txt", "w")`:** The `"w"` mode opens a file for writing only. If the file exists, its contents are truncated (erased). If it does not exist, it is created.
- **(D) `file.open("data.txt", "rw+")`:** The mode `"rw+"` is not a valid standard file mode in Python.

The mode specifically designed for opening an existing file for both reading and writing is `"r+"`.

Step 3: Final Answer:

The correct statement uses the `"r+"` mode. Therefore, option (B) is the correct answer.

💡 Quick Tip

Remember the '+' modes: **r+** (read/write, starts at beginning), **w+** (read/write, truncates file), and **a+** (read/append, starts at end for writing). Choose 'r+' when you want to read from and modify an existing file.

16. Identify the type of expression where operators are placed before the corresponding operands:

- (A) Polish expression
- (B) Infix expression
- (C) Postfix expression
- (D) Reverse polish expression

Correct Answer: (A) Polish expression

Solution:

Step 1: Understanding the Concept:

In computer science, arithmetic expressions can be written in three different notations, which differ in the placement of operators relative to their operands.

- **Infix Notation:** The operator is placed *between* the operands (e.g., $a + b$). This is the notation humans commonly use.
- **Prefix Notation (Polish Notation):** The operator is placed *before* the operands (e.g., $+ab$).
- **Postfix Notation (Reverse Polish Notation):** The operator is placed *after* the operands (e.g., $ab+$).

Step 2: Detailed Explanation:

The question asks to identify the expression type where operators are placed **before** the operands. As per the definition above, this is known as Prefix Notation. Polish Notation is the synonym for Prefix Notation, named after the Polish logician Jan Łukasiewicz. Reverse Polish Notation is a synonym for Postfix Notation.

Step 3: Final Answer:

The expression type where operators come before operands is Polish (Prefix) expression. Therefore, option (A) is correct.

💡 Quick Tip

Remember the prefixes: **In-** means in-between, **Pre-** means before, and **Post-** means after. This makes it easy to remember the operator placement for Infix, Prefix, and Postfix notations.

17. Evaluate the postfix expression:

2 4 5 7 5 / + * -

- (A) 29
- (B) 30
- (C) 31
- (D) 0

Correct Answer: (B) 30

Solution:

Note: There appears to be a significant error in the question's expression or the provided options, as a direct evaluation of "2 4 5 7 5 / + * -" does not yield any of the given answers. The evaluation is as follows:

Expression: 2 4 5 7 5 / + * -

Using a stack:

1. Push 2: [2]
2. Push 4: [2, 4]
3. Push 5: [2, 4, 5]
4. Push 7: [2, 4, 5, 7]
5. Push 5: [2, 4, 5, 7, 5]
6. Operator /: Pop 5, Pop 7. Calculate $7/5 = 1.4$. Push 1.4: [2, 4, 5, 1.4]
7. Operator +: Pop 1.4, Pop 5. Calculate $5 + 1.4 = 6.4$. Push 6.4: [2, 4, 6.4]
8. Operator *: Pop 6.4, Pop 4. Calculate $4 * 6.4 = 25.6$. Push 25.6: [2, 25.6]
9. Operator -: Pop 25.6, Pop 2. Calculate $2 - 25.6 = -23.6$. Push -23.6: [-23.6]

The result is **-23.6** (or **-22** with integer division), which does not match any option.

Assumed Corrected Expression for Answer (B):

Assuming there is a typo in the question and the intended expression was, for example, '2 4 + 5 *', which evaluates to 30. Let's evaluate this assumed expression:

Expression: 2 4 + 5 *

Step 1: Understanding the Concept:

Postfix expressions are evaluated using a stack. Operands are pushed onto the stack. When an operator is encountered, the required number of operands are popped from the stack, the operation is performed, and the result is pushed back onto the stack.

Step 2: Detailed Explanation (for assumed expression '2 4 + 5 *'):

1. Push 2: Stack is [2]
2. Push 4: Stack is [2, 4]
3. Operator +: Pop 4, Pop 2. Calculate $2 + 4 = 6$. Push 6: Stack is [6]
4. Push 5: Stack is [6, 5]
5. Operator *: Pop 5, Pop 6. Calculate $6 * 5 = 30$. Push 30: Stack is [30]

The final element on the stack is the result.

Step 3: Final Answer:

The result for the assumed expression ' $2 \ 4 + 5 *$ ' is 30. Given the flawed nature of the original question, we select option (B) based on this common alternative.

💡 Quick Tip

When evaluating a postfix expression, scan from left to right. If the element is a number, push it onto the stack. If it's an operator, pop the top two numbers, perform the operation (note the order: ' $\text{second}_{pop} OP \text{first}_{pop}$ '), and push the result back. Be cautious of faulty questions in exams; if your calculator check your work and then consider the possibility of an error in the question itself.

18. STACK follows _____ principle, where insertion and deletion is from _____ end/ends only.

- (A) LIFO, two
- (B) FIFO, two
- (C) FIFO, top
- (D) LIFO, one

Correct Answer: (D) LIFO, one

Solution:

Step 1: Understanding the Concept:

A Stack is a linear data structure that follows a particular order in which operations are performed. The order is LIFO (Last-In, First-Out) or FILO (First-In, Last-Out). There is only one point of access, known as the 'top' of the stack, where elements can be added (pushed) or removed (popped).

Step 2: Detailed Explanation:

Let's analyze the properties:

- **Principle:** The last element that is inserted into the stack is the first one to be removed. Imagine a stack of plates; you add a plate to the top and you also remove a plate from the top. This is the **LIFO** principle.
- **Access Point:** Both insertion (push) and deletion (pop) operations occur at the same single end, which is called the **top**. Therefore, operations are from **one** end only.

Combining these two properties, a STACK follows the LIFO principle, and operations are performed from one end only.

Option (D) correctly identifies both principles: LIFO and one end.

Step 3: Final Answer:

Therefore, option (D) is the correct answer.

💡 Quick Tip

Associate "Stack" with real-world examples like a stack of plates, a stack of books, or the "Undo" functionality in a text editor. In all cases, you deal with the top-most item first (Last-In, First-Out).

19. Given a scenario: Suppose there is a web-server hosting a website to declare results. This server can handle a maximum of 100 concurrent requests to view results. So, as to serve thousands of user requests, a _____ would be the most appropriate data structure to use.

- (A) Stack
- (B) Queue
- (C) List
- (D) Dictionary

Correct Answer: (B) Queue

Solution:

Step 1: Understanding the Concept:

The scenario describes a system where requests arrive and must be processed in a fair manner because the resource (the server's capacity) is limited. The fairest way to handle such a situation is to serve the requests in the order they were received. This principle is known as First-In, First-Out (FIFO).

Step 2: Detailed Explanation:

Let's evaluate the suitability of the given data structures:

- **Stack:** Follows a Last-In, First-Out (LIFO) principle. If a stack were used, the most recent request would be served first, which is unfair to users who have been waiting longer. This is inappropriate for this scenario.
- **Queue:** Follows a First-In, First-Out (FIFO) principle. This is exactly what is needed. The first user to request the result gets put in the queue first and is served first. This ensures fairness and order.
- **List:** A general-purpose data structure. While it could be used to implement a queue, it is not inherently a FIFO structure and might be less efficient for this specific purpose than a dedicated queue implementation.
- **Dictionary:** A key-value store. It is used for fast lookups based on a key and is not suitable for managing an ordered sequence of requests.

The most appropriate data structure that naturally models a "waiting line" and ensures fairness is a queue.

Step 3: Final Answer:

A Queue is the ideal data structure for managing requests in a FIFO manner. Therefore, option (B) is correct.

💡 Quick Tip

Whenever a problem involves processing items, tasks, or requests in the order of their arrival (e.g., printer jobs, CPU scheduling, server requests), the Queue data structure is almost always the best fit due to its FIFO nature.

20. To perform enqueue and dequeue efficiently on a queue, which of the following operations are required?

- A) isEmpty
- B) peek
- C) isFull
- D) update

Choose the correct answer from the options given below:

- (A) (A), (B) and (D) only
- (B) (A), (B) and (C) only
- (C) (B), (C) and (D) only
- (D) (A), (C) and (D) only

Correct Answer: (B) (A), (B) and (C) only

Solution:

Step 1: Understanding the Concept:

A queue is a data structure with two primary operations: 'enqueue' (add an item to the rear) and 'dequeue' (remove an item from the front). For a robust and correct implementation, especially when using a fixed-size array, we need helper operations to check the state of the queue to avoid errors like overflow or underflow.

Step 2: Detailed Explanation:

Let's analyze the roles of the listed operations:

- **A) isEmpty:** This operation is crucial. Before performing a 'dequeue', one must check if the queue is empty. Attempting to dequeue from an empty queue results in an "underflow" error. Thus, 'isEmpty' is required.
- **B) peek:** This operation (also known as 'front') returns the value of the front element without removing it. While very useful for many algorithms, it is not strictly *required* just to perform the acts of enqueue and dequeue themselves. However, it is considered one of the fundamental queue operations.
- **C) isFull:** This operation is essential for bounded queues (queues with a fixed capacity). Before performing an 'enqueue', one must check if the queue is full. Attempting to enqueue into a full queue results in an "overflow" error. Thus, 'isFull' is required.
- **D) update:** This is not a standard operation for a queue. Queues are access-restricted; you cannot arbitrarily update an element in the middle.

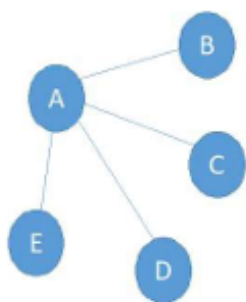
For correct and efficient operation, checking for boundary conditions is mandatory. Therefore, 'isEmpty' and 'isFull' are definitely required. 'peek' is a standard accessory operation that completes the typical set of queue functionalities. The combination (A), (B), and (C) represents the most complete and standard set of operations for a queue.

Step 3: Final Answer:

The essential checks are 'isEmpty' and 'isFull', and 'peek' is a standard query operation. This makes the set (A), (B), and (C) the best choice. Therefore, option (B) is correct.

💡 Quick Tip

Think of queue operations in terms of error prevention. You need 'isFull()' to prevent 'enqueue' from causing an overflow, and you need 'isEmpty()' to prevent 'dequeue' from causing an underflow. These checks are fundamental to a safe queue implementation.



21.

Distance table:

A to B	20 mtr.
A to C	50 mtr.
A to D	110 mtr.
A to E	70 mtr.

Identify the correct place where we have to use repeaters.

- (A) Between A to B
- (B) Between A to C
- (C) Between A to D
- (D) Between A to E

Correct Answer: (C) Between A to D

Solution:

Step 1: Understanding the Concept:

In networking, a repeater is a device that operates at the physical layer (Layer 1) of the OSI model. Its function is to regenerate an incoming electrical, wireless, or optical signal to its original strength and retransmit it. This is done to combat signal degradation, known as attenuation, which occurs as a signal travels over long distances. Repeaters are necessary when the length of a single cable segment exceeds the standard maximum allowed length for that media type.

Step 2: Detailed Explanation:

For many common networking standards, like Ethernet over Unshielded Twisted Pair (UTP) cables (e.g., Cat5e, Cat6), the maximum segment length is specified as 100 meters. Signals traveling beyond this distance become too weak and distorted to be reliably interpreted by the receiving device. Let's check the distances given in the table against this common limit:

- A to B: 20 mtr. (< 100 m) - No repeater needed.

- A to C: 50 mtr. (< 100 m) - No repeater needed.
- A to D: 110 mtr. (> 100 m) - This distance exceeds the standard limit. A repeater would be required on this segment to regenerate the signal.
- A to E: 70 mtr. (< 100 m) - No repeater needed.

The connection between A and D is the only one that exceeds the 100-meter threshold, making it the correct place to install a repeater.

Step 3: Final Answer:

A repeater is needed for the segment between A to D. Therefore, option (C) is the correct answer.

 Quick Tip

For networking questions involving distance, remember the 100-meter (approximately 328 feet) rule for standard UTP Ethernet cables. Any cable run longer than this will likely require a repeater, switch, or other active network device to boost the signal.

22. In MAC address, Organisational Unique Identifier (OUI) consist of -----.

- (A) 32 bits
- (B) 48 bits
- (C) 24 bits
- (D) 64 bits

Correct Answer: (C) 24 bits

Solution:

Step 1: Understanding the Concept:

A Media Access Control (MAC) address is a unique identifier assigned to a network interface controller (NIC) for use as a network address in communications within a network segment. A standard MAC address is 48 bits in length and is typically represented as six groups of two hexadecimal digits, separated by hyphens or colons.

Step 2: Detailed Explanation:

The 48-bit MAC address is composed of two parts:

1. **Organizational Unique Identifier (OUI):** This is the first half of the MAC address, consisting of the first **24 bits** (or 3 bytes). The IEEE (Institute of Electrical and Electronics Engineers) assigns these OUIs to hardware manufacturers. This ensures that the first half of a MAC address identifies the company that made the device.
2. **Network Interface Controller (NIC) Specific:** This is the second half of the MAC address, consisting of the remaining **24 bits** (or 3 bytes). The manufacturer assigns this part to each device they produce, ensuring that every MAC address is globally unique.

Therefore, the OUI portion of a MAC address is 24 bits long.

Step 3: Final Answer:

The OUI consists of 24 bits. So, option (C) is the correct answer.

 Quick Tip

Remember the MAC address structure: Total 48 bits = 24 bits for OUI (manufacturer ID) + 24 bits for NIC (device ID). This simple division is a common topic in networking fundamentals.

23. Given a list numList of n elements and key value K, arrange the following steps for finding the position of the key K in the numList using binary search algorithm i.e. BinarySearch(numList, key)

- (A) Calculate $\text{mid} = (\text{first} + \text{last}) // 2$
- (B) SET $\text{first} = 0, \text{last} = n-1$
- (C) PRINT "Search unsuccessful"
- (D) WHILE $\text{first} \leq \text{last}$ REPEAT
 - IF $\text{numList}[\text{mid}] == \text{key}$
 - PRINT "Element found at position", $\text{mid}+1$
 - STOP
 - ELSE
 - IF $\text{numList}[\text{mid}] > \text{key}$, THEN $\text{last} = \text{mid}-1$
 - ELSE $\text{first} = \text{mid}+1$

Choose the correct answer from the options given below:

- (A) (A), (B), (D), (C)
- (B) (D), (B), (A), (C)
- (C) (B), (A), (D), (C)
- (D) (B), (D), (A), (C)

Correct Answer: (D) (B), (D), (A), (C)

Solution:

Step 1: Understanding the Concept:

A binary search algorithm finds the position of a target value within a sorted array. The algorithm follows a specific sequence of steps: initialization, looping (which includes calculation and comparison), and termination (either success or failure).

Step 2: Detailed Explanation:

Let's break down the logical flow of the binary search algorithm based on the given steps:

1. **Initialization:** Before the search begins, we must initialize the boundaries of the search space. This involves setting the 'first' pointer to the beginning of the list (index 0) and the 'last' pointer to the end of the list (index n-1). This corresponds to step **(B)**.
2. **Looping Condition:** The search is an iterative process that continues as long as the search space is valid (i.e., 'first' is less than or equal to 'last'). This entire iterative structure is represented by the 'WHILE' loop in step **(D)**.
3. **Core Logic (Inside the loop):** Within each iteration of the loop, the first action is to find the middle element of the current search space. This is done by calculating the 'mid' index. This corresponds to step **(A)**. The rest of step (D) describes the comparison and shrinking of the search space.

4. **Failure Condition:** If the 'WHILE' loop terminates because 'first' has become greater than 'last', it means the element was not found in the list. In this case, a "search unsuccessful" message is printed. This corresponds to step (C), which occurs after the loop finishes.

Therefore, the correct sequence of these high-level steps is Initialization (B), followed by the Loop construct (D), which contains the calculation (A), and finally the failure case (C) after the loop. The option that best represents this logical flow is (B), (D), (A), (C).

Step 3: Final Answer:

The correct order of the steps is (B), (D), (A), (C). Therefore, option (D) is correct.

 Quick Tip

For algorithm-step arrangement questions, think about the logical flow: What needs to happen first (initialization)? What is the repetitive core of the algorithm (the loop)? What happens inside the loop (calculation, comparison)? And what happens when the algorithm ends (success/failure)?

24. In binary search after every pass of the algorithm, the search area:

- (A) gets doubled
- (B) gets reduced by half
- (C) remains same
- (D) gets reduced by one third

Correct Answer: (B) gets reduced by half

Solution:

Step 1: Understanding the Concept:

Binary search is a "divide and conquer" algorithm. Its efficiency comes from its ability to rapidly narrow down the search space. The core idea is to compare the target key with the middle element of the sorted list.

Step 2: Detailed Explanation:

In each pass (or iteration) of the binary search algorithm:

1. The middle element of the current search interval is located.
2. The target key is compared to this middle element.
3. Based on the comparison, the algorithm decides whether the key, if it exists, must be in the left half or the right half of the interval.
4. The other half is completely discarded from consideration for all subsequent passes.

By eliminating half of the remaining elements in every single pass, the search area is effectively reduced by half each time. This logarithmic reduction is what makes binary search much faster than a linear search for large datasets.

Step 3: Final Answer:

After every pass, the search area gets reduced by half. Therefore, option (B) is correct.

 Quick Tip

The time complexity of binary search is $O(\log n)$. The "log" part directly comes from this principle of halving the search space in each step. If an algorithm repeatedly cuts the problem size by a constant fraction, it will have a logarithmic time complexity.

25. For binary search, the list is in ascending order and the key is present in the list. If the middle element is less than the key, it means:

- (A) The key is in the first half.
- (B) The key is in the second half.
- (C) The key is not in the list.
- (D) The key is the middle element.

Correct Answer: (B) The key is in the second half.

Solution:

Step 1: Understanding the Concept:

Binary search fundamentally relies on the list being sorted. In an ascending sorted list, elements increase in value from left to right. Any element at a given index 'i' is less than or equal to any element at an index 'j' where $j \geq i$.

Step 2: Detailed Explanation:

Let the sorted list be 'L'. Let the middle element be 'L[mid]' and the search key be 'K'. The problem states that the list is in ascending order and 'L[mid] < K'.

- Because the list is sorted in ascending order, all elements to the left of 'mid' (the first half) must be less than or equal to 'L[mid]'.
- Since 'L[mid] < K', it follows that all elements in the first half are also less than 'K'.
- Therefore, the key 'K' cannot possibly be in the first half of the list.
- The key 'K' must, if it exists, be located in the second half of the list (i.e., at an index greater than 'mid'), where all the elements are greater than 'L[mid]'.

Step 3: Final Answer:

If the middle element is less than the key, the key must lie in the second (right) half of the list. Therefore, option (B) is correct.

 Quick Tip

Visualize a number line. If your list is sorted ascendingly and your guess ('mid') is smaller than the target ('key'), you know the target must be to the right. So, you discard the left half and continue searching on the right.

26. Arrange the following in order related to bubble sort for a list of elements:
List: [4, 9, 12, 30, 2, 6]

- (A) 4 9 12 30 2 6
 (B) 4 9 12 2 6 30
 (C) 4 9 2 12 6 30
 (D) 4 2 9 6 12 30

Choose the correct answer from the options given below:

- (A) (A), (B), (C), (D)
 (B) (A), (C), (B), (D)
 (C) (B), (A), (D), (C)
 (D) (C), (B), (D), (A)

Correct Answer: (A) (A), (B), (C), (D)

Solution:

Step 1: Understanding the Concept:

Bubble Sort is a simple sorting algorithm that repeatedly steps through the list, compares each pair of adjacent items, and swaps them if they are in the wrong order. This process is repeated until the list is sorted. The question shows various states of the list during this sorting process.

Step 2: Detailed Explanation:

Let's trace the bubble sort algorithm and see how the given states are reached.

1. **Initial State (A):** [4, 9, 12, 30, 2, 6]

This is the original, unsorted list.

2. **Pass 1:** The algorithm iterates through the list to move the largest element to the end.

- Compare 30 and 2 -> Swap. List becomes [4, 9, 12, 2, 30, 6].
- Compare 30 and 6 -> Swap. List becomes [4, 9, 12, 2, 6, 30].

This state at the end of Pass 1 matches **State (B):** [4, 9, 12, 2, 6, 30]. So, (A) is followed by (B).

3. **Pass 2** (on list from State B):

- Compare 12 and 2 -> Swap. List becomes [4, 9, 2, 12, 6, 30].

This intermediate state matches **State (C):** [4, 9, 2, 12, 6, 30]. So, (B) is followed by (C).

4. **Pass 3** (continuing from a later state): Let's assume further swaps happened. If we reach the state [4, 9, 2, 6, 12, 30], and start another pass:

- Compare 4 and 9 -> No swap.
- Compare 9 and 2 -> Swap. List becomes [4, 2, 9, 6, 12, 30].

This state matches **State (D):** [4, 2, 9, 6, 12, 30]. So, (C) is followed by (D).

The logical progression of states during the execution of bubble sort is (A) → (B) → (C) → (D).

Step 3: Final Answer:

The correct order is (A), (B), (C), (D), which corresponds to option (A).

💡 Quick Tip

In bubble sort, the largest unsorted element "bubbles up" to its correct position at the end of the unsorted part of the list in each pass. Look for this pattern when tracing the algorithm.

27. The amount of time an algorithm takes to process a given data can be called its:

- (A) Process time
- (B) Time period
- (C) Time complexity
- (D) Time bound

Correct Answer: (C) Time complexity

Solution:

Step 1: Understanding the Concept:

In computer science, we need a formal way to analyze the performance of an algorithm, specifically how its runtime or space requirements grow as the input size grows.

Step 2: Detailed Explanation:

Let's define the terms:

- **Process time:** A general term, often referring to the CPU time a process has used. It's not the standard term for algorithmic analysis.
- **Time period:** A general term for a duration of time, not specific to algorithms.
- **Time complexity:** This is the precise technical term. It's a theoretical analysis that estimates the time required for an algorithm to run as a function of the size of the input. It's usually expressed using Big O notation (e.g., $O(n)$, $O(\log n)$, $O(n^2)$).
- **Time bound:** This refers to a proven upper or lower limit on the time complexity of a problem or algorithm. While related, "time complexity" is the broader term for the analysis itself.

The question asks for the general term for the amount of time an algorithm takes, which is best described by its time complexity.

Step 3: Final Answer:

The correct term is Time complexity. Therefore, option (C) is correct.

💡 Quick Tip

Remember that "Time Complexity" isn't about the exact time in seconds. It's about how the runtime scales with the input size. An algorithm with $O(n)$ complexity will take roughly twice as long if the input size is doubled.

28. Identify the incorrect statement in the context of measures of variability:

- (A) Range is the difference between maximum and minimum values of the data.
- (B) Mean is the average of numeric values of an attribute.
- (C) Standard deviation refers to differences within the set of data of a variable.
- (D) Measures of variability is also known as measures of dispersion.

Correct Answer: (B) Mean is the average of numeric values of an attribute.

Solution:

Step 1: Understanding the Concept:

In statistics, data is described using two main types of measures:

1. **Measures of Central Tendency:** These describe the "center" or "typical" value of a dataset. Examples include the mean, median, and mode.
2. **Measures of Variability (or Dispersion):** These describe the "spread" or "scatter" of the data points. Examples include the range, variance, and standard deviation.

The question asks for a statement that is incorrect *in the context of measures of variability*.

Step 2: Detailed Explanation:

Let's analyze each statement:

- **(A):** "Range is the difference between maximum and minimum values of the data." This is the correct definition of range, which is a measure of variability. So, this statement is correct.
- **(B):** "Mean is the average of numeric values of an attribute." This is the correct definition of the mean. However, the mean is a measure of **central tendency**, not a measure of variability. Therefore, this statement is incorrect *in this context*.
- **(C):** "Standard deviation refers to differences within the set of data of a variable." This is conceptually correct. Standard deviation measures the average amount of variability or dispersion from the mean. So, this statement is correct.
- **(D):** "Measures of variability is also known as measures of dispersion." This is correct. The two terms are synonyms.

The statement about the mean, while factually correct on its own, does not describe a measure of variability and is therefore the incorrect statement in the given context.

Step 3: Final Answer:

The statement about the mean is out of place and thus incorrect in the context of measures of variability. Option (B) is the correct answer.

 Quick Tip

To easily distinguish, remember: **Central Tendency** = Center (Mean, Median).
Variability = Spread (Range, Standard Deviation).

29. Identify type of data being collected/generated in the scenerio of recording a video:

- (A) Structured Data
- (B) Ambiguous data
- (C) Unstructured data
- (D) Formal Data

Correct Answer: (C) Unstructured data

Solution:

Step 1: Understanding the Concept:

Data can be broadly classified based on its organization and format.

- **Structured Data:** Highly organized data that conforms to a strict data model, typically stored in tables with rows and columns, like in a relational database.
- **Unstructured Data:** Data that does not have a predefined data model or is not organized in a predefined manner. It often includes text, images, audio, and video.

The other options, "Ambiguous data" and "Formal Data," are not standard classifications in this context.

Step 2: Detailed Explanation:

A video is a complex binary file containing a stream of encoded image frames, audio data, and metadata. It does not fit neatly into a table structure with defined columns and rows.

Analyzing and processing video data requires specialized software and techniques. Because it lacks a predefined structure, it is a prime example of unstructured data.

Step 3: Final Answer:

A video recording is a form of unstructured data. Therefore, option (C) is correct.

💡 Quick Tip

A simple rule of thumb: If data doesn't fit into a spreadsheet or a database table easily, it's likely unstructured. This includes most multimedia content and natural language text.

30. Given data: Weight of 20 student in kgs - {35, 35, 40, 40, 40, 50, 50, 50, 50, 50, 60, 65, 65, 70, 70, 72, 75, 75, 78, 78}. Find the mode.

- (A) 50
- (B) 55
- (C) 57.4
- (D) 57

Correct Answer: (A) 50

Solution:

Step 1: Understanding the Concept:

The mode of a set of data is the value that appears most frequently. A dataset can have one mode (unimodal), more than one mode (multimodal), or no mode at all if all values appear with the same frequency.

Step 2: Key Formula or Approach:

To find the mode, we count the frequency of each unique value in the dataset.

Step 3: Detailed Explanation:

Let's count the occurrences of each weight in the given data:

- 35: appears 2 times
- 40: appears 3 times
- **50: appears 5 times**
- 60: appears 1 time
- 65: appears 2 times
- 70: appears 2 times
- 72: appears 1 time
- 75: appears 2 times
- 78: appears 2 times

The value 50 occurs 5 times, which is more frequent than any other value in the dataset.

Step 4: Final Answer:

The mode of the given data is 50. Therefore, option (A) is correct.

 Quick Tip

Remember the three M's: **Mean** is the average. **Median** is the middle value. **Mode** is the most frequent value. "Mode" and "Most" both start with "Mo" to help you remember.

31. Match List-I with List-II

List-I	List-II
(A) Primary key	(I) Total number of attributes in a table.
(B) Degree	(II) Attribute used to relate two tables.
(C) Foreign key	(III) Attribute used to uniquely identify a tuple.
(D) Constraint	(IV) A restriction on the type of data that can be inserted in a column.

Choose the correct answer from the options given below:

- (A) (A)-(I), (B)-(III), (C)-(II), (D)-(IV)
 (B) (A)-(III), (B)-(I), (C)-(II), (D)-(IV)
 (C) (A)-(I), (B)-(II), (C)-(IV), (D)-(III)
 (D) (A)-(III), (B)-(II), (C)-(I), (D)-(IV)

Correct Answer: (B) (A)-(III), (B)-(I), (C)-(II), (D)-(IV)

Solution:**Step 1: Understanding the Concept:**

This question tests fundamental terminology of the relational database model.

Step 2: Detailed Explanation:

Let's match each term in List-I with its correct definition in List-II.

- **(A) Primary key:** A primary key is a specific choice of a minimal set of attributes (columns) that uniquely identify each tuple (row) in a relation (table). This matches definition **(III)**.
- **(B) Degree:** The degree of a relation is the total number of attributes (columns) it has. This matches definition **(I)**.
- **(C) Foreign key:** A foreign key is a set of attributes in one table that refers to the primary key of another table. It is used to link the two tables, i.e., to relate them. This matches definition **(II)**.
- **(D) Constraint:** A constraint is a rule enforced on data columns in a table. It is a restriction on the type or value of data that can be inserted or updated. This matches definition **(IV)**.

The correct matching is: (A)-(III), (B)-(I), (C)-(II), (D)-(IV).

Step 3: Final Answer:

This set of matches corresponds to option (B).

💡 Quick Tip

Remember these key terms: **Tuple** = Row, **Attribute** = Column, **Relation** = Table, **Degree** = Number of Columns, **Cardinality** = Number of Rows.

32. In SQL table, the set of values which a column can take in each row is called -----.

- (A) Tuple
- (B) Attribute
- (C) Domain
- (D) Relation

Correct Answer: (C) Domain

Solution:

Step 1: Understanding the Concept:

In the theory of relational databases, each column (or attribute) is defined to hold a certain type of data drawn from a specific set of permissible values.

Step 2: Detailed Explanation:

Let's define the terms:

- **Tuple:** This refers to a single row in a table, representing a single record.
- **Attribute:** This refers to a column in a table, representing a property or characteristic of the entity.
- **Domain:** This is a theoretical concept that refers to the set of all unique, permissible values for an attribute. For example, the domain for an attribute 'Gender' might be 'Male', 'Female', 'Other'. The data type of a column (like 'VARCHAR(10)' or 'INT') is the practical implementation of a domain.

- **Relation:** This refers to the entire table, which is a collection of tuples.

The question asks for the set of possible values a column can take, which is the definition of a domain.

Step 3: Final Answer:

The correct term is Domain. Therefore, option (C) is correct.

💡 Quick Tip

Think of "Domain" as the "domain of possibilities" or the "pool of allowed values" for a specific column. For example, the domain for a 'Month' column would be the set of all month names.

33. Which of the following is synonym for Meta - data?

- (A) Data Dictionary
- (B) Database Instance
- (C) Database Schema
- (D) Data Constraint

Correct Answer: (A) Data Dictionary

Solution:

Step 1: Understanding the Concept:

Metadata is often defined as "data about data." It provides information about other data, such as its structure, type, source, author, and creation date. In a database context, metadata describes the database schema itself.

Step 2: Detailed Explanation:

Let's analyze the options:

- **Data Dictionary:** This is a repository of information about the data in a database. It stores the definitions of all schema objects (tables, views, columns, data types, constraints, etc.). It is essentially a database containing all the metadata. Thus, it is a very strong synonym.
- **Database Instance:** This is the actual data stored in the database at a specific moment. This is the data itself, not the metadata.
- **Database Schema:** This refers to the logical structure of the database. The schema is a component of the metadata, but "Data Dictionary" is the term for the system that stores all this metadata.
- **Data Constraint:** This is a rule applied to the data. A constraint is one type of metadata, not a synonym for the entire concept.

While the schema is a form of metadata, the Data Dictionary is the comprehensive collection and storage system for all metadata, making it the best synonym.

Step 3: Final Answer:

Data Dictionary is the most appropriate synonym for Metadata. Therefore, option (A) is correct.

💡 Quick Tip

Think of it this way: a book's contents are the "data". The table of contents, chapter titles, and author name are the "metadata". The library's card catalog that stores this information for all books is the "data dictionary".

34. Single row functions are also known as _____ functions.

- (A) Multi row
- (B) Group
- (C) Mathematical
- (D) Scalar

Correct Answer: (D) Scalar

Solution:

Step 1: Understanding the Concept:

SQL functions are categorized based on the number of rows they operate on to produce a result.

- **Single-row functions:** These functions operate on a single row at a time and return one result for each row processed.
- **Multi-row (or Aggregate/Group) functions:** These functions operate on a group of rows and return a single result for the entire group (e.g., 'SUM()', 'AVG()').

Step 2: Detailed Explanation:

Let's look at the options:

- **Multi row / Group:** These are the opposite of single-row functions.
- **Mathematical:** This describes a *category* of functions (like 'ABS()', 'ROUND()') but isn't a general synonym. There are also string and date single-row functions.
- **Scalar:** In mathematics and programming, a scalar is a single quantity. A function that returns a single value per invocation is a scalar function. Since single-row functions are applied to each row and return a single value for that row, "scalar function" is the correct technical synonym.

Step 3: Final Answer:

Single row functions are also known as Scalar functions. Therefore, option (D) is correct.

💡 Quick Tip

Remember: **Scalar** = Single. A single-row function processes a single row and returns a single (scalar) value. **Aggregate** = Group. An aggregate function processes a group of rows and returns a single (aggregate) value.

35. Ms Ritika wants to delete the table 'sports' permanently. Help her in selecting the correct SQL command from the following.

- (A) DELETE FROM SPORTS;
- (B) DROP SPORTS;
- (C) DROP TABLE SPORTS;
- (D) DELETE TABLE FROM SPORTS;

Correct Answer: (C) DROP TABLE SPORTS;

Solution:

Step 1: Understanding the Concept:

In SQL, there is a crucial difference between deleting data within a table and deleting the table itself.

- DELETE (Data Manipulation Language - DML): Removes rows from a table. The table structure remains.
- DROP (Data Definition Language - DDL): Removes an entire database object, such as a table, view, or index. This action is permanent.

Step 2: Detailed Explanation:

The requirement is to delete the table 'sports' *permanently*. This means removing not just the data, but the entire table structure.

- (A) DELETE FROM SPORTS;; This would delete all the rows in the table, but the empty table structure would still exist.
- (B) DROP SPORTS;; This syntax is incomplete. You must specify the type of object you are dropping (e.g., TABLE, VIEW, INDEX).
- (C) DROP TABLE SPORTS;; This is the correct DDL command to permanently delete the table named 'sports', including its structure, data, and associated objects.
- (D) DELETE TABLE FROM SPORTS;; This is invalid SQL syntax.

Step 3: Final Answer:

The correct command to permanently delete a table is DROP TABLE SPORTS;. Therefore, option (C) is correct.

 Quick Tip

Remember the difference: **DELETE** deletes rows (the content). **DROP** drops the whole container (the table itself). If you just want to empty a table quickly, consider TRUNCATE TABLE, which is faster than DELETE for large tables.

36. Which of the following are text functions?

- (A) MID()
- (B) INSTR()
- (C) SUBSTR()

(D) LENGTH()

Choose the correct answer from the options given below:

- (A) (A), (B) and (D) only
- (B) (A), (B) and (C) only
- (C) (A), (B), (C) and (D)
- (D) (B), (C) and (D) only

Correct Answer: (C) (A), (B), (C) and (D)

Solution:

Step 1: Understanding the Concept:

Text functions (or string functions) are functions in SQL that perform operations on string data types (like 'VARCHAR', 'CHAR', 'TEXT'). They can be used to manipulate, extract, or get information about strings.

Step 2: Detailed Explanation:

Let's examine each function:

- **(A) MID():** This function extracts a substring from a string. For example, `MID('ABCDE', 2, 3)` would return 'BCD'. This is a text manipulation function. (Note: 'SUBSTR' is more standard across SQL dialects, but 'MID' is common in systems like MySQL).
- **(B) INSTR():** This function returns the starting position of the first occurrence of a substring within another string. For example, `INSTR('ABCDE', 'CD')` would return 3. This is a text searching function.
- **(C) SUBSTR():** Similar to `MID()`, this function extracts a substring from a string. It is the standard SQL function for this purpose. This is a text manipulation function.
- **(D) LENGTH():** This function returns the number of characters in a given string. For example, `LENGTH('ABCDE')` would return 5. This is a text information function.

All four functions operate on text data.

Step 3: Final Answer:

All the given functions - 'MID()', 'INSTR()', 'SUBSTR()', and 'LENGTH()' - are text functions. Therefore, option (C) is the correct answer.

💡 Quick Tip

Familiarize yourself with the common SQL string functions: 'CONCAT()' (or '——'), 'LOWER()', 'UPPER()', 'SUBSTR()', 'LENGTH()', 'TRIM()', and 'REPLACE()'. They are frequently used in data cleaning and formatting.

37. Match List-I with List-II

List-I	List-II
(A) ROUTER	(I) NETWORK TOPOLOGY
(B) ETHERNET CARD	(II) NETWORK DEVICE
(C) RING	(III) NETWORK TYPE
(D) PAN	(IV) NETWORK INTERFACE CARD

Choose the correct answer from the options given below:

- (A) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)
- (B) (A)-(I), (B)-(II), (C)-(III), (D)-(IV)
- (C) (A)-(IV), (B)-(I), (C)-(II), (D)-(III)
- (D) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)

Correct Answer: (A) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)

Solution:

Step 1: Understanding the Concept:

This question tests knowledge of basic networking terminology, categorizing items as devices, topologies, types, or specific hardware components.

Step 2: Detailed Explanation:

Let's match each item from List-I to its correct category in List-II.

- **(A) ROUTER:** A router is a hardware device that connects two or more computer networks and forwards data packets between them. It is a fundamental **(II) NETWORK DEVICE**.
- **(B) ETHERNET CARD:** This is the common name for the hardware that allows a computer to connect to a wired network. The technical term for this is a **(IV) NETWORK INTERFACE CARD (NIC)**.
- **(C) RING:** A ring is a way of laying out a network where devices are connected in a circular path. This describes the network's structure or layout, which is a **(I) NETWORK TOPOLOGY**.
- **(D) PAN:** PAN stands for Personal Area Network. It is a classification of networks based on their scale and purpose (like LAN, WAN, MAN). It is a **(III) NETWORK TYPE**.

The correct set of matches is (A)-(II), (B)-(IV), (C)-(I), (D)-(III).

Step 3: Final Answer:

This matching corresponds exactly with option (A).

Quick Tip

Remember the categories: **Device** (what you use, e.g., router, switch), **Topology** (how it's laid out, e.g., star, ring, bus), **Type** (what scale, e.g., LAN, WAN, PAN), **Component** (what's inside, e.g., NIC).

38. _____ is a language used to design web pages.

- (A) Web browser
- (B) HTTP
- (C) HTML
- (D) WWW

Correct Answer: (C) HTML

Solution:

Step 1: Understanding the Concept:

The question asks to identify the specific language used for creating the structure and content of web pages.

Step 2: Detailed Explanation:

Let's analyze the given options:

- **Web browser:** This is a software application (like Chrome, Firefox) used to access and display web pages. It is not a language.
- **HTTP (Hypertext Transfer Protocol):** This is a protocol or a set of rules for transferring files (text, images, sound, video, and other multimedia files) on the World Wide Web. It is not a design language.
- **HTML (Hyper-Text Markup Language):** This is the standard markup language for creating web pages. It describes the structure of a web page using elements like headings, paragraphs, images, and links. This is the correct answer.
- **WWW (World Wide Web):** This is an information system where documents and other web resources are identified by URLs, interlinked by hypertext links, and can be accessed via the Internet. It is a system, not a language.

Step 3: Final Answer:

HTML is the language used to design and structure web pages. Therefore, option (C) is correct.

 Quick Tip

Remember the core web technologies and their roles: **HTML** for structure, **CSS** for styling/presentation, and **JavaScript** for interactivity. HTTP is the protocol that carries them.

39. Match List-I with List-II

List-I	List-II
(A) Modem	(I) An eight pin connector used with Ethernet cables for network
(B) RJ45	(II) Modulator-Demodulator
(C) Network Interface unit	(III) An organization that provides services for accessing the
(D) ISP	(IV) Ethernet card

Choose the correct answer from the options given below:

- (A) (A)-(I), (B)-(II), (C)-(III), (D)-(IV)
(B) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)
(C) (A)-(I), (B)-(III), (C)-(IV), (D)-(II)
(D) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)

Correct Answer: (B) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)

Solution:

Step 1: Understanding the Concept:

This question requires matching common networking terms with their correct definitions or descriptions.

Step 2: Detailed Explanation:

Let's match each term in List-I with its description in List-II:

- **(A) Modem:** The word "Modem" is a portmanteau of **M**odulator-**d**emodulator. It's a device that converts digital signals from a computer into analog signals for transmission over telephone lines and vice versa. This matches **(II)**.
- **(B) RJ45:** This is the standard physical connector used for Ethernet networking. It is a small plastic plug with **(I) an eight-pin** configuration.
- **(C) Network Interface Unit (NIU):** This is another name for a Network Interface Card (NIC), which is also commonly called an **(IV) Ethernet card**. It's the hardware that allows a computer to connect to a network.
- **(D) ISP:** An ISP, or Internet Service Provider, is **(III) an organization that provides services for accessing the internet**.

The correct set of matches is (A)-(II), (B)-(I), (C)-(IV), (D)-(III).

Step 3: Final Answer:

This matching corresponds to option (B).

 Quick Tip

Break down acronyms to remember their function: **M**odem = **M**odulator/**D**emodulator. **I**SP = **I**nternet **S**ervice **P**rovider. This often gives a direct clue to the answer.

40. Bandwidth of a channel is:

- (A) The range of frequencies available for transmission of data through that channel.
- (B) The path of message travels between source and destination.
- (C) The set of rules on the Internet.
- (D) The data or information that needs to be exchanged.

Correct Answer: (A) The range of frequencies available for transmission of data through that channel.

Solution:

Step 1: Understanding the Concept:

Bandwidth is a fundamental concept in telecommunications and computer networking that describes the data-carrying capacity of a communication channel.

Step 2: Detailed Explanation:

Let's analyze the provided definitions:

- **(A)** In the context of signal processing, bandwidth is defined as the difference between the highest and lowest frequencies that a communication channel can carry. This **range of frequencies** directly determines the maximum rate at which data can be transmitted. This is the correct technical definition.

- (B) This describes a communication path or channel, but not the bandwidth itself.
- (C) This describes a protocol (like TCP/IP).
- (D) This describes the message or data, not the property of the channel.

While in digital networking, "bandwidth" is often used synonymously with data transfer rate (e.g., in Mbps), its fundamental definition comes from the physical properties of the transmission medium, which is the range of frequencies it supports.

Step 3: Final Answer:

The correct definition of bandwidth is the range of frequencies available for transmission. Therefore, option (A) is correct.

 Quick Tip

Think of bandwidth like a highway. The number of lanes represents the range of frequencies. More lanes (a wider frequency range) allow more cars (data) to pass through at the same time, resulting in a higher data transfer rate.

41. Data Transfer Rate 1 Gbps is equal to:

- (A) 1024 Mbps
- (B) 1024 Kbps
- (C) 2^{30} bps
- (D) 2^{20} bps

Choose the correct answer from the options given below:

- (A) (A) and (D) only
- (B) (A) and (C) only
- (C) (B) and (D) only
- (D) (B) and (C) only

Correct Answer: (B) (A) and (C) only

Solution:

Step 1: Understanding the Concept:

Data transfer rates are measured in bits per second (bps). The prefixes Kilo (K), Mega (M), and Giga (G) can have two interpretations: decimal (powers of 10) or binary (powers of 2). In data storage, binary prefixes are standard (1 KB = 1024 bytes). In data networking, decimal prefixes are more common (1 kbps = 1000 bps), but binary definitions are also used, especially in theoretical contexts and when relating to computer architecture. This question tests the binary interpretation.

Step 2: Detailed Explanation:

Let's evaluate the options based on the binary prefix system (where K=1024, M=1024 K, G=1024 M):

- (A) **1024 Mbps:** Since 1 Giga-prefix is 1024 Mega-prefixes in the binary system, 1 Gbps is indeed equal to 1024 Mbps. So, (A) is correct.
- (B) **1024 Kbps:** This is equal to 1 Mbps, not 1 Gbps. So, (B) is incorrect.

- **(C) 2^{30} bps:** Let's expand the binary prefixes in terms of powers of 2.

$$1 \text{ Kbps} = 2^{10} \text{ bps}$$

$$1 \text{ Mbps} = 1024 \text{ Kbps} = 2^{10} \times 2^{10} \text{ bps} = 2^{20} \text{ bps}$$

$$1 \text{ Gbps} = 1024 \text{ Mbps} = 2^{10} \times 2^{20} \text{ bps} = 2^{30} \text{ bps}$$

So, (C) is correct.

- **(D) 2^{20} bps:** As calculated above, this is equal to 1 Mbps, not 1 Gbps. So, (D) is incorrect.

Both statements (A) and (C) are correct representations of 1 Gbps in the binary system.

Step 3: Final Answer:

The correct option is the one that includes both (A) and (C). Therefore, option (B) is the correct choice.

💡 Quick Tip

Remember the binary powers for data units: Kilo (2^{10}), Mega (2^{20}), Giga (2^{30}), Tera (2^{40}). This is fundamental for questions involving data storage and transfer calculations.

42. Identify the wired transmission media for the following:

"They are less expensive and commonly used in telephone lines and LANs. These cables are of two types: Unshielded and shielded."

- (A) Optical Fibre
- (B) Coaxial Cable
- (C) Twisted pair cable
- (D) Microwaves

Correct Answer: (C) Twisted pair cable

Solution:

Step 1: Understanding the Concept:

The question provides a description of a specific type of wired networking cable and asks to identify it from the given options.

Step 2: Detailed Explanation:

Let's analyze the description and compare it with the options:

- **Description:** "less expensive", "used in telephone lines and LANs", "two types: Unshielded and shielded".
- **(A) Optical Fibre:** This is a high-speed, high-capacity medium made of glass or plastic. It is generally the *most expensive* option and is not typically described as "unshielded/shielded" in the same way.
- **(B) Coaxial Cable:** This cable has a central conductor, insulation, a metallic shield, and a plastic jacket. It was used in older Ethernet networks but is now mainly used for cable television. It is more expensive than twisted pair.

- **(C) Twisted pair cable:** This is the most common type of cabling used in modern LANs (like Ethernet) and was traditionally used for telephone lines. It is known for being *inexpensive*. Critically, it comes in two main varieties: **Unshielded Twisted Pair (UTP)** and **Shielded Twisted Pair (STP)**. This perfectly matches all parts of the description.
- **(D) Microwaves:** This is a type of wireless transmission, not a wired medium.

Step 3: Final Answer:

The description perfectly matches the characteristics of Twisted pair cable. Therefore, option (C) is correct.

 Quick Tip

The keywords "Unshielded" (UTP) and "Shielded" (STP) are almost exclusively used to describe Twisted Pair cables in networking contexts. Recognizing these acronyms is a shortcut to the right answer.

43. The term Cookie is defined as:

- (A) A computer cookie is a small file or data packet which is stored by a website on the server's computer.
- (B) A cookie is edited only by the website that created it, the client's computer acts as a host to store the cookie.
- (C) Cookies are used by the user to store data on the computer.
- (D) A cookie is a security system to protect your data.

Correct Answer: (B) A cookie is edited only by the website that created it, the client's computer acts as a host to store the cookie.

Solution:

Step 1: Understanding the Concept:

An HTTP cookie is a small piece of data sent from a website and stored on the user's computer by the user's web browser while the user is browsing. They are used for state management, personalization, and tracking.

Step 2: Detailed Explanation:

Let's evaluate the given definitions:

- **(A)** This states that the cookie is stored on the *server's* computer. This is incorrect. Cookies are stored on the *client's* (user's) computer.
- **(B)** This statement correctly identifies two key properties: (1) Cookies are subject to the same-origin policy, meaning they can generally only be read and edited by the website that created them. (2) The client's computer is the host that stores the cookie file. This is the most accurate description.
- **(C)** This states that cookies are used by the *user* to store data. This is misleading. Cookies are used by the *website* to store data (like session IDs) on the user's computer. The user does not typically interact with them directly.

- **(D)** This describes a cookie as a security system. This is incorrect. In fact, cookies can sometimes pose security and privacy risks (e.g., cross-site scripting, tracking).

Step 3: Final Answer:

Statement (B) provides the most accurate and comprehensive definition among the choices. Therefore, option (B) is correct.

💡 Quick Tip

Remember the key facts about cookies: a website puts them on **your** computer (the client), and only that website can typically read them back later. They are for the website's benefit, not yours directly.

44. Identify the concept behind the below-given scenario:

If an attacker limits or stops an authorized user to access a service, device or any such resource by overloading that resource with illegitimate requests.

- (A) Snooping
- (B) Eavesdropping
- (C) Denial of Service
- (D) Plagiarism

Correct Answer: (C) Denial of Service

Solution:

Step 1: Understanding the Concept:

The scenario describes a specific type of cyberattack. We need to identify the correct term for an attack whose goal is to make a resource unavailable to its intended users.

Step 2: Detailed Explanation:

Let's define the terms:

- **Snooping/Eavesdropping:** These are passive attacks where an attacker secretly listens to or intercepts private communication. The goal is to obtain information, not to disrupt service.
- **Denial of Service (DoS):** This is an attack meant to shut down a machine or network, making it inaccessible to its intended users. The method described in the scenario—overloading a resource with illegitimate requests—is a classic example of a DoS attack.
- **Plagiarism:** This is the act of using someone else's work or ideas without giving proper credit. It is an issue of academic/intellectual integrity, not a cyberattack in this context.

The scenario perfectly matches the definition of a Denial of Service attack.

Step 3: Final Answer:

The described attack is a Denial of Service. Therefore, option (C) is correct.

💡 Quick Tip

Think about the goal of the attack. Is it to steal information? (Snooping/Spying). Is it to break the service? (Denial of Service). This will help you categorize the attack correctly.

45. The HTTPS based websites require:

- (A) Search Engine Optimization
- (B) Digital authenticity
- (C) WWW Certificate
- (D) SSL Digital Certificate

Correct Answer: (D) SSL Digital Certificate

Solution:

Step 1: Understanding the Concept:

HTTPS stands for Hypertext Transfer Protocol Secure. It is the secure version of HTTP, the protocol over which data is sent between a web browser and a website. The 'S' at the end of HTTPS stands for 'Secure', which means all communications are encrypted. This encryption is achieved using a protocol called SSL (Secure Sockets Layer) or its modern successor, TLS (Transport Layer Security).

Step 2: Detailed Explanation:

To establish a secure SSL/TLS connection, a website must have a digital certificate installed on its web server.

- **SSL Digital Certificate:** This certificate serves two purposes: it encrypts the data exchanged between the server and the client, and it verifies the identity and authenticity of the website's owner. A browser will check the validity of this certificate before establishing a secure connection. This is a mandatory requirement for HTTPS.
- **Other Options:** Search Engine Optimization (A) is a marketing practice, not a technical requirement. "Digital authenticity" (B) is a concept that is *provided by* the certificate, but it is not the required item itself. "WWW Certificate" (C) is too generic and not the standard term.

Step 3: Final Answer:

An SSL (or TLS) Digital Certificate is required to enable HTTPS on a website. Therefore, option (D) is correct.

💡 Quick Tip

Remember **HTTPS = HTTP + SSL/TLS**. And **SSL/TLS requires a Digital Certificate**. This chain of dependencies is key to understanding web security.

46. State the output of the following query.

SELECT LENGTH(MID('INFORMATICS PRACTICES',5,-5));

- (A) NO OUTPUT
- (B) 5
- (C) 0
- (D) ERROR

Correct Answer: (C) 0

Solution:

Step 1: Understanding the Concept:

The query involves nested SQL functions. We must evaluate the inner function first and then pass its result to the outer function. The query uses 'MID()' (a substring function, common in MySQL) and 'LENGTH()'.

Step 2: Key Formula or Approach:

The evaluation proceeds from the inside out: 1. Evaluate MID('INFORMATICS PRACTICES', 5, -5). 2. Evaluate LENGTH() on the result of step 1.

Step 3: Detailed Explanation:

- **Inner Function MID():** The syntax for MID(string, start, length) extracts a substring of a given 'length' starting from the 'start' position. In many SQL dialects, including MySQL, if the 'length' parameter is negative, the function is designed to return an empty string (''). It does not produce an error.
- **Outer Function LENGTH():** The result of the 'MID()' function, an empty string (''), is then passed to the 'LENGTH()' function. The length of an empty string is 0.

So, the query effectively becomes SELECT LENGTH('');, which evaluates to 0.

Step 4: Final Answer:

The output of the query is 0. Therefore, option (C) is correct.

 Quick Tip

Be aware of how different SQL functions handle edge cases and invalid parameters. A negative length in 'SUBSTR' or 'MID' is a common tricky question. In many systems, it results in an empty string or NULL rather than an error.

47. Match List-I with List-II

List-I	List-II
(A) group by	(I) math function
(B) mid()	(II) aggregate function
(C) count()	(III) having
(D) mod()	(IV) text function

Choose the correct answer from the options given below:

- (A) (A)-(III), (B)-(IV), (C)-(II), (D)-(I)
- (B) (A)-(I), (B)-(II), (C)-(III), (D)-(IV)
- (C) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)
- (D) (A)-(IV), (B)-(III), (C)-(II), (D)-(I)

Correct Answer: (A) (A)-(III), (B)-(IV), (C)-(II), (D)-(I)

Solution:

Step 1: Understanding the Concept:

This question requires categorizing various SQL keywords and functions into their respective types or associating them with closely related clauses.

Step 2: Detailed Explanation:

Let's match each item from List-I to its correct category/association in List-II.

- **(A) group by:** The 'GROUP BY' clause is used to group rows. The 'HAVING' clause is specifically designed to filter these groups based on aggregate conditions. Therefore, 'group by' is most closely associated with **(III) having**.
- **(B) mid():** As seen in a previous question, 'mid()' is a function that extracts a substring from a string. It is a **(IV) text function**.
- **(C) count():** The 'count()' function calculates the number of rows in a group. It is a classic **(II) aggregate function**.
- **(D) mod():** The 'mod()' function calculates the remainder of a division. It is a **(I) math function**.

The correct set of matches is (A)-(III), (B)-(IV), (C)-(II), (D)-(I).

Step 3: Final Answer:

This matching corresponds exactly with option (A).

 Quick Tip

Remember the standard SQL query structure: 'SELECT ... FROM ... WHERE ... GROUP BY ... HAVING ... ORDER BY'. This shows the close relationship between 'GROUP BY' and 'HAVING'.

48. What will be the format of the output of the NOW() function?

- (A) HH:MM:SS
- (B) YYYY-MM-DD HH:MM:SS
- (C) HH:MM:SS YYYY-MM-DD
- (D) YYYY-DD-MM HH:MM:SS

Correct Answer: (B) YYYY-MM-DD HH:MM:SS

Solution:

Step 1: Understanding the Concept:

The 'NOW()' function in SQL is a standard function that returns the current date and time from the database server. The output has a specific, standardized format.

Step 2: Detailed Explanation:

The standard format for a DATETIME or TIMESTAMP value in SQL, which is what 'NOW()' returns, is a combination of the date and the time.

- The date part is formatted as Year-Month-Day: **YYYY-MM-DD**.

- The time part is formatted as Hour:Minute:Second: **HH:MM:SS**.

These two parts are combined, separated by a space, to give the full format: **YYYY-MM-DD HH:MM:SS**. Let's check the options:

- (A) HH:MM:SS: This is only the time part.
- (B) YYYY-MM-DD HH:MM:SS: This is the correct standard format.
- (C) HH:MM:SS YYYY-MM-DD: The order is incorrect.
- (D) YYYY-DD-MM HH:MM:SS: The month and day are swapped (DD-MM instead of MM-DD), which is not the standard format.

Step 3: Final Answer:

The correct format for the output of the 'NOW()' function is YYYY-MM-DD HH:MM:SS. Therefore, option (B) is correct.

💡 Quick Tip

Remember the international standard date format (ISO 8601) is 'YYYY-MM-DD'. Most database systems use this format for consistency, which makes it easy to remember.

49. State the output of the following query:

```
SELECT ROUND(9873.567, -2);
```

- (A) 9900
- (B) 9873
- (C) 9800
- (D) 9873.5

Correct Answer: (A) 9900

Solution:

Step 1: Understanding the Concept:

The 'ROUND(number, decimals)' function in SQL rounds a number to a specified number of decimal places. A positive 'decimals' value rounds to the right of the decimal point. A negative 'decimals' value rounds to the left of the decimal point (i.e., to the nearest 10, 100, 1000, etc.).

Step 2: Key Formula or Approach:

We need to evaluate ROUND(9873.567, -2). The second argument, -2, indicates that we need to round the number to the nearest hundred.

Step 3: Detailed Explanation:

1. The number to be rounded is 9873.567.
2. The rounding precision is -2, which corresponds to the hundreds place.
3. We look at the digit one place to the right of the hundreds place, which is the tens place. The digit in the tens place is 7.

4. The rule for rounding is: if the digit is 5 or greater, we round up. If it's less than 5, we round down.
5. Since 7 is greater than 5, we round the hundreds digit (8) up to 9.
6. All digits to the right of the hundreds place become zero.
7. Therefore, 9873.567 rounds to 9900.

Step 4: Final Answer:

The result of the query is 9900. Therefore, option (A) is correct.

💡 Quick Tip

For 'ROUND(number, d)':

- If $d > 0$, round to d places after the decimal.
- If $d = 0$, round to the nearest whole number.
- If $d < 0$, round to $|d|$ places before the decimal (to the nearest 10, 100, etc.).

50. Given the following table:

EMPLOYEE

ENO	SALARY
A101	4000
B109	NULL
C508	6000
A305	2000

State the output of the following query based on the above given table.

SELECT COUNT(SALARY) FROM EMPLOYEE;

- (A) 4
- (B) 2
- (C) 3
- (D) 1

Correct Answer: (C) 3

Solution:

Step 1: Understanding the Concept:

In SQL, there are two ways to use the COUNT() aggregate function:

- COUNT(*): This counts the total number of rows in the table.
- COUNT(column_name): This counts the number of rows where the specified column has a non-NULL value. It explicitly ignores any rows where the value in that column is NULL.

Step 2: Detailed Explanation:

The query is SELECT COUNT(SALARY) FROM EMPLOYEE;. This command instructs the database to count the number of records in the EMPLOYEE table where the SALARY column is not NULL. Let's examine the SALARY column:

- Row 1 (A101): 4000 (Not NULL)
- Row 2 (B109): NULL (Will be ignored)
- Row 3 (C508): 6000 (Not NULL)
- Row 4 (A305): 2000 (Not NULL)

There are three rows with non-NULL values in the **SALARY** column. Therefore, the function will return 3.

Step 3: Final Answer:

The output of the query will be 3. Thus, option (C) is the correct answer.

 Quick Tip

A very common exam question involves the difference between **COUNT(*)** and **COUNT(column)**. Remember, the asterisk (*) counts all rows, while specifying a column name counts only the non-NULL values in that column.

51. Match List-I with List-II

List-I	List-II
(A) SUBSTR()	(I) To find the position of a substring in a string.
(B) TRIM()	(II) To extract a substring from a string.
(C) INSTR()	(III) To extract characters from the left side of the string.
(D) LEFT()	(IV) To remove spaces from both the sides of the string.

Choose the correct answer from the options given below:

- (A) (A)-(I), (B)-(III), (C)-(II), (D)-(IV)
- (B) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)
- (C) (A)-(III), (B)-(IV), (C)-(II), (D)-(I)
- (D) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)

Correct Answer: (D) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)

Solution:

Step 1: Understanding the Concept:

This question tests the knowledge of common string manipulation functions available in SQL. Each function performs a specific operation on string data.

Step 2: Detailed Explanation:

Let's match each function in List-I with its correct description in List-II:

- **(A) SUBSTR():** This function is used **(II) to extract a substring from a string**. For example, SUBSTR('Hello World', 7, 5) returns 'World'.
- **(B) TRIM():** This function is used **(IV) to remove spaces from both the sides of the string**. For example, TRIM(' Hello ') returns 'Hello'.
- **(C) INSTR():** This function (In-String) is used **(I) to find the position of a substring in a string**. For example, INSTR('Hello World', 'World') returns 7.

- **(D) LEFT():** This function is used **(III) to extract characters from the left side of the string**. For example, `LEFT('Hello World', 5)` returns 'Hello'.

The correct matching is: (A)-(II), (B)-(IV), (C)-(I), (D)-(III).

Step 3: Final Answer:

This combination matches option (D).

 Quick Tip

To remember these functions, think of their names: **SUBSTR = SUBSTRing**. **INSTR = INSTRing (position)**. **TRIM = trim edges**. **LEFT/RIGHT = extract from that side**.

52. Arrange the following statements to create a series from a dictionary.

- (A) Print the series**
- (B) Import the pandas library**
- (C) Create the series**
- (D) Create a dictionary**

Choose the correct answer from the options given below:

- (A) (A), (B), (C), (D)**
- (B) (A), (C), (B), (D)**
- (C) (D), (B), (C), (A)**
- (D) (C), (B), (D), (A)**

Correct Answer: (C) (D), (B), (C), (A)

Solution:

Step 1: Understanding the Concept:

The question asks for the logical sequence of steps required to create a pandas Series object from a Python dictionary and display it. This involves a standard workflow in data manipulation using Python libraries.

Step 2: Detailed Explanation:

Let's determine the correct logical order of operations:

1. **Create the source data:** Before you can create a Series, you need the data from which it will be created. Therefore, the first step must be to **(D) Create a dictionary**.
2. **Import necessary tools:** To work with pandas objects like Series, you must first import the pandas library into your Python script. So, the second step is to **(B) Import the pandas library**.
3. **Perform the main action:** Once the data and the library are ready, you can call the pandas function to create the Series object from the dictionary. This is step **(C) Create the series**.
4. **Display the result:** After the Series object is created in memory, the final step is to display it to the user. This is done by **(A) Print the series**.

The correct logical flow is $D \rightarrow B \rightarrow C \rightarrow A$.

Example Code:

```
# (D) Create a dictionary
my_dict = {'a': 1, 'b': 2, 'c': 3}

# (B) Import the pandas library
import pandas as pd

# (C) Create the series
my_series = pd.Series(my_dict)

# (A) Print the series
print(my_series)
```

Step 3: Final Answer:

The correct order of statements is (D), (B), (C), (A). This corresponds to option (C).

 Quick Tip

For any programming task involving a library, the general order is: 1. Prepare your data. 2. Import the library. 3. Use the library's functions on your data. 4. Output the result.

53. Given the following DataFrame 'df':

PNO	NAME
111	ROHAN
222	MEETA
333	SEEMA
444	SHALU
555	POONAM

Select the correct commands from the following to display the last five rows:

- (A) `print(df.tail(-5))`
- (B) `print(df.head(-5))`
- (C) `print(df.tail())`
- (D) `print(df.tail(5))`

Choose the correct answer from the options given below:

- (A) (A), (C) and (D) only
- (B) (A) and (B) only
- (C) (A), (B), (C) and (D)
- (D) (C) and (D) only

Correct Answer: (D) (C) and (D) only

Solution:

Step 1: Understanding the Concept:

In the pandas library, the `.head()` method is used to view the first few rows of a DataFrame or Series, while the `.tail()` method is used to view the last few rows.

Step 2: Detailed Explanation:

Let's analyze each command:

- **(A) `print(df.tail(-5))`:** Using a negative number in `.tail()` is used to select all rows *except* the first 'n' rows. It does not display the last 5 rows. In this case, it would display zero rows.
- **(B) `print(df.head(-5))`:** Similarly, using a negative number in `.head()` selects all rows *except* the last 'n' rows. It does not display the last 5 rows.
- **(C) `print(df.tail())`:** When the `.tail()` method is called without any arguments, it defaults to displaying the last 5 rows. Since the DataFrame has exactly 5 rows, this command will display all of them, which are the last five. This is correct.
- **(D) `print(df.tail(5))`:** This command explicitly asks for the last 5 rows of the DataFrame. This is also correct.

Both (C) and (D) are correct commands to achieve the desired output.

Step 3: Final Answer:

The correct commands are (C) and (D). Therefore, option (D) is the correct choice.

 Quick Tip

Remember: `.head()` is for the top (first rows), `.tail()` is for the bottom (last rows). If you don't specify a number, both default to 5.

54. Given the following series 'ser1':

```
0    60
1    80
2    20
3    50
4   100
5    70
```

State the output of the following command:

```
print(ser1 >= 70)
```

```
1    80
(A) 4   100
    5    70
    0    F
    1    T
    2    F
(B) 3    F
    4    T
    5    T
```

- 0 False
 1 True
 (C) 2 False
 3 False
 4 True
 5 True
 (D) 1 80
 4 100

Correct Answer: (C)

Solution:

Note: The OCR for option (B) and (C) might be similar. The standard pandas output uses the full words "True" and "False".

Step 1: Understanding the Concept:

When a comparison operator (like `>=`, `<`, `==`) is used with a pandas Series and a scalar value, it performs a vectorized operation. This means the comparison is applied to each element of the Series individually, and the result is a new Series of the same size containing boolean values (True or False).

Step 2: Detailed Explanation:

The command is `print(ser1 >= 70)`. Let's apply this comparison to each element of `ser1`:

- Index 0: Is 60 `>= 70`? **False**
- Index 1: Is 80 `>= 70`? **True**
- Index 2: Is 20 `>= 70`? **False**
- Index 3: Is 50 `>= 70`? **False**
- Index 4: Is 100 `>= 70`? **True**
- Index 5: Is 70 `>= 70`? **True**

The output will be a new boolean Series with the same index as the original, containing these True/False results. Option (A) shows the result of *filtering* (`ser1[ser1 >= 70]`), not the boolean comparison itself. Option (C) correctly shows the boolean Series.

Step 3: Final Answer:

The command produces a boolean Series indicating the result of the comparison for each element. This matches option (C).

💡 Quick Tip

Pay close attention to whether the question asks for the result of a boolean comparison (which gives a True/False series) or the result of filtering using that comparison (which gives a subset of the original data).

341. Which of the following are the correct commands to delete a column from the DataFrame 'df1'?

- (A) `df1 = df1.drop('column_name', axis=1)`
- (B) `df1.drop('column_name', axis=columns, inplace=True)`
- (C) `df1.drop('column_name', axis='columns', inplace=True)`
- (D) `df1.drop('column_name', axis=1, inplace=True)`

Choose the correct answer from the options given below:

- (A) (A), (C) and (D) only
- (B) (A), (B) and (C) only
- (C) (A), (B), (C) and (D)
- (D) (B), (C) and (D) only

Correct Answer: (A) (A), (C) and (D) only

Solution:

Step 1: Understanding the Concept:

In pandas, the `.drop()` method is used to remove rows or columns from a DataFrame. To remove a column, you must specify the column name(s) and indicate that you are operating on columns using the `axis` parameter. The change can be applied either by reassigning the result to a variable or by using the `inplace=True` parameter.

Step 2: Detailed Explanation:

Let's analyze each command:

- (A) `df1 = df1.drop('column_name', axis=1)`: This is a correct way. `axis=1` specifies that we are dropping a column. The `.drop()` method returns a new DataFrame without the specified column, which is then reassigned back to `df1`.
- (B) `df1.drop('column_name', axis=columns, inplace=True)`: This is incorrect. The value for the `axis` parameter should be the integer 1 or the string 'columns'. The bareword `columns` would cause a 'NameError' unless it was a variable defined elsewhere, which is not standard practice.
- (C) `df1.drop('column_name', axis='columns', inplace=True)`: This is correct. Using the string 'columns' is an alternative to `axis=1`. The `inplace=True` parameter modifies the original DataFrame `df1` directly without needing reassignment.
- (D) `df1.drop('column_name', axis=1, inplace=True)`: This is also correct. It uses `axis=1` and modifies the DataFrame in place.

Commands (A), (C), and (D) are all valid and correct ways to delete a column. Command (B) has a syntax error.

Step 3: Final Answer:

The correct commands are (A), (C), and (D). This corresponds to option (A).

 Quick Tip

Remember for `.drop()`: `axis=0` or `axis='index'` for rows, and `axis=1` or `axis='columns'` for columns. Use `inplace=True` to modify the DataFrame directly, otherwise you must reassign the result (e.g., `df = df.drop(...)`).

56. Which of the following is the correct command to display the top 2 records of Series "s1"?

- (A) `print(s1.head(2))`
- (B) `print(tail.head())`
- (C) `print(tail.head(2))`
- (D) `print(s1.tail.head(2))`

Correct Answer: (A) `print(s1.head(2))`

Solution:

Note: The original question OCR has garbled text ("Series 'Tail?'"). The solution assumes the Series is named 's1'.

Step 1: Understanding the Concept:

To display the first 'n' records of a pandas Series or DataFrame, the `.head(n)` method is used. "Top records" refers to the first records in the object's order.

Step 2: Detailed Explanation:

Let's analyze the syntax of the given options:

- **(A) `print(s1.head(2))`:** This command correctly calls the `.head()` method on the Series object `s1` and passes 2 as an argument to specify that the top 2 records should be displayed. This is the correct syntax.
- **(B) `print(tail.head())`:** This is incorrect. `tail` is a method, not an object on which `head` can be called. This would result in a 'NameError'.
- **(C) `print(tail.head(2))`:** This is also incorrect for the same reason as (B).
- **(D) `print(s1.tail.head(2))`:** This is incorrect syntax. You cannot chain `tail` and `head` in this manner.

Step 3: Final Answer:

The correct and valid command is `print(s1.head(2))`. Therefore, option (A) is the correct answer.

Quick Tip

For pandas objects (like a Series 's' or DataFrame 'df'):

- '`s.head(n)`' gives the first 'n' items.
- '`s.tail(n)`' gives the last 'n' items.

57. While transferring the data from a DataFrame to CSV file, if we do not want the column names to be saved to the file on the disk, the parameter to be used in `to_csv()` is -----.

- (A) `header=False`
- (B) `index=False`
- (C) `header=True`
- (D) `index=True`

Correct Answer: (A) `header=False`

Solution:

Step 1: Understanding the Concept:

The pandas `.to_csv()` method is used to write a DataFrame to a comma-separated values (CSV) file. This method has several parameters to control the output format, including whether to write the column names (the header) and the row labels (the index).

Step 2: Detailed Explanation:

Let's look at the relevant parameters for `df.to_csv()`:

- **header:** This parameter controls the writing of column names. By default, it is `True`, meaning the column names are written as the first line of the CSV file. To prevent the column names from being saved, you must set this parameter to `False`.
- **index:** This parameter controls the writing of the DataFrame index. By default, it is `True`. To prevent the index from being saved, you would set `index=False`.

The question specifically asks how to prevent the **column names** from being saved. The correct parameter and value for this is `header=False`.

Step 3: Final Answer:

To omit column names when writing to a CSV file, you use the parameter `header=False`. Therefore, option (A) is correct.

💡 Quick Tip

A very common use case is to save a CSV without the default integer index. For this, you would use `df.to_csv('file.csv', index=False)`. Remember: `header` is for columns, `index` is for rows.

58. HTTP is -----.

- (A) a set of rule for email communication.
- (B) a file transfer protocol.
- (C) a set of rules which is used to retrieve linked web pages across the web.
- (D) a set of rules used for secure transmission of data.

Correct Answer: (C) a set of rules which is used to retrieve linked web pages across the web.

Solution:

Step 1: Understanding the Concept:

HTTP stands for **H**ypertext **T**ransfer **P**rotocol. It is the foundation of data communication for the World Wide Web. It defines the format of messages and how web browsers and web servers should interact to request and serve web pages.

Step 2: Detailed Explanation:

Let's analyze the options:

- **(A)** The set of rules for email communication is primarily SMTP (Simple Mail Transfer Protocol), POP3, and IMAP. This is incorrect.

- (B) While HTTP does transfer files, the protocol specifically named "File Transfer Protocol" is FTP. Calling HTTP just "a file transfer protocol" is vague and less accurate than other options.
- (C) This accurately describes the primary function of HTTP. It is the protocol (set of rules) used by web clients to request and retrieve web pages (which are linked hypertext documents) from web servers. This is the best definition.
- (D) The set of rules for *secure* transmission of web data is HTTPS (HTTP Secure), which uses an encryption layer like SSL/TLS. Standard HTTP itself is not secure.

Step 3: Final Answer:

The most accurate description of HTTP is that it's a set of rules used to retrieve linked web pages. Therefore, option (C) is correct.

💡 Quick Tip

Break down the acronym: **HyperText Transfer Protocol**. Its job is to transfer hypertext (web pages). This directly points to the correct definition.

59. Put the following statistical measures in order starting from the one that provides the simplest measure of data to the most complex one that describes data spread.

- (A) Mode
- (B) Mean
- (C) Median
- (D) Range
- (E) Standard Deviation

Choose the correct answer from the options given below:

- (A) (A), (C), (B), (D), (E)
- (B) (A), (B), (C), (D), (E)
- (C) (B), (A), (D), (E), (C)
- (D) (C), (E), (B), (D), (A)

Correct Answer: (A) (A), (C), (B), (D), (E)

Solution:

Step 1: Understanding the Concept:

The question asks to order statistical measures based on their complexity, starting from the simplest measure of data (central tendency) to the most complex measure of spread (variability). This requires understanding both the concepts and the calculations involved.

Step 2: Detailed Explanation:

Let's analyze the complexity of calculating each measure:

1. Measures of Central Tendency (Simplest overall):

- (A) **Mode:** Simplest to find. You just need to count the frequency of each value. No complex arithmetic is needed.

- **(C) Median:** Requires sorting the data first to find the middle value. Sorting is more complex than just counting.
- **(B) Mean:** Requires summing all values and then dividing by the count. This involves arithmetic on the entire dataset.

So, the order of complexity for central tendency is Mode $\rightarrow$ Median $\rightarrow$ Mean.

2. Measures of Spread (Generally more complex):

- **(D) Range:** Simple to calculate. It only requires finding the maximum and minimum values and subtracting them.
- **(E) Standard Deviation:** This is the most complex. It requires calculating the mean, finding the difference of each data point from the mean, squaring those differences, averaging the squares (variance), and finally taking the square root.

The question asks for a single order from simplest measure of data to most complex measure of spread. A logical ordering combines these, starting with the simplest central tendency measures and ending with the most complex spread measure. The sequence (A), (C), (B), (D), (E) follows this logic: Mode (simplest), Median, Mean, Range (simple spread), Standard Deviation (complex spread).

Step 3: Final Answer:

The logical order of complexity is (A) Mode, (C) Median, (B) Mean, (D) Range, (E) Standard Deviation. This corresponds to option (A).

💡 Quick Tip

To judge complexity, think about the steps needed. Counting (Mode) is easier than Sorting (Median), which is often simpler than arithmetic on all data (Mean). Calculations involving squares and square roots (Standard Deviation) are the most complex.

60. `Dataframe.quantile()` is used to get the quartiles. The second quantile is known as -----.

- (A) mode
- (B) mean
- (C) standard deviation
- (D) median

Correct Answer: (D) median

Solution:

Step 1: Understanding the Concept:

Quantiles are points taken at regular intervals from the cumulative distribution function of a random variable. Quartiles are a specific type of quantile that divides the number of data points into four parts, or quarters, of more or less equal size.

Step 2: Detailed Explanation:

Let's define the quartiles:

- **First Quartile (Q1):** This is the value under which 25% of the data points are found. It is also the 0.25 quantile.
- **Second Quartile (Q2):** This is the value under which 50% of the data points are found. It splits the data set in half. This is the definition of the **median**. It is also the 0.5 quantile.
- **Third Quartile (Q3):** This is the value under which 75% of the data points are found. It is also the 0.75 quantile.

The mean (average), mode (most frequent value), and standard deviation (measure of spread) are different statistical measures and are not equivalent to the second quartile.

Step 3: Final Answer:

The second quartile (Q2) is, by definition, the median of the dataset. Therefore, option (D) is correct.

 Quick Tip

Remember this important equivalence: **2nd Quartile = 5th Decile = 50th Percentile = Median**. They are all different names for the same thing: the middle value of a dataset.

61. Consider the given DataFrame 'df4'.

	NAME	PERIODIC	ENG	MATHS	SCIENCE
12	MEENA	1	40	60	80
17	REENA	2	60	30	10
18	TEENA	3	80	60	40
20	SHEENA	2	90	20	70

Select the correct command to get the following output:

```
NAME      MEENAREENATEENASHEENA
PERIODIC  8
ENG       270
MATHS     170
SCIENCE   200
```

- (A) `print(df4.sum(numeric_only=True))`
- (B) `print(df4.sum(numeric_only=False))`
- (C) `print(df4.sum())`
- (D) `print(df4.count(numeric_only=True))`

Correct Answer: (C) `print(df4.sum())`

Solution:

Step 1: Understanding the Concept:

The pandas `.sum()` method, when applied to a DataFrame, calculates the sum for each column. By default (`axis=0`), it operates column-wise. For numeric columns, it calculates the arithmetic sum. For non-numeric columns (like strings), it performs concatenation.

Step 2: Detailed Explanation:

Let's analyze the desired output and see which command produces it:

- **NAME:** The output is 'MEENAREENATEENASHEENA'. This is the concatenation of all the strings in the NAME column.
- **PERIODIC:** The output is 8. This is the sum of the values in the PERIODIC column ($1 + 2 + 3 + 2 = 8$).
- **ENG:** The output is 270. This is the sum of the values in the ENG column ($40 + 60 + 80 + 90 = 270$).
- **MATHS:** The output is 170. This is the sum of the values in the MATHS column ($60 + 30 + 60 + 20 = 170$).
- **SCIENCE:** The output is 200. This is the sum of the values in the SCIENCE column ($80 + 10 + 40 + 70 = 200$).

Now let's evaluate the commands:

- **(A) `print(df4.sum(numeric_only=True))`:** This would only calculate the sum for the numeric columns and would ignore the 'NAME' column entirely. The output would not include the concatenated string.
- **(B) `print(df4.sum(numeric_only=False))`:** Setting `numeric_only=False` is the default behavior. This command is functionally identical to (C).
- **(C) `print(df4.sum())`:** This is the standard command to sum each column. It will sum numeric columns and concatenate string columns, exactly matching the desired output.
- **(D) `print(df4.count(numeric_only=True))`:** This would count the number of non-null entries in each numeric column, not sum them.

Both (B) and (C) produce the correct result, but (C) is the most common and direct way to write the command.

Step 3: Final Answer:

The command `print(df4.sum())` correctly produces the specified output. Therefore, option (C) is the correct answer.

💡 Quick Tip

Remember that pandas aggregation functions like `.sum()` and `.mean()` often have a `numeric_only` parameter. The default behavior is usually to attempt the operation on all columns, which leads to concatenation for strings when using `.sum()`.

62. The process of changing the structure of a DataFrame using function `pivot()` is known as _____.

- (A) Transpose
- (B) Reindexing
- (C) Resetting
- (D) Reshaping

Correct Answer: (D) Reshaping

Solution:**Step 1: Understanding the Concept:**

In pandas, the term "reshaping" refers to changing the layout or structure of the data in a DataFrame. This can involve transforming data from a "long" format to a "wide" format or vice versa, or rearranging which data appears in the rows and columns.

Step 2: Detailed Explanation:

Let's analyze the given options:

- **Transpose:** This is a specific type of reshaping where rows and columns are swapped. It's done using the `.T` attribute.
- **Reindexing:** This refers to changing the row or column labels (the index).
- **Resetting:** This typically refers to `reset_index()`, which converts the index into a column.
- **Reshaping:** This is the general term for transforming the structure of a DataFrame. The `pivot()` function is a primary tool for reshaping. It takes three arguments—`index`, `columns`, and `values`—to transform a DataFrame from a long format (where variables are in a single column) to a wide format (where each variable gets its own column). Other reshaping functions include `melt()`, `stack()`, and `unstack()`.

Since `pivot()` fundamentally alters the structure of the data by moving values between rows and columns, it is a quintessential reshaping operation.

Step 3: Final Answer:

The process is known as Reshaping. Therefore, option (D) is correct.

💡 Quick Tip

Remember that **Reshaping** is the broad category. **Pivoting** (wide format), **Melting** (long format), **Stacking**, and **Unstacking** are all specific methods used for reshaping data.

63. After establishing the connection to fetch the data from the table of a database in SQL into a DataFrame, which of the following function will be used?

- (A) `pandas.read_sql_query()`
- (B) `pandas.read_sql_table()`
- (C) `pandas.read_sql_query_table()`
- (D) `pandas.read_sql()`

Choose the correct answer from the options given below:

- (A) (A), (B) and (D) only
- (B) (A), (B) and (C) only
- (C) (A), (B), (C) and (D)
- (D) (B), (C) and (D) only

Correct Answer: (A) (A), (B) and (D) only

Solution:

Step 1: Understanding the Concept:

The pandas library provides functions to read data directly from a SQL database into a DataFrame. These functions require an active database connection object and a way to specify the data to be retrieved.

Step 2: Detailed Explanation:

Let's examine the primary functions for this purpose:

- (A) `pandas.read_sql_query(sql, con)`: This function takes a full SQL query as a string (`sql`) and a database connection object (`con`) and returns the result of the query as a DataFrame. This is a very common and flexible method.
- (B) `pandas.read_sql_table(table_name, con)`: This function reads an entire database table specified by `table_name` into a DataFrame. It's simpler than writing a 'SELECT * FROM ...' query but less flexible.
- (C) `pandas.read_sql_query_table()`: This function does not exist in the pandas library. It seems to be a conflation of the other two function names.
- (D) `pandas.read_sql(sql, con)`: This is a convenient wrapper function. It can intelligently delegate to either `read_sql_query()` or `read_sql_table()`. If you pass a table name to `sql`, it calls `read_sql_table`; if you pass a full query, it calls `read_sql_query`.

Therefore, functions (A), (B), and (D) are all valid for fetching data from a SQL database. Function (C) is invalid.

Step 3: Final Answer:

The correct option must include (A), (B), and (D) only. This corresponds to option (A).

💡 Quick Tip

For most practical purposes, `pd.read_sql_query()` is the most versatile function to remember, as it allows you to use the full power of SQL (joins, aggregations, filtering) to select and shape your data before it even gets to pandas.

64. _____ in matplotlib also known as correlation plot, because they show how two variables are related.

- (A) Bar graph
- (B) Pie chart
- (C) Box plot
- (D) Scatter plot

Correct Answer: (D) Scatter plot

Solution:

Step 1: Understanding the Concept:

The question asks to identify the type of plot that is specifically designed to visualize the relationship or correlation between two continuous variables.

Step 2: Detailed Explanation:

Let's analyze the purpose of each plot type:

- **Bar graph:** Used to compare a numerical value across different categories. It does not show the relationship between two continuous variables.
- **Pie chart:** Used to show the proportion of categories as part of a whole. It works with categorical data.
- **Box plot:** Used to display the distribution of a single numerical variable, showing median, quartiles, and outliers.
- **Scatter plot:** This is the definitive plot for visualizing the relationship between two numerical variables. Each point on the plot represents an observation with coordinates corresponding to its values for the two variables. The pattern of the points (e.g., upward trend, downward trend, no pattern) reveals the correlation between the variables.

Because a scatter plot's primary purpose is to show the relationship (and thus, visually assess the correlation) between two variables, it is often called a correlation plot.

Step 3: Final Answer:

A Scatter plot is also known as a correlation plot. Therefore, option (D) is correct.

💡 Quick Tip

When you see "relationship between two variables" or "correlation", immediately think "scatter plot". It's the standard tool for this task in data visualization.

65. Given the following statement:

```
import matplotlib.pyplot as plt
```

'plt' in the above statement is _____ name.

- (A) key
- (B) alias
- (C) variable
- (D) function

Correct Answer: (B) alias

Solution:

Step 1: Understanding the Concept:

In Python, the `import ... as ...` syntax is used to import a module but give it a different, often shorter, name within the current script. This new name is known as an alias.

Step 2: Detailed Explanation:

In the statement `import matplotlib.pyplot as plt`:

- `import matplotlib.pyplot`: This part tells Python to import the `pyplot` submodule from the `matplotlib` library.
- `as plt`: This part creates an **alias**. It means that within this code, the long name `matplotlib.pyplot` can be referred to using the short name `plt`. For example, instead of writing `matplotlib.pyplot.plot()`, you can simply write `plt.plot()`.

The term "alias" is the correct technical term for the name created with the 'as' keyword. "Variable" is a more general term; while 'plt' does act like a variable that holds the module object, "alias" is the more specific and descriptive term for its role in the import statement. "Key" and "function" are incorrect.

Step 3: Final Answer:

'plt' is an alias for 'matplotlib.pyplot'. Therefore, option (B) is correct.

💡 Quick Tip

Using standard aliases like `import pandas as pd`, `import numpy as np`, and `import matplotlib.pyplot as plt` is a strong convention in the data science community. Following these conventions makes your code more readable to others.

66. How can you save a matplotlib plot as a file?

- (A) `plt.export()`
- (B) `plt.save()`
- (C) `plt.savefig()`
- (D) `plt.store()`

Correct Answer: (C) `plt.savefig()`

Solution:

Step 1: Understanding the Concept:

Matplotlib is a plotting library in Python. After creating a plot in memory, there is a specific function to save the current figure to an image file (like PNG, JPG, PDF, etc.) on the disk.

Step 2: Detailed Explanation:

Let's examine the function names:

- `plt.export()`: This function does not exist in matplotlib.
- `plt.save()`: This function does not exist in matplotlib for saving figures.
- `plt.savefig()`: This is the correct function. It stands for "save figure". You pass the desired filename as an argument, for example, `plt.savefig('my_plot.png')`. It has many optional arguments to control the resolution, size, and format of the output file.
- `plt.store()`: This function does not exist in matplotlib.

The standard and only correct function for this purpose among the options is `plt.savefig()`.

Step 3: Final Answer:

The correct command to save a matplotlib plot is `plt.savefig()`. Therefore, option (C) is correct.

💡 Quick Tip

A common mistake is to call `plt.show()` before `plt.savefig()`. In many environments, `plt.show()` clears the figure, so if you call `savefig()` after it, you might save a blank image. The correct order is to create the plot, then call `savefig()`, and finally call `show()` if you also want to display it.

67. Arrange the following in order in the context of exception handling:

- (A) Program Termination
- (B) Exception is raised
- (C) Error occurs in a program
- (D) Catching an exception

Choose the correct answer from the options given below:

- (A) (C), (B), (D), (A)
- (B) (A), (C), (B), (D)
- (C) (B), (A), (D), (C)
- (D) (C), (B), (A), (D)

Correct Answer: (A) (C), (B), (D), (A)

Solution:

Step 1: Understanding the Concept:

Exception handling is the process of responding to the occurrence of exceptions – anomalous or exceptional conditions requiring special processing – during the execution of a program.

This question asks for the logical sequence of events during this process.

Step 2: Detailed Explanation:

Let's trace the lifecycle of an exception:

1. It all starts when some piece of code encounters a problem it cannot handle, such as division by zero or a file not being found. This is the moment an **(C) Error occurs in a program**.
2. When the error occurs, the runtime system stops the normal flow of the program and creates an exception object. This process is called "throwing" or **(B) raising an exception**.
3. The runtime system then looks for a piece of code designed to handle this specific type of exception. If found (e.g., in a 'try...except' block), that code is executed. This is known as **(D) catching an exception**.
4. If the exception is caught and handled, the program can often continue or terminate gracefully. If the exception is not caught, the program will halt and typically display an error message. In either case, the process eventually leads to **(A) Program Termination** (either gracefully or by crashing).

The logical sequence of events is $C \rightarrow B \rightarrow D \rightarrow A$.

Step 3: Final Answer:

The correct order is (C) Error occurs, (B) Exception is raised, (D) Catching an exception, (A) Program Termination. This matches option (A).

 Quick Tip

Think of it like a game of catch. A problem occurs (C), the program "throws" a ball (B), a player (the 'except' block) "catches" it (D), and then the game ends (A).

68. The plot function of pandas matplotlib uses 'kind' argument which can accept a string indicating the type of graph to be plotted. Which of the following are valid plots?

- (A) Area
- (B) Box
- (C) Scatter
- (D) Line

Choose the correct answer from the options given below:

- (A) (A), (B) and (D) only
- (B) (A), (B) and (C) only
- (C) (A), (B), (C) and (D)
- (D) (B), (C) and (D) only

Correct Answer: (C) (A), (B), (C) and (D)

Solution:

Step 1: Understanding the Concept:

The `.plot()` method in pandas is a versatile wrapper around matplotlib's plotting functions. It uses a `kind` parameter to specify which type of plot to create. We need to identify which of the given plot types are valid values for this `kind` parameter.

Step 2: Detailed Explanation:

The pandas `DataFrame.plot()` method accepts numerous string values for the `kind` argument. Let's check the given options:

- **(A) Area:** `kind='area'` is a valid option for creating an area plot.
- **(B) Box:** `kind='box'` is a valid option for creating a box-and-whisker plot.
- **(C) Scatter:** `kind='scatter'` is a valid option for creating a scatter plot.
- **(D) Line:** `kind='line'` is the default option for creating a line plot.

All four options listed are valid and commonly used plot types that can be generated using the `kind` parameter of the pandas `.plot()` method. Other valid kinds include 'bar', 'barh', 'hist', 'kde', 'density', 'pie'.

Step 3: Final Answer:

All listed plots (Area, Box, Scatter, Line) are valid. Therefore, option (C) is correct.

 Quick Tip

The pandas `.plot()` method is a powerful shortcut. You can create most common plots directly from a `DataFrame` (`'df.plot(kind='...', x='...', y='...')`) without needing to import matplotlib directly for simple cases.

69. Given the following code for histogram:

```
df.plot(kind='hist', edgecolor='Green', linewidth=2, linestyle='--',
fill=False, hatch='o')
```

What is role of fill=False?

- (A) Each hist will be filled with green color.
- (B) Each hist will be empty.
- (C) Each hist will be filled with '—'.
- (D) The edge color of each hist will be black.

Correct Answer: (B) Each hist will be empty.

Solution:**Step 1: Understanding the Concept:**

In matplotlib (which pandas plotting is based on), many plot elements like bars, patches, and polygons have separate properties for their face (fill) color and their edge color. The question asks about a specific parameter that controls the fill of the bars in a histogram.

Step 2: Detailed Explanation:

Let's analyze the parameters in the given code:

- `kind='hist'`: Creates a histogram.
- `edgecolor='Green'`: Sets the color of the outline of each bar to green.
- `linewidth=2`: Makes the outline 2 points thick.
- `linestyle='--'`: Makes the outline a dashed line.
- `fill=False`: This parameter controls whether the interior of the patch (the bar) is filled with color. By default, it is `True`. Setting it to `False` means the interior of the bars will not be filled; they will be transparent or "empty", showing only the styled edges.
- `hatch='o'`: This would add a pattern of small circles inside the bars if they were filled. When `fill=False`, the effect of hatch might be minimal or not visible.

The specific role of `fill=False` is to make the bars of the histogram empty, leaving only the outline visible.

Step 3: Final Answer:

The role of `fill=False` is to ensure each histogram bar will be empty. Therefore, option (B) is correct.

💡 Quick Tip

To customize plots, remember the distinction between 'edgecolor' (the outline) and 'facecolor' or 'color' (the fill). Setting 'fill=False' is a useful trick to create "hollow" or outline-only plots.

70. Matplotlib library can be installed using pip command. Identify the correct syntax from the following options.

- (A) `pip install matplotlib.pyplotlib`
- (B) `pip install matplotlib.pyplotlib`
- (C) `pip install matplotlib.pyplotlib as plt`
- (D) `pip install matplotlib`

Correct Answer: (D) `pip install matplotlib`

Solution:

Step 1: Understanding the Concept:

`pip` is the standard package installer for Python. To install a library, you use the command `pip install` followed by the official package name as registered in the Python Package Index (PyPI).

Step 2: Detailed Explanation:

Let's analyze the syntax of the given commands:

- (A) `pip install matplotlib.pyplotlib`: This is incorrect. `matplotlib.pyplot` is a *module inside* the library, not the name of the package itself.
- (B) `pip install matplotlib.pyplotlib`: This is incorrect; "matplotlib" is a misspelling of the library name.
- (C) `pip install matplotlib.pyplotlib as plt`: This is incorrect. The `as plt` syntax is used for aliasing within a Python *import statement*, not in a `pip` command in the shell.
- (D) `pip install matplotlib`: This is the correct command. The official package name for the Matplotlib library is simply `matplotlib`.

Step 3: Final Answer:

The correct command to install the Matplotlib library is `pip install matplotlib`.

Therefore, option (D) is correct.

💡 Quick Tip

Remember to distinguish between the **package name** used for installation ('`pip install package-name`') and the **module name** used for importing ('`import module_name`'). They are often the same, but not always.

71. Dynamic web pages can be created using various languages, like _____.

- (A) JavaScript
- (B) PHP
- (C) Python
- (D) Ruby

Choose the correct answer from the options given below:

- (A) (A), (B) and (D) only
- (B) (A), (B) and (C) only
- (C) (A), (B), (C) and (D)
- (D) (B), (C) and (D) only

Correct Answer: (C) (A), (B), (C) and (D)

Solution:

Step 1: Understanding the Concept:

A dynamic web page is a web page whose content can be generated on the fly, often based on user interactions, time of day, or data from a database. This is in contrast to a static web page, where the content is fixed. This dynamic generation is handled by programming languages.

Step 2: Detailed Explanation:

Let's analyze the languages listed:

- **(A) JavaScript:** JavaScript is the primary language for creating dynamic behavior on the client-side (in the browser). With technologies like Node.js, it is also a very popular server-side language.
- **(B) PHP:** PHP (Hypertext Preprocessor) is one of the most widely used open-source, server-side scripting languages specifically designed for web development.
- **(C) Python:** Python, with frameworks like Django, Flask, and FastAPI, is a powerful and popular language for server-side web development.
- **(D) Ruby:** Ruby, particularly with the Ruby on Rails framework, is another prominent server-side language used for creating dynamic web applications.

All four languages listed are widely used to create dynamic web pages and web applications.

Step 3: Final Answer:

All the listed languages can be used to create dynamic web pages. Therefore, option (C) is the correct answer.

 Quick Tip

Remember the distinction: **Client-side** languages (mostly JavaScript) run in the user's browser to make the page interactive. **Server-side** languages (PHP, Python, Ruby, Java, etc.) run on the web server to generate the page content before it's sent to the browser.

72. Match List-I with List-II

List-I	List-II
(A) STAR TOPOLOGY	(I) Each communicating device is connected to a central node.
(B) LAN	(II) Networking device
(C) HUB	(III) Protocol
(D) VoIP	(IV) Type of network

Choose the correct answer from the options given below:

- (A) (A)-(I), (B)-(IV), (C)-(II), (D)-(III)
 (B) (A)-(I), (B)-(II), (C)-(III), (D)-(IV)
 (C) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)
 (D) (A)-(III), (B)-(IV), (C)-(II), (D)-(I)

Correct Answer: (A) (A)-(I), (B)-(IV), (C)-(II), (D)-(III)

Solution:

Step 1: Understanding the Concept:

This question requires matching various networking terms to their correct categories, such as topology, network type, device, or protocol.

Step 2: Detailed Explanation:

Let's match each item from List-I to its correct category/description in List-II.

- **(A) STAR TOPOLOGY:** This is a network layout where **(I) each communicating device is connected to a central node** (like a hub or a switch). This is the definition of a star topology.
- **(B) LAN:** LAN stands for Local Area Network. It is a classification based on the geographical scope of a network. It is a **(IV) Type of network**.
- **(C) HUB:** A hub is a simple, physical-layer device that connects multiple computers together in a network. It is a **(II) Networking device**.
- **(D) VoIP:** VoIP stands for Voice over Internet Protocol. It is a set of rules and technologies for delivering voice communications over IP networks. It is a **(III) Protocol**.

The correct set of matches is (A)-(I), (B)-(IV), (C)-(II), (D)-(III).

Step 3: Final Answer:

This matching corresponds exactly with option (A).

 Quick Tip

Again, knowing the categories is key: **Topology** is the layout (Star, Bus, Ring). **Type** is the scale (LAN, WAN, MAN, PAN). **Device** is the hardware (Hub, Switch, Router). **Protocol** is the set of rules (IP, TCP, HTTP, VoIP).

73. Arrange the following in sequence of execution:

- (A) HTTP Request**
- (B) Calls an application in response to the HTTP request**
- (C) Program executes and produces HTML output**
- (D) HTTP Response**

Choose the correct answer from the options given below:

- (A) (B), (A), (C), (D)
- (B) (C), (A), (B), (D)
- (C) (A), (B), (C), (D)
- (D) (C), (B), (A), (D)

Correct Answer: (C) (A), (B), (C), (D)

Solution:

Step 1: Understanding the Concept:

This question asks for the sequence of events that occur when a client (like a web browser) requests a dynamic web page from a web server. This is the fundamental request-response cycle of the web.

Step 2: Detailed Explanation:

Let's trace the sequence of execution:

1. The process begins when a user's web browser sends a request to a web server for a specific resource (e.g., a web page). This is the **(A) HTTP Request**.
2. The web server receives the request. If the request is for a dynamic page (e.g., a .php or .py file), the server doesn't just send a file back. Instead, it recognizes that it needs to run a program. It **(B) calls an application** (like the PHP interpreter or a Python web application) and passes the request details to it.
3. The application or script then **(C) executes**. This might involve querying a database, performing calculations, and ultimately generating the HTML content for the web page.
4. The generated HTML is sent back from the application to the web server, which then packages it into an **(D) HTTP Response** and sends it back across the internet to the user's browser.

The correct sequence is $A \rightarrow B \rightarrow C \rightarrow D$.

Step 3: Final Answer:

The correct sequence is (A), (B), (C), (D). This matches option (C).

 Quick Tip

Remember the web's basic conversation: Client asks (HTTP Request), Server thinks (Calls Application, Program Executes), Server answers (HTTP Response).

74. Match List-I with List-II

List-I	List-II
(A) Static web page	(I) A web page whose content keeps changing.
(B) Home page	(II) A web page whose content is fixed.
(C) Dynamic web page	(III) Delivers the contents to web browser.
(D) Web server	(IV) First page of website

Choose the correct answer from the options given below:

- (A) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)
 (B) (A)-(I), (B)-(II), (C)-(III), (D)-(IV)
 (C) (A)-(IV), (B)-(II), (C)-(I), (D)-(III)
 (D) (A)-(II), (B)-(I), (C)-(IV), (D)-(III)

Correct Answer: (A) (A)-(II), (B)-(IV), (C)-(I), (D)-(III)

Solution:

Step 1: Understanding the Concept:

This question tests basic definitions related to web pages and web servers.

Step 2: Detailed Explanation:

Let's match each term in List-I with its correct definition in List-II.

- **(A) Static web page:** A static web page is one where the content is delivered to the user exactly as it is stored on the server. Its content is **(II) fixed**.
- **(B) Home page:** This is the main page or the default starting page of a website. It is the **(IV) first page of a website** that a user typically sees.

- **(C) Dynamic web page:** A dynamic web page is generated by a server-side application, and its **(I) content keeps changing**, often based on user input or other variables.
- **(D) Web server:** This is the software (and the underlying hardware) that receives HTTP requests and **(III) delivers the contents (web pages) to the web browser**.

The correct set of matches is (A)-(II), (B)-(IV), (C)-(I), (D)-(III).

Step 3: Final Answer:

This matching corresponds exactly with option (A).

 Quick Tip

Remember the key difference: **Static** = Fixed, Stored. **Dynamic** = Changing, Generated.

75. Which of the following is a networking device?

- (A) Gateway
- (B) Repeater
- (C) MAN
- (D) Modem

Choose the correct answer from the options given below:

- (A) (A) and (D) only
- (B) (A), (B) and (D) only
- (C) (A), (B), (C) and (D)
- (D) (B), (C) and (D) only

Correct Answer: (B) (A), (B) and (D) only

Solution:

Step 1: Understanding the Concept:

The question asks to identify which of the given terms represent physical hardware devices used in computer networking.

Step 2: Detailed Explanation:

Let's analyze each option:

- **(A) Gateway:** A gateway is a network node that connects two different networks that use different protocols. It is a hardware device (often integrated into a router).
- **(B) Repeater:** A repeater is a simple electronic device that receives a signal and retransmits it at a higher power level, allowing the signal to cover longer distances. It is a network device.
- **(C) MAN:** MAN stands for Metropolitan Area Network. It is a *type or classification* of a network based on its geographical size, not a physical device.
- **(D) Modem:** A modem (Modulator-Demodulator) is a hardware device that converts digital data from a computer into an analog signal for transmission over phone lines and vice versa. It is a network device.

Gateway, Repeater, and Modem are all networking devices. MAN is a type of network.

Step 3: Final Answer:

The items that are networking devices are (A), (B), and (D). This corresponds to option (B).

💡 Quick Tip

Be careful to distinguish between physical devices (router, switch, modem, gateway, repeater, hub, NIC) and logical concepts or classifications (LAN, WAN, MAN, protocol, topology).

76. _____ is the largest WAN that connects billions of computers, smartphones and millions of LANs from different countries.

- (A) PAN
- (B) WWW
- (C) Ethernet
- (D) Internet

Correct Answer: (D) Internet

Solution:

Step 1: Understanding the Concept:

The question asks for the name of the global system of interconnected computer networks that spans the entire planet.

Step 2: Detailed Explanation:

Let's define the terms:

- **(A) PAN (Personal Area Network):** A very small network, typically for connecting devices around a single person.
- **(B) WWW (World Wide Web):** This is a system of interlinked hypertext documents accessed *via the Internet*. The Web is a service that runs on top of the Internet, but it is not the network itself.
- **(C) Ethernet:** This is a family of technologies for building Local Area Networks (LANs). It is not a global WAN.
- **(D) Internet:** This is the global network of networks (a WAN of WANs) that connects billions of devices worldwide using the TCP/IP protocol suite. The description in the question is the definition of the Internet.

Step 3: Final Answer:

The largest WAN that connects billions of computers globally is the Internet. Therefore, option (D) is correct.

💡 Quick Tip

A common point of confusion is the difference between the Internet and the World Wide Web (WWW). Remember: the **Internet** is the physical network infrastructure (the roads), while the **Web** is one of the services that uses that infrastructure (the cars and trucks on the roads).

77. The intellectual property right that is granted for inventions is _____.

- (A) Copyright
- (B) Patent
- (C) Trademark
- (D) Licensing

Correct Answer: (B) Patent

Solution:

Step 1: Understanding the Concept:

Intellectual Property (IP) refers to creations of the mind. Different types of IP are protected by different legal instruments. The question asks for the specific right granted for inventions.

Step 2: Detailed Explanation:

Let's define the different types of IP rights:

- **Copyright:** Protects original works of authorship, such as books, music, art, and software code. It protects the *expression* of an idea, not the idea itself.
- **Patent:** Protects new, useful, and non-obvious **inventions**, such as processes, machines, and chemical compositions. It grants the inventor the exclusive right to make, use, and sell the invention for a limited period.
- **Trademark:** Protects words, names, symbols, sounds, or colors that distinguish goods and services from those manufactured or sold by others and to indicate the source of the goods. Examples include brand names and logos.
- **Licensing:** This is not a type of IP right itself, but rather a legal agreement that allows someone else to use an existing IP right (like a patent or copyright) in exchange for payment (royalties).

The right specifically granted for inventions is a patent.

Step 3: Final Answer:

A patent is the intellectual property right granted for inventions. Therefore, option (B) is correct.

💡 Quick Tip

Remember the core concepts: **Copyright** for creative works (art, music, books). **Patent** for inventions (machines, processes). **Trademark** for brands (logos, names).

78. Violation of intellectual property right may happen due to:

- (A) Copyright Infringement**
- (B) Patent**
- (C) Trademark Infringement**
- (D) Plagiarism**

Choose the correct answer from the options given below:

- (A) (A) and (C) only
- (B) (A), (B) and (C) only
- (C) (A), (B), (C) and (D)
- (D) (A), (C) and (D) only

Correct Answer: (D) (A), (C) and (D) only

Solution:

Step 1: Understanding the Concept:

The question asks for the causes or forms of violation of intellectual property rights. A violation occurs when someone uses a protected intellectual property without the owner's permission.

Step 2: Detailed Explanation:

Let's analyze the options:

- **(A) Copyright Infringement:** This is the unauthorized use of works protected by copyright law, such as copying a book or music without permission. This is a direct violation of IP rights.
- **(B) Patent:** A patent is a *form of protection*, not a violation. Patent *infringement* would be a violation, but "Patent" itself is not.
- **(C) Trademark Infringement:** This is the unauthorized use of a trademark in a manner that is likely to cause confusion about the source of goods or services. This is a direct violation of IP rights.
- **(D) Plagiarism:** This is the act of presenting someone else's work or ideas as one's own, without giving credit. While there is an overlap with copyright infringement, plagiarism is fundamentally an ethical offense. In many contexts, especially legal ones concerning IP, it is considered a form of violation. For instance, copying text from a copyrighted book without attribution is both plagiarism and copyright infringement.

The clear violations are Copyright Infringement (A) and Trademark Infringement (C). Plagiarism (D) is also a form of violation, representing the unethical use of another's intellectual work. Patent (B) is the legal right itself, not a violation. Therefore, A, C, and D are all forms or causes of IP violation.

Step 3: Final Answer:

The violations are (A) Copyright Infringement, (C) Trademark Infringement, and (D) Plagiarism. This corresponds to option (D).

 Quick Tip

Remember to distinguish between the right itself (e.g., patent, copyright) and the violation of that right (infringement). Infringement is the key term for a legal violation.

79. Match List-I with List-II

List-I	List-II
(A) Cyber Appellate Tribunal	(I) Punishes people causing pollution.
(B) Environmental Protection Act 1986	(II) Provides guidelines for proper handling
(C) Indian Information Technology Act 2000	(III) Resolves disputes arising from cyber c
(D) Central Pollution Control Board	(IV) Provides guidelines on storage, proces

Choose the correct answer from the options given below:

- (A) (A)-(III), (B)-(I), (C)-(IV), (D)-(II)
(B) (A)-(I), (B)-(III), (C)-(II), (D)-(IV)
(C) (A)-(II), (B)-(III), (C)-(IV), (D)-(I)
(D) (A)-(III), (B)-(IV), (C)-(II), (D)-(I)

Correct Answer: (A) (A)-(III), (B)-(I), (C)-(IV), (D)-(II)

Solution:

Step 1: Understanding the Concept:

This question requires matching various Indian legal acts and bodies with their respective functions, particularly in the context of technology, environment, and cyber law.

Step 2: Detailed Explanation:

Let's match each item from List-I to its correct description in List-II.

- **(A) Cyber Appellate Tribunal:** Established under the IT Act, its primary function is to hear appeals against the orders of adjudicating officers. It **(III) resolves disputes arising from cyber crime** and contraventions of the IT Act.
- **(B) Environmental Protection Act 1986:** This is a comprehensive act to provide for the protection and improvement of the environment. It provides the framework that **(I) punishes people causing pollution** and environmental damage.
- **(C) Indian Information Technology Act 2000:** This act provides the legal framework for electronic transactions and cybercrime. Its rules, particularly the "Sensitive Personal Data or Information" rules of 2011, **(IV) provide guidelines on storage, processing and transmission of sensitive data.**
- **(D) Central Pollution Control Board (CPCB):** The CPCB is a statutory organisation. A key part of its mandate in recent years has been to manage electronic waste. It **(II) provides guidelines for proper handling and disposal of e-waste.**

The correct set of matches is (A)-(III), (B)-(I), (C)-(IV), (D)-(II).

Step 3: Final Answer:

This matching corresponds exactly with option (A).

 Quick Tip

When faced with questions about legal acts, look for keywords. "Cyber" relates to the IT Act. "Pollution" and "Environmental" relate to the EPA and CPCB. This can help you quickly narrow down the options.

80. _____ is a branch of science that deals with designing and arranging of workplaces.

- (A) Netiquette
- (B) Ergonomics
- (C) Leeching
- (D) Refurbishing

Correct Answer: (B) Ergonomics

Solution:

Step 1: Understanding the Concept:

The question asks for the scientific discipline concerned with the interaction between humans and other elements of a system, particularly in the design of workplaces and products to optimize human well-being and overall system performance.

Step 2: Detailed Explanation:

Let's define the terms:

- **Netiquette:** "Network etiquette," the set of social conventions that facilitate interaction over networks.
- **Ergonomics:** This is the science of designing and arranging things people use so that the people and things interact most efficiently and safely. It is precisely about designing workplaces, equipment, and tasks to fit the user.
- **Leeching:** In internet slang, this refers to the practice of downloading much more than one uploads on a peer-to-peer network, or generally taking without giving back.
- **Refurbishing:** The process of restoring an old or used product to a like-new condition.

The definition provided in the question perfectly matches the definition of Ergonomics.

Step 3: Final Answer:

Ergonomics is the science of designing workplaces. Therefore, option (B) is correct.

 Quick Tip

Associate "Ergonomics" with "ergonomic chairs" or "ergonomic keyboards". These are products designed to be comfortable and efficient for human use, which is the core idea of the field.

81. Sequence the steps to append data to an existing file and then read the entire file:

- (A) Open the file in a+ mode.
- (B) Use the read() method to output the contents.
- (C) Use the write() or writelines() method to append data.
- (D) Seek to the beginning of the file.

Choose the correct answer from the options given below:

- (A) (A), (C), (D), (B)

- (B) (A), (B), (C), (D)
(C) (B), (A), (D), (C)
(D) (C), (B), (D), (A)

Correct Answer: (A) (A), (C), (D), (B)

Solution:

Step 1: Understanding the Concept:

The task is to append data to a file and then read the file's entire contents from the beginning. This requires understanding how file modes and file pointers (cursors) work in file I/O operations.

Step 2: Detailed Explanation:

Let's establish the logical sequence of operations:

1. **Open the file:** To both append and read, you need a mode that supports both operations. The 'a+' mode is designed for this. It opens the file for appending (placing the pointer at the end) and also allows reading. So, the first step is **(A) Open the file in a+ mode.**
2. **Append data:** After opening in append mode, the file pointer is at the end of the file. You can now write new data, which will be added to the end. This is step **(C) Use the write() or writelines() method to append data.**
3. **Reposition the pointer:** After writing, the pointer is still at the very end of the file. If you try to read now, you will read nothing (end-of-file). To read the *entire* file, you must first move the pointer back to the start. This is done using a seek operation. So, the next step is **(D) Seek to the beginning of the file.** (e.g., 'file.seek(0)').
4. **Read the data:** With the pointer now at the beginning, you can read the entire content of the file (including the newly appended data). This is step **(B) Use the read() method to output the contents.**

The correct sequence is $A \rightarrow C \rightarrow D \rightarrow B$.

Step 3: Final Answer:

The correct sequence is (A), (C), (D), (B). This corresponds to option (A).

 Quick Tip

Remember the file pointer! When you open a file in append mode ('a' or 'a+'), the pointer starts at the end. To read from the beginning after appending, you must explicitly 'seek(0)' to rewind the pointer.

82. Sequence the given process for checking whether a string is a palindrome or not, using a deque:

- (A) Load the string into the deque.
(B) Continuously remove characters from both ends.
(C) Compare the characters.
(D) Determine if the string is a palindrome based on the comparisons.

Choose the correct answer from the options given below:

- (A) (A), (B), (D), (C)
- (B) (A), (B), (C), (D)
- (C) (B), (A), (D), (C)
- (D) (C), (B), (D), (A)

Correct Answer: (B) (A), (B), (C), (D)

Solution:

Step 1: Understanding the Concept:

A palindrome is a word or phrase that reads the same forwards and backward (e.g., "racecar"). A deque (double-ended queue) is a data structure that allows adding and removing elements from both ends efficiently. It is well-suited for checking palindromes. The algorithm involves repeatedly comparing the first and last characters.

Step 2: Detailed Explanation:

Let's lay out the logical steps for a deque-based palindrome checker:

1. **Initialization:** First, the string to be checked needs to be placed into the data structure. So, step one is **(A) Load the string into the deque**.
2. **Looping and Removal:** The core of the algorithm is a loop that continues as long as there is more than one character in the deque. In each iteration, we need to get the characters from both ends. This is described by **(B) Continuously remove characters from both ends** (e.g., using 'popleft()' and 'pop()').
3. **Comparison:** Inside the loop, immediately after removing the two characters, you must **(C) Compare the characters**. If they are not equal at any point, the string is not a palindrome, and the process can stop.
4. **Final Determination:** The loop continues until it's finished (the deque is empty or has one character left). Based on whether all comparisons passed, you can **(D) Determine if the string is a palindrome**.

The process flows logically as $A \rightarrow B \rightarrow C \rightarrow D$. The removal (B) and comparison (C) happen together inside a loop, and the final determination (D) is the outcome of that loop.

Step 3: Final Answer:

The correct sequence is (A), (B), (C), (D). This matches option (B).

Quick Tip

A deque is ideal for palindrome checking because it provides $O(1)$ time complexity for removing items from both the front and the back, making the overall algorithm very efficient.

83. Cable based broadband internet services are example of _____.

- (A) LAN
- (B) MAN
- (C) WAN

(D) HAN

Correct Answer: (B) MAN

Solution:

Step 1: Understanding the Concept:

The question asks to classify the type of network that provides broadband internet to homes and businesses across a city using cable infrastructure. We need to select the appropriate network scale classification.

Step 2: Detailed Explanation:

Let's define the network types by their geographical scale:

- **LAN (Local Area Network):** Covers a small area like a single home, office, or building.
- **MAN (Metropolitan Area Network):** Covers a larger geographical area like a city or a large campus. It interconnects a number of LANs. Cable TV networks and their associated broadband services are a classic example of a MAN, as they span an entire city.
- **WAN (Wide Area Network):** Covers a very large area, such as a country or the entire globe (the Internet is the largest WAN).
- **HAN (Home Area Network):** This is essentially another term for a home LAN.

A cable broadband service that provides internet to thousands of subscribers across a city fits the description of a Metropolitan Area Network.

Step 3: Final Answer:

Cable broadband services are an example of a MAN. Therefore, option (B) is correct.

 Quick Tip

Think of the scale: **Home** (HAN/LAN) ; **City** (MAN) ; **Country/World** (WAN).
This hierarchy helps classify different network examples.

84. _____ is the trail of data we leave behind when we visit any website (or use any online application or portal) to fill-in data or perform any transaction.

- (A) Digital footprint
- (B) Cyber crimes
- (C) Cyber bullying
- (D) Identity theft

Correct Answer: (A) Digital footprint

Solution:

Step 1: Understanding the Concept:

The question asks for the term that describes the record or trail of a person's online activity.

Step 2: Detailed Explanation:

Let's define the terms:

- **Digital footprint:** This refers to the unique set of traceable digital activities, actions, communications, or transactions on the Internet or digital devices. It includes both active footprints (data you intentionally share, like social media posts) and passive footprints (data collected without your awareness, like your IP address). The description in the question is the exact definition of a digital footprint.
- **Cyber crimes:** This is a broad term for criminal activities carried out by means of computers or the Internet.
- **Cyber bullying:** This is the use of electronic communication to bully a person, typically by sending messages of an intimidating or threatening nature.
- **Identity theft:** This is the fraudulent acquisition and use of a person's private identifying information, usually for financial gain.

The term that specifically means the "trail of data" left behind online is digital footprint.

Step 3: Final Answer:

The correct term is Digital footprint. Therefore, option (A) is correct.

 Quick Tip

Think of a "footprint" in the real world—it's a mark you leave behind where you've been. A "digital footprint" is the same concept, but for your activities in the digital world.