



black N White



NAME	
ROLL NUMBER	
SEMESTER	2 nd
COURSE CODE	DCA1206
COURSE NAME	BCA
Subject Name	BASIC STATISTICS AND PROBABILITY

SET - I

Q.1.a) Define statistics and discuss its scope across different fields with examples.

Answer :-

Statistics is the branch of mathematics that deals with the collection, organization, analysis, interpretation, and presentation of data. It helps in making informed decisions based on numerical evidence, especially when dealing with uncertain or variable information.

Scope of Statistics in Different Fields:

Statistics is widely applied across various disciplines:

1. Business and Economics:

In business, statistics helps in market analysis, demand forecasting, quality control, and financial planning. For example, companies use statistical tools to analyze customer behavior and improve marketing strategies.

2. Health and Medicine:

In healthcare, statistics is essential for medical research, clinical trials, and analyzing patient data. For example, statistical analysis is used to determine the effectiveness of new drugs.

3. Education:

Educators use statistics to assess student performance, evaluate teaching methods, and analyze enrollment trends. For instance, exam scores are statistically analyzed to determine pass rates and overall performance.

4. Government and Public Policy:

Governments use statistics for census data, unemployment rates, economic planning, and policy formulation. For example, statistical data helps in planning budgets and social welfare programs.

5. Social Sciences:

In fields like psychology, sociology, and political science, statistics is used to interpret surveys, conduct experiments, and study human behavior.

Q1.b) Explain the process and importance of data classification, including the distinction between attributes and variables.

Answer :- Data classification is the process of organizing data into categories for effective analysis and interpretation. It involves grouping similar types of data together to simplify complex information, making it easier to understand, compare, and draw conclusions.

Process of Data Classification:

- 1. Collection of Data:** Gather raw data through surveys, experiments, or records.
- 2. Editing and Cleaning:** Remove errors, incomplete or irrelevant data.
- 3. Categorization:** Group data based on common characteristics.
- 4. Tabulation:** Arrange classified data in tables for easier analysis.

Importance of Data Classification:

- **Simplifies complex data:** Makes large volumes of data manageable.
- **Improves accuracy:** Helps in identifying patterns and trends.
- **Enhances decision-making:** Organized data supports better analysis and predictions.
- **Saves time:** Makes data retrieval and interpretation faster.

Attributes vs. Variables:

- **Attributes:** These are qualitative characteristics that cannot be measured numerically. Example: Gender, Religion, Marital Status.
- **Variables:** These are quantitative characteristics that can be measured. Example: Age, Income, Height.

Data classification is essential for transforming raw data into meaningful information, and understanding the difference between attributes and variables helps in choosing the right method of analysis.

Q1.c) Describe the concept and advantages of frequency distribution in summarizing large datasets

Answer :-

Frequency distribution is a method used to organize large sets of data into a structured form by displaying the number of times each value or range of values occurs. It helps in converting raw data into an understandable format by grouping data into classes or intervals and counting how often each class appears.

Advantages:

1. **Simplifies Data:** It makes large and unorganized data easier to understand by grouping similar values together.
 2. **Identifies Patterns:** Helps to recognize trends, patterns, and central tendencies within the data.
 3. **Facilitates Comparison:** Enables comparison between different data sets or categories.
 4. **Supports Visualization:** Forms the basis for charts like histograms and bar graphs, aiding visual analysis.
 5. **Improves Decision Making:** Well-organized data helps analysts make informed decisions more effectively.
- Frequency distribution is a vital tool in statistics that helps in summarizing and analyzing data concisely.

Q.2.a) What is central tendency? Explain its purpose and the characteristics of a good measure.

Answer :- Central tendency refers to a statistical measure that identifies the center or average value of a dataset. The most common measures of central tendency are **mean**, **median**, and **mode**. These values give a single representative figure around which all other data points tend to cluster.

Purpose:

The main purpose of central tendency is to summarize a large set of data with a single value that represents the entire distribution. It helps in understanding the general behavior of the data and allows for easy comparison between different datasets.

Characteristics of a Good Measure of Central Tendency:

1. **Simple and Easy to Understand:** It should be easy to compute and interpret.
2. **Based on All Observations:** A good measure considers every value in the dataset (e.g., mean).
3. **Not Affected by Extreme Values:** It should remain stable in the presence of outliers (e.g., median).

4. **Mathematically Useful:** It should allow for further statistical analysis.
5. **Consistent and Reliable:** It should give consistent results across similar datasets.
- Central tendency provides a quick overview of data and helps in effective decision-making and analysis.

Q.2.b) Illustrate the calculation of median and mode for grouped and ungrouped data with suitable examples.

Answer :-

1. Ungrouped Data (Simple list of numbers)

Median (Ungrouped):

When you have a simple list of numbers:

- First, arrange the numbers in ascending (small to large) order.
- If the number of items is **odd**, the middle number is the **median**.
- If the number of items is **even**, take the average of the two middle numbers.

Example:

Data: 10, 12, 15, 18, 20

This list has 5 numbers (odd), so the middle one is **15** — that's the **median**.

Mode (Ungrouped):

The **mode** is the number that appears **most frequently** in the data.

Example:

Data: 2, 4, 4, 5, 6

Here, **4** appears twice, so the **mode** is 4.

2. Grouped Data (Class Intervals)

Median (Grouped):

When data is given in class intervals like 0–10, 10–20, etc.:

- Add up all the frequencies to get the total.
- Divide total by 2 to get half.
- Find the class interval where this halfway value falls — this is the **median class**.
- Then use a formula to calculate the **median**, based on the class limit, frequency, and class width.

Example:

Classes: 0–10, 10–20, 20–30, 30–40

Frequencies: 5, 8, 12, 10

Total = 35 → Half = 17.5

The 17.5th value falls in class 20–30 → this is the **median class**.

Then we apply a formula to get the exact median value.

Mode (Grouped):

- Look for the class interval with the **highest frequency** — this is the **modal class**.
- Then use a formula to calculate the mode using the modal class and the surrounding frequencies.

The **mode** gives an estimate of the value that occurs most frequently in the grouped data.

Q.2.c) Differentiate between arithmetic mean and weighted mean, and discuss their applications.

Answer :- Arithmetic Mean:

Arithmetic mean is the simple average of a set of numbers. It is calculated by adding all the values and dividing by the number of values.

Formula:

Arithmetic Mean = (Sum of all values) / (Number of values)

Example:

If a student scores 60, 70, and 80 in three subjects:

$$\text{Mean} = (60 + 70 + 80) / 3 = 70$$

Weighted Mean:

Weighted mean is used when different values have different levels of importance or weights. Each value is multiplied by its weight, and the total is divided by the sum of weights.

Formula:

$$\text{Weighted Mean} = (\sum wx) / (\sum w),$$

where w is the weight and x is the value.

Example:

Marks in subjects: 60 (weight 2), 70 (weight 3), 80 (weight 5)

$$\text{Weighted Mean} = (60 \times 2 + 70 \times 3 + 80 \times 5) / (2 + 3 + 5) = (120 + 210 + 400) / 10 = 73$$

Applications:

- **Arithmetic Mean:** Used in general situations like calculating average marks, temperatures, or incomes.
- **Weighted Mean:** Used when some items are more significant, like calculating GPA, stock indexes, or economic indicators.

Q.3.a) What is dispersion? Discuss its significance and the difference between absolute and relative measures.**Answer :-**

Dispersion refers to the degree to which data values spread out or vary from the central value (like the mean or median). It shows how much the values in a dataset differ from each other and helps understand the consistency or variability within the data.

Significance of Dispersion:

1. **Understanding Variability:** It tells how scattered the data is, which is important for analyzing reliability.
2. **Comparing Data Sets:** Helps to compare two or more datasets even if they have the same average.
3. **Improves Decision-Making:** In business, economics, and research, knowing data variability helps in better planning and predictions.
4. **Identifying Risks:** In finance, higher dispersion means higher risk.

Difference Between Absolute and Relative Measures:

Aspect	Absolute Measure	Relative Measure
Definition	Shows dispersion in original units (e.g., kg, ₹)	Shows dispersion as a ratio or percentage
Examples	Range, Variance, Standard Deviation	Coefficient of Variation, Relative Range
Use	Used when units are the same	Used to compare datasets with different units

Conclusion:

Dispersion is a key concept in statistics that reveals the spread of data. Absolute measures give exact values, while relative measures are useful for comparison.

Q.3.b.) Explain how to calculate range, variance, and standard deviation using an example each.

Answer :-

1. Range

Range is the simplest measure of dispersion. It tells you how spread out the data is by showing the difference between the largest and the smallest value in a dataset.

Example:

If your data is 10, 20, 30, 40, and 50 —

The smallest value is 10 and the largest is 50.

So, the range tells us the data is spread over a gap of 40.

Use: It quickly shows how wide the data is but doesn't tell much about how the values are spread in between.

2. Variance

Variance shows how far each number in the dataset is from the average value. It tells us whether the numbers are close to the mean or spread out over a wide range.

Example:

If all values are very close to the average, the variance will be small. But if values are very different from the average, the variance will be large.

Use: It helps in understanding the consistency or variation in data.

3. Standard Deviation

Standard deviation is closely related to variance. It is simply the average distance of each value from the mean. But unlike variance, it is in the same units as the original data, so it is easier to understand.

Example:

If a student's test scores are close to each other, the standard deviation will be low, meaning the performance is consistent. If the scores are very different, the standard deviation will be high, showing inconsistency.

In short:

- **Range** shows the total spread.
- **Variance** shows how values differ from the average.
- **Standard Deviation** tells how much values differ in a more understandable way.

Q.3.c) Describe one practical scenario each where standard deviation and relative variance are important in decision-making.

Answer :-

Standard Deviation - Practical Scenario:

Investment Risk Assessment:

An investor wants to choose between two stocks. Both have the same average return, but the investor needs to know how risky each one is. By calculating the **standard deviation** of their past returns, the investor can see which stock's returns vary more. A higher standard deviation means the stock price is more volatile and riskier. This helps the investor decide based on their risk tolerance.

**Relative Variance (Coefficient of Variation) - Practical Scenario:
Comparing Production Processes:**

A manufacturing company runs two production lines making different products with different average production times. To decide which line is more consistent, the company uses **relative variance** (coefficient of variation), which shows the variation as a percentage of the average time. This helps compare variability fairly despite differences in average production times and improve overall efficiency.

SET - II

Q.4.a) Define the classical, empirical, and subjective approaches to probability. Explain their significance in analyzing uncertainty.

Answer :-

1. Classical Approach to Probability

This approach defines probability based on equally likely outcomes. If an experiment has a fixed number of possible outcomes, and all outcomes are equally likely, the probability of an event is the ratio of favorable outcomes to total outcomes.

Example:

Rolling a fair six-sided die — the probability of getting a 3 is 1 out of 6.

Significance:

It's useful when dealing with simple, well-defined experiments like games of chance, where outcomes are symmetric and known.

2. Empirical (or Statistical) Approach to Probability

This approach calculates probability based on observed data or past experiments. The probability of an event is estimated by the relative frequency of that event occurring in repeated trials.

Example:

If it rained 30 days out of 100 days observed, the probability of rain on any day is estimated as $30/100 = 0.3$.

Significance:

It helps analyze real-world situations where outcomes are uncertain but historical data is available, such as weather forecasting or quality control.

3. Subjective Approach to Probability

This approach defines probability as a personal belief or judgment about how likely an event is, based on available information, intuition, or experience rather than exact calculations.

Example:

A doctor estimates the chance of a patient recovering based on medical knowledge and experience.

Significance:

It is useful when no exact data is available, allowing decision-makers to include expert opinion or personal beliefs to handle uncertainty.

Overall Significance in Analyzing Uncertainty:

These three approaches provide different ways to measure and understand uncertainty in various contexts.

- The **classical approach** works well with well-defined random experiments.
- The **empirical approach** relies on data and is practical in many real-life situations.
- The **subjective approach** helps when data is missing or incomplete but decisions still need to be made.

Q.4.b) What is a sample space? Distinguish between finite, infinite, discrete, and continuous sample spaces with examples.

Answer :- A **sample space** is the set of all possible outcomes of a random experiment. It represents every result that can occur when an experiment is performed.

There are different types of sample spaces:

- **Finite Sample Space:** Contains a limited number of outcomes.
Example: Tossing a coin results in {Heads, Tails}.
- **Infinite Sample Space:** Has unlimited outcomes.
Example: Counting the number of times a die is rolled until a 6 appears (can go on indefinitely).
- **Discrete Sample Space:** Consists of distinct, separate outcomes that can be counted.
Example: Number of students present in a class {0, 1, 2, ...}.
- **Continuous Sample Space:** Includes outcomes that take any value within a range, often measured.
Example: Measuring the height of students, which can be any value in a range like 150.5 cm, 150.51 cm, etc.

Understanding sample spaces helps in calculating probabilities by clearly defining all possible outcomes of an experiment.

Q.4.c) Define an event and discuss its role in probability theory, using a real-life situation.

Answer :- An **event** in probability theory is any specific outcome or a set of outcomes from a random experiment. It represents something that may or may not happen when the experiment is performed.

Role in Probability Theory:

Events are central to probability because probability measures how likely an event is to occur. By defining events clearly, we can assign probabilities to them and make predictions or decisions based on those chances.

Real-life Example:

Consider the event “It will rain tomorrow.” The weather forecast is based on this event’s probability. If the chance of rain is high, people may carry umbrellas or cancel outdoor plans. Here, the event helps individuals prepare for uncertainty by understanding the likelihood of rain.

Events allow us to focus on specific outcomes and evaluate their chances, making probability a practical tool for managing uncertainty in everyday life.

Q.5.a.) Differentiate deterministic, non-deterministic, and hybrid experiments with examples. How are these experiments used in probability analysis?

Answer :-

Deterministic experiments are those where the outcome is always certain and predictable. Given the initial conditions, the result is fixed and no randomness is involved. For example, turning on a light switch in a properly wired circuit will always result in the light turning on. There is no uncertainty or variation in the outcome.

Non-deterministic experiments, on the other hand, involve randomness or uncertainty, meaning the outcome cannot be predicted exactly. Each time the experiment is conducted, the result may vary. Tossing a fair coin is a classic example—there are two possible outcomes (heads or tails), and which one occurs is uncertain.

Hybrid experiments combine both deterministic and non-deterministic elements. Part of the process has a fixed outcome, while another part involves randomness. For example, a factory machine (deterministic part) produces items, but the number of defective items in each batch (non-deterministic part) varies randomly.

Use in Probability Analysis

In probability theory, **non-deterministic and hybrid experiments** are the primary focus since they involve uncertainty. Probability helps quantify the likelihood of various outcomes in these experiments. Deterministic experiments, having no uncertainty, generally don't require probability analysis. Hybrid experiments are analyzed by separating the deterministic and random components to assess the overall behavior under uncertainty.

Q.5.b) Explain the concept of expected value (EV). How is EV used in evaluating decision-making scenarios? Provide an example.

Answer :-

Concept of Expected Value (EV)

Expected Value (EV) is a statistical measure that represents the average outcome of a random experiment if it were repeated many times. It is calculated by multiplying each possible outcome by its probability and then adding all these products. Essentially, EV tells us what to expect on average in the long run.

Use of EV in Decision-Making

EV helps in making informed decisions under uncertainty by quantifying the average benefit or cost of different choices. Decision-makers use EV to compare alternatives and select the option with the best expected outcome, balancing risks and rewards.

Example

Suppose you enter a game where you can win \$100 with a 20% chance or win nothing with an 80% chance. The EV of playing the game is:

$$\bullet \quad EV = (0.2 \times \$100) + (0.8 \times \$0) = \$20 + \$0 = \$20$$

This means that on average, you can expect to win \$20 per game over many plays. If the cost to play is less than \$20, it might be a good decision to play; otherwise, it might be better to avoid the game.

EV provides a powerful way to evaluate uncertain outcomes and make rational decisions based on expected results.

Q.5.c) What are equally likely and exhaustive events? Illustrate how they influence probability calculations.

Answer :-

Equally Likely Events

Equally likely events are outcomes in a sample space that have the same chance of occurring. For example, when you roll a fair six-sided die, each face (1 to 6) has an equal probability of $1/6$. Because all outcomes are equally likely, calculating probability becomes straightforward: it's just the number of favorable outcomes divided by the total number of outcomes.

Exhaustive Events

Exhaustive events are a set of outcomes that include all possible results of an experiment. Together, they cover the entire sample space. For example, when tossing a coin, the events "Heads" and "Tails" are exhaustive because they include every possible outcome.

Influence on Probability Calculations

- For **equally likely events**, probability is simple to find because each outcome has the same chance, allowing use of the formula:
$$\text{Probability of event} = (\text{Number of favorable outcomes}) / (\text{Total outcomes})$$
- When events are **exhaustive**, their total probabilities add up to 1, meaning something from that set will definitely happen. This helps in checking if all possible outcomes are considered and ensures probabilities are correctly assigned.

Equally likely and exhaustive events simplify probability calculations and ensure complete coverage of possible outcomes in an experiment.