

DUMPS ARENA

VPC2 Data-Driven Decision Making C207

WGU Data-Driven-Decision-Making

Version Demo

Total Demo Questions: 14

Total Premium Questions: 141

Buy Premium PDF

<https://dumpsarena.co>

sales@dumpsarena.co

sales@dumpsarena.co
dumpsarena.co

QUESTION NO: 1

What is the formula for calculating a simple index number?

- A. The current price divided by the base year's price times 100
- B. The base year's price divided by the current price times 100
- C. The base year's price divided by the current price
- D. The current price divided by the base year's price

ANSWER: A

Explanation:

The current price divided by the base year's price times 100 is the standard formula for a simple price index: $(\text{current-period price} \div \text{base-period price}) \times 100$. This calculation expresses the current value relative to a selected reference period, conventionally assigning the base period an index value of 100. For example, if an item cost \$50 in the base year and \$60 currently, its simple index is $(60 \div 50) \times 100 = 120$. The resulting index of 120 means the current price is 120% of its base-year price, or 20% higher than the base-period level. This standardized form supports clear comparisons over time even when the original values use different units or scales. Simple index numbers can be used for prices, quantities, sales values, and other individual measures; the numerator is the period being compared and the denominator is the reference period. Index-number methods and their use in measuring changes relative to a base period are described by [Encyclopaedia Britannica](#). The U.S. Bureau of Labor Statistics also explains how price indexes measure average price change relative to a reference period in its [Consumer Price Index questions and answers](#).

QUESTION NO: 2

What describes fact-based decision-making according to quality management principles?

Choose 2 answers.

- A. Decisions foster trust in plans.
- B. Decisions are based on the instincts of experienced leaders.
- C. Decisions have a marginal effect on relationships with suppliers.
- D. Decisions reduce external bias.

ANSWER: A D

Explanation:

Fact-based decision-making, also called evidence-based decision making in ISO quality-management terminology, means making choices from the analysis and evaluation of reliable data and information. This approach makes plans more credible because stakeholders can understand the objective evidence supporting a decision. Accordingly, "Decisions foster trust in plans." accurately describes an important outcome of using facts, measurements, and analysis in management decisions. Data-supported decisions provide a transparent basis for selecting priorities, allocating resources, monitoring results, and making improvements.

"Decisions reduce external bias." is also correct because using verified, relevant information helps limit the influence of unsupported opinions, assumptions, preferences, and external pressure. The principle does not claim that data eliminates all judgment; rather, it requires leaders to use sound analysis alongside their experience and organizational context. ISO identifies evidence-based decision making as a quality-management principle and notes that decisions based on analysis and evaluation of data and information are more likely to produce desired results. The American Society for Quality likewise describes fact-based decision making as using factual information and data analysis to support effective quality decisions. See [ISO Quality Management Principles](#) and [ASQ Quality Management Principles](#).

QUESTION NO: 3

An analyst combines customer records from an e-commerce system and a customer-service system. Before calculating retention rates, the analyst finds that some customers appear more than once and that dates are stored in different formats.

Which two actions should the analyst take before performing the retention analysis?

Choose 2 answers.

- A. Identify and resolve duplicate customer records.
- B. Standardize dates into a consistent format.
- C. Delete all records containing unusual values.
- D. Average the values in records that appear duplicated.

ANSWER: A B

Explanation:

The analyst should identify and resolve duplicate customer records so that a single customer is not counted multiple times in the retention calculation. The analyst should also standardize date formats so that transactions and service interactions can be placed in the correct time period and compared consistently.

Deleting all records with unusual values could remove valid customer activity, and averaging the duplicate records does not establish which records belong to the same customer. These data-quality issues should be documented and addressed before the analysis is used for decision-making.

QUESTION NO: 4

Which two tools make it easier to detect an out-of-range error?

Choose 2 answers.

- A. Relational databases
- B. Experimental studies
- C. Spreadsheets
- D. Observational studies

ANSWER: A C

Explanation:

Out-of-range errors are data-quality problems in which a recorded value falls outside the valid limits established for a field, such as an age below zero, a percentage above 100, or an order date outside the organization's reporting period. **Relational databases** make these errors easier to detect because they can define data types and enforce integrity rules through constraints. For example, a database can use a CHECK constraint to require a quantity to fall within an approved minimum and maximum, preventing invalid values or identifying records that need review. SQL queries can also filter for values outside expected boundaries.

Spreadsheets support the same type of review through data-validation rules, formulas, filters, sorting, and conditional formatting. A validation rule can restrict entry to a specified numeric range, while conditional formatting can visibly flag values outside that range in an existing dataset. These capabilities let analysts locate exceptions quickly before using the data for summaries, visualizations, or decisions. Microsoft describes database constraints in its SQL documentation and describes spreadsheet validation features in its Excel support guidance: [SQL Server CHECK constraints](#) and [Excel data validation](#).

QUESTION NO: 5

A researcher seeks to pass a bond issue and asks a sample of respondents who have a bachelor's degree if they are voting in favor of the bond because it would be beneficial to the county.

Which type of error does this represent?

- A. Faulty operationalization
- B. Response bias
- C. Confusion of association and causality
- D. Selection bias

ANSWER: D

Explanation:

Selection bias is present because the researcher limits the sample to people who have earned a bachelor's degree while attempting to draw a conclusion about support for a county bond issue. The relevant population is the broader group of eligible county voters, not solely voters with a particular educational credential. Restricting participation in this way creates a sample that may differ systematically from the voting population in characteristics associated with opinions, civic participation, household circumstances, or priorities for public spending.

For a poll to provide useful evidence about the likely outcome of the bond vote, its sampling process should give the relevant voter population an appropriate opportunity to be represented. Educational attainment is a meaningful demographic characteristic in voting research, so a sample composed only of bachelor's-degree holders cannot reliably estimate countywide support without a sound design that represents other education groups as well. This is a problem introduced at the point the respondents are selected, which is the defining feature of selection bias. The U.S. Census Bureau documents educational attainment as a standard population characteristic used in survey analysis: [U.S. Census Bureau—Educational Attainment](#). Guidance on probability sampling and representative samples is also available from [Penn State STAT 506](#).

QUESTION NO: 6

Which graphical representation depicts the steps that compose a process?

- A. Histogram
- B. Control chart
- C. Flowchart
- D. Scatterplot

ANSWER: C

Explanation:

Flowchart is correct because it graphically maps the sequence of activities, decisions, inputs, and outputs that make up a process. A process flowchart uses connected symbols and directional arrows to show how work moves from its starting point through each action or decision to a final outcome. This visual representation makes a process easier to document, communicate, and analyze consistently among stakeholders.

In data-driven decision-making and process-improvement work, creating a flowchart helps a team establish a shared understanding of the current process before collecting data or proposing changes. It can reveal the handoffs, decision points, rework loops, and possible bottlenecks that affect performance. The resulting map also provides a practical foundation for identifying where measures should be gathered and where an improvement effort may have the greatest impact. Standard flowchart notation commonly distinguishes process actions, decisions, and data or input/output elements, enabling readers to follow the workflow logically. The American Society for Quality describes flowcharts as a way to display the steps of a process sequentially, and the National Institute of Standards and Technology identifies process mapping as a useful tool for understanding and improving organizational processes. See [ASQ's flowchart resource](#) and [NIST Baldrige self-assessment resources](#).

QUESTION NO: 7

A county government must increase trust among voters that their tallying machines are accurately calibrated to count their votes. Each department is tasked with creating an online marketing campaign; however, the budget for these campaigns is limited.

How can the county apply data analytic approaches to allocate funds to each department?

- A. By measuring the number of voter complaints per department
- B. By benchmarking the voter turnout rates in each county
- C. By surveying the county controllers
- D. By surveying employees on polling strategies

ANSWER: A

Explanation:

By measuring the number of voter complaints per department is correct because complaint data provide a relevant, quantifiable indicator of where voter confidence concerns are concentrated. The county can aggregate complaints by department, standardize the measure when appropriate (for example, by the number of voters served or ballots processed), and identify departments with the greatest demonstrated need for communication and trust-building outreach. Decision-makers can then direct a larger share of the limited online-marketing budget to areas where voters are reporting the most concerns about the voting process or equipment calibration. This is a practical data-driven allocation approach: it uses operational data tied directly to the campaign objective, rather than allocating funds solely through opinion or an unrelated outcome measure. The county should also define consistent complaint categories and review trends over time so that resources address recurring or increasing confidence issues. Performance-management guidance supports using meaningful measures connected to an agency's goals to inform management decisions, while election-administration guidance emphasizes collecting and using data to improve voter-facing processes. See the [GAO performance-measurement guide](#) and the [U.S. Election Assistance Commission's Election Administration and Voting Survey](#).

QUESTION NO: 8

A hospital system wants to survey patient satisfaction across its emergency departments. The system operates urban, suburban, and rural locations, and leaders want each location type represented in the results.

Which two actions would strengthen the representativeness of the sample?

Choose 2 answers.

- A. Randomly select patients within each location type.
- B. Survey only patients who submit online comments.
- C. Create separate samples for urban, suburban, and rural locations.
- D. Survey patients at the highest-volume location only.

ANSWER: A C

Explanation:

Creating separate groups for urban, suburban, and rural emergency departments and randomly selecting patients within each group is a stratified sampling approach. It ensures that each important location type is represented rather than allowing a larger group to dominate the survey results.

Using random selection within each group also reduces the risk that staff members select only convenient patients or patients who are expected to provide favorable responses. Surveying only patients who submit online comments creates voluntary-response bias, while surveying a single high-volume location does not represent the full hospital system.

QUESTION NO: 9

Which type of study is also known as a quasi-experimental study?

- A. Observational study
- B. Blind study
- C. Hypothesis testing
- D. Content validity

ANSWER: A

Explanation:

Observational study is the correct answer in the terminology used for this data-driven decision-making context. A quasi-experimental design examines the effect of an exposure, program, policy, or treatment when participants have not been randomly assigned by the researcher. Instead, the analyst works with naturally occurring groups or conditions and compares outcomes using available data. For example, an organization could assess performance measures before and after a policy change, or compare results between locations that adopted a new process at different times.

Because assignment is not randomized, careful study design and analysis are needed to account for meaningful preexisting differences between groups. Analysts may use comparison groups, matching, regression, or before-and-after measurements to make the evidence more useful for a decision. These studies are especially valuable when random assignment would be impractical, disruptive, or unethical. The U.S. National Library of Medicine describes quasi-experimental studies as intervention studies without random assignment, and CDC guidance recognizes quasi-experimental approaches as useful for evaluating real-world programs and policies. See [NCBI Bookshelf](#) and [CDC Program Evaluation Framework](#).

QUESTION NO: 10

When researchers are studying the effect of new drug treatments on patients, bias can be introduced by patients if they are aware of who receives the placebo.

Which type of research design eliminates this type of bias?

- A. Observational study
- B. Prospective cohort study
- C. Time series study
- D. Blind study

ANSWER: D

Explanation:

Blind study is correct because blinding prevents participants from knowing whether they received the investigational drug or a placebo. That knowledge can influence patients' expectations, symptom reporting, adherence, and other behavior during the trial. For example, a patient who believes they are receiving an active treatment may perceive or report improvement differently from a patient who knows they received a placebo. Keeping treatment assignment concealed from participants reduces this expectation-driven influence, commonly called participant bias or placebo-related bias, and makes observed differences between treatment groups more credible. In a drug trial, a single-blind design specifically addresses bias introduced by patients because the participants are unaware of their assigned treatment. Researchers may also use double-blind designs, in which both participants and investigators assessing outcomes are unaware of assignments, to guard against additional sources of bias. Blinding is therefore a core feature of well-designed controlled clinical trials when subjective outcomes, behavioral responses, or expectations could affect the results. The U.S. Food and Drug Administration describes blinding as withholding treatment-assignment information to minimize bias in clinical investigations. See the FDA's [E9 Statistical Principles for Clinical Trials](#) and the National Cancer Institute's [definition of a blinded study](#).

QUESTION NO: 11

What results from starting an analysis with flawed data?

Choose 2 answers.

- A. Spreadsheets must be used to increase the likelihood of analyzing the flawed data.
- B. More time is spent managing data than analyzing data.
- C. Data must be put in a table or a chart so that errors can be more easily detected.
- D. Missing data tend to skew the results of the analysis.

ANSWER: B D

Explanation:

More time is spent managing data than analyzing data because flawed input requires analysts to identify quality problems, reconcile inconsistent values, investigate anomalies, and clean or validate records before the data can support trustworthy analysis. These preparation activities are essential parts of a sound analytical process, but poor initial quality increases their scope and delays the point at which useful interpretation and decision support can begin. The National Institute of Standards and Technology describes data quality as involving characteristics such as accuracy, completeness, consistency, and validity; deficiencies in those characteristics must be addressed to ensure data are fit for their intended use. See [NIST Big Data Interoperability Framework](#).

Missing data tend to skew the results of the analysis because absent observations can alter distributions, summary measures, estimated relationships, and model results. In particular, when records are missing systematically rather than randomly, the remaining observations may no longer represent the population or process being studied. Analysts therefore need to assess the amount and pattern of missingness and use an appropriate treatment before drawing conclusions. The U.S. National Library of Medicine discusses how missing data can introduce bias and affect statistical inferences in [Missing data: the importance and the effect of missing data from clinical research](#).

QUESTION NO: 12

A school board has nine voting members. Five members need to be chosen each year for the finance committee. Why should the combination technique be used when selecting committee members?

- A. Because there is a correlation between committee member meeting attendance and selection
- B. Because the selection order of committee members is not important
- C. Because the selection order of committee members is important
- D. Because there is not a correlation between committee member meeting attendance and selection

ANSWER: B

Explanation:

Because the selection order of committee members is not important. A combination is the appropriate counting method when the outcome is a group or set and rearranging the same selected people does not produce a new outcome. Here, the finance committee consists of five of the nine voting members. Selecting the same five people in a different sequence—for example, naming one member first and another second—creates exactly the same committee. Thus, the calculation concerns the number of unique groups of five members, expressed as “9 choose 5,” rather than the number of ordered arrangements. The number of possible committees is $\binom{9}{5}=126$. This distinction is fundamental in data-driven decision making because choosing the correct counting approach ensures that possible outcomes and probabilities are calculated accurately. The order-independent nature of committee membership is a standard application of combinations. For additional guidance on combinations and how they differ from permutations, see [Khan Academy's combinations review](#) and the [OpenStax explanation of combinations and permutations](#).

QUESTION NO: 13

Which two results occur when the null hypothesis is accepted using an F-test?

Choose 2 answers.

- A. The test statistic is less than the critical value.
- B. The test statistic is greater than the critical value.
- C. There appears to be a difference between the two samples.
- D. There appears to be no difference between the two samples.

ANSWER: A D

Explanation:

In common introductory F-test terminology, accepting the null hypothesis means the observed F statistic does not fall in the rejection region. For the usual upper-tailed F decision rule, this occurs when the test statistic is less than the critical value. The observed ratio of variation is therefore not large enough, at the selected significance level, to provide statistically significant evidence against the null hypothesis.

Accordingly, there appears to be no difference between the two samples in the population parameter being evaluated by the particular F-test. For example, an F-test may evaluate whether population variances differ, while analysis of variance uses an F statistic to evaluate whether group means differ. In either use, the result means the data do not provide sufficient statistical evidence to conclude a difference exists. More precise statistical language is usually to *fail to reject* the null hypothesis, rather than to prove or definitively accept it; nevertheless, this question uses “accepted” to describe that same non-rejection outcome. The decision is based on comparing the calculated statistic with the applicable critical value for the test’s degrees of freedom and significance level. See the [NIST discussion of F tests](#) and [NIST guidance on hypothesis-test decisions](#).

QUESTION NO: 14

What is true about outliers?

Choose 2 answers.

- A. All outliers are statistically significant when using a normal distribution.
- B. All observed outliers should be eliminated from a study prior to analysis.
- C. Outliers that are miskeyed can be corrected prior to analysis.
- D. Outliers detected in a study are useful in determining if something does not belong in the study.

ANSWER: C D

Explanation:

Outliers that are miskeyed can be corrected prior to analysis because a value known to result from a data-entry, transcription, measurement-recording, or coding mistake does not represent the actual observation. The appropriate practice is to verify the original source and correct the recorded value while retaining a documented audit trail of the change. This protects data quality without arbitrarily changing valid observations.

Outliers detected in a study are useful in determining if something does not belong in the study because unusual values prompt investigation of the observation, its source, and its relationship to the target population and defined study scope. That investigation may identify a record from the wrong population, an invalid measurement, a process exception, or a rare but genuine event. Consequently, outlier review is a data-validation and analytical step rather than an automatic deletion rule. NIST notes that exploratory data analysis includes identifying unusual observations and investigating their possible causes, while its data-quality guidance emphasizes checking and documenting data before analysis. See the [NIST Engineering Statistics Handbook discussion of outliers](#) and the [NIST guidance on data screening and analysis](#).