

CUET PG 2026 Bioinformatics Question Paper with Solutions

Time Allowed :1 Hour 30 Minutes	Maximum Marks :300	Total Questions :75
---------------------------------	--------------------	---------------------

General Instructions

Read the following instructions very carefully and strictly follow them:

- The Question Paper consists of 75 Multiple Choice Questions (MCQs).
- Total marks for the exam is 300.
- The total duration is 90 minutes; a timer on the screen will display the remaining time, and the exam will automatically submit when the time reaches zero.
- For every correct answer, 4 marks (+4) will be awarded to the candidate, and a 1 mark (-1) will be deducted for every incorrect answer.
- Question papers are available in both English and Hindi.

1. Which algorithm is used for performing Global Sequence Alignment?

- (A) Smith–Waterman Algorithm
- (B) Needleman–Wunsch Algorithm
- (C) BLAST Algorithm
- (D) FASTA Algorithm

Correct Answer: (B) Needleman–Wunsch Algorithm

Solution:

This question addresses a core concept in bioinformatics: sequence alignment. The goal of sequence alignment is to arrange DNA, RNA, or protein sequences to identify regions of similarity that may be a consequence of functional, structural, or evolutionary relationships. Alignments can be performed globally, comparing sequences across their entire length, or locally, to find the best-matching segments within them.

Understanding the Question The question asks to identify the specific algorithm designed for performing a global alignment of two sequences.

Key Concepts and Approach The key is to distinguish between global and local alignment algorithms. The approach involves matching the type of alignment with its corresponding foundational algorithm.

Detailed Solution

1. **Global Alignment Goal:** Global sequence alignment attempts to find the best possible alignment between two sequences from end to end. This method is most suitable when comparing two sequences that are similar in length and are expected to share similarity across their entire span.

2. **The Needleman–Wunsch Algorithm:** This is the classic dynamic programming algorithm developed specifically for global alignment. It constructs a matrix to score all possible alignments and then traces back from the final cell to find the optimal alignment path that maximizes the score.
3. **Contrasting Other Options:**
 - The **Smith–Waterman algorithm** is used for **local alignment**, finding the most similar subsequences.
 - **BLAST** and **FASTA** are heuristic algorithms designed for rapid similarity searches in large databases, not for computing the mathematically optimal global alignment.
4. **Conclusion:** Therefore, the Needleman–Wunsch algorithm is the standard method for global sequence alignment.

💡 Quick Tip

For optimal sequence alignment: **Global Alignment** uses the **Needleman–Wunsch** algorithm, while **Local Alignment** uses the **Smith–Waterman** algorithm.

2. What is the primary difference between PAM and BLOSUM scoring matrices?

- (A) PAM matrices are based on global alignments, while BLOSUM matrices are based on local alignments
- (B) PAM matrices are used only for DNA sequences, while BLOSUM matrices are used for proteins
- (C) PAM matrices are derived from highly divergent sequences, while BLOSUM matrices use closely related sequences
- (D) PAM matrices are heuristic methods, while BLOSUM matrices use dynamic programming

Correct Answer: (A) PAM matrices are based on global alignments, while BLOSUM matrices are based on local alignments

Solution:

This question focuses on the scoring matrices used in protein sequence alignment. These matrices are essential because they assign a score to every possible substitution of one amino acid for another, reflecting the likelihood of that mutation occurring through evolution. PAM and BLOSUM are the two most widely used families of substitution matrices, but they are derived using different methodologies.

Understanding the Question The question asks for the fundamental distinction in the way PAM and BLOSUM matrices are constructed.

Key Concepts and Approach The core concepts are the derivation methods for PAM and BLOSUM matrices. The approach is to compare the source of alignment data used to calculate the substitution frequencies for each matrix type.

Detailed Solution

1. **PAM (Point Accepted Mutation) Matrices:** These matrices are based on an explicit evolutionary model. They were constructed by observing amino acid substitutions in **global alignments** of very closely related proteins (over 85% identical). The probabilities for more distant relationships (e.g., PAM250) are extrapolated from these initial observations.
2. **BLOSUM (BLOcks SUBstitution Matrix) Matrices:** These matrices are derived empirically from observing substitutions in ungapped, conserved regions (**local alignments** or "blocks") of more distantly related proteins. For example, BLOSUM62 is derived from proteins that share no more than 62% identity, making it suitable for finding similarities between more divergent sequences.
3. **The Core Difference:** The foundational difference is the alignment data used. PAM relies on an evolutionary model built from **global alignments** of highly similar sequences, while BLOSUM is based on direct observation of substitutions within conserved **local alignment blocks** from a broader range of sequences.

💡 Quick Tip

PAM matrices are derived from an evolutionary model based on **global alignments**.
BLOSUM matrices are derived empirically from conserved **local alignment blocks**.

3. Which biological database is considered a curated, secondary database for protein sequences?

- (A) GenBank
- (B) UniProtKB/Swiss-Prot
- (C) EMBL
- (D) DDBJ

Correct Answer: (B) UniProtKB/Swiss-Prot

Solution:

This question distinguishes between different types of biological databases. Primary databases serve as archives, storing raw sequence data as submitted by researchers. In contrast, secondary databases add a layer of value by processing, integrating, and manually reviewing (curating) information from primary sources to provide high-quality, annotated entries.

Understanding the Question The question asks to identify which of the listed options is a secondary database known for its expert curation of protein sequence data.

Key Concepts and Approach The key is to differentiate between primary (raw data repositories) and secondary (curated) databases. The approach involves categorizing each of the given database options.

Detailed Solution

1. **Primary Nucleotide Databases:** **GenBank** (USA), **EMBL** (Europe), and **DDBJ** (Japan) are the three major primary databases for nucleotide sequences. They are part of the International Nucleotide Sequence Database Collaboration (INSDC) and primarily contain raw DNA and RNA sequences with annotations provided by the submitters.
2. **A Curated Secondary Protein Database:** **UniProtKB/Swiss-Prot** is a high-quality, manually annotated, and non-redundant protein sequence database. Expert curators review scientific literature and computational analyses to add extensive information to each entry, including protein function, domains, modifications, and structural features.
3. **Conclusion:** While the other options are primary archives, UniProtKB/Swiss-Prot is renowned for being a curated secondary database, making it a more reliable source for detailed protein information.

💡 Quick Tip

Primary databases like **GenBank** and **EMBL** are raw data archives. A secondary database like **UniProtKB/Swiss-Prot** provides high-quality, manually reviewed (curated) protein information.

4. In the Michaelis-Menten equation, what does the constant K_m represent?

- (A) Maximum reaction velocity
- (B) Substrate concentration at half of V_{max}
- (C) Enzyme concentration
- (D) Rate of enzyme degradation

Correct Answer: (B) Substrate concentration at half of V_{max}

Solution:

This question pertains to enzyme kinetics, specifically the Michaelis-Menten model, which is a cornerstone for understanding enzyme function. The model describes how the initial rate of an enzyme-catalyzed reaction (v) is affected by the concentration of its substrate ($[S]$). The equation includes two key constants: V_{max} (the maximum reaction rate) and K_m (the Michaelis constant).

Understanding the Question The question asks for the physical or operational definition of the Michaelis constant, K_m .

Key Formula and Approach The approach is to analyze the Michaelis-Menten equation to understand the relationship between v , $[S]$, and K_m . The equation is:

$$v = \frac{V_{max}[S]}{K_m + [S]}$$

Detailed Solution

1. **Setting the Condition:** To define K_m , let's consider the specific condition where the reaction velocity v is exactly half of the maximum velocity, i.e., $v = \frac{V_{max}}{2}$.

2. **Substituting into the Equation:** We substitute this condition into the Michaelis-Menten equation:

$$\frac{V_{max}}{2} = \frac{V_{max}[S]}{K_m + [S]}$$

3. **Solving for $[S]$:** We can cancel V_{max} from both sides:

$$\frac{1}{2} = \frac{[S]}{K_m + [S]}$$

Cross-multiplying gives:

$$K_m + [S] = 2[S]$$

Subtracting $[S]$ from both sides yields:

$$K_m = [S]$$

4. **Conclusion:** This derivation shows that K_m is numerically equal to the substrate concentration at which the reaction rate is half of its maximum. It is also an inverse measure of the enzyme's affinity for its substrate (a lower K_m implies higher affinity).

💡 Quick Tip

The Michaelis constant, K_m , is the substrate concentration $[S]$ needed for an enzyme to operate at half of its maximum velocity (V_{max}).

5. Which amino acid is responsible for forming disulfide bridges in a protein structure?

- (A) Glycine
- (B) Cysteine
- (C) Lysine
- (D) Alanine

Correct Answer: (B) Cysteine

Solution:

This question deals with protein chemistry and the forces that stabilize a protein's three-dimensional structure. While many non-covalent interactions (like hydrogen bonds) hold a protein together, disulfide bridges are strong covalent bonds that can form cross-links within a polypeptide chain or between different chains, significantly reinforcing the protein's structure.

Understanding the Question The question asks to identify the specific amino acid whose chemical properties allow it to form these disulfide bonds.

Key Concepts and Approach The key concept is the chemical structure of amino acid side chains. The approach involves identifying which of the 20 standard amino acids contains a sulfur group capable of forming a disulfide (-S-S-) bond.

Detailed Solution

1. **What is a Disulfide Bridge?:** A disulfide bridge is a covalent bond formed between two sulfur atoms. This bond is created through an oxidation reaction involving two thiol groups (-SH).
2. **Identifying the Right Amino Acid:** Among the standard amino acids, only **Cysteine** has a side chain that contains a thiol group, also known as a sulfhydryl group (-SH).
3. **The Formation Reaction:** When two Cysteine residues are brought into close proximity during protein folding, their sulfhydryl groups can be oxidized, losing their hydrogen atoms and forming a covalent bond between the two sulfur atoms. This (-S-S-) linkage is the disulfide bridge.
4. **Structural Importance:** These bonds are crucial for stabilizing the tertiary and quaternary structures of many proteins, especially those that are secreted from the cell, such as antibodies and insulin.

💡 Quick Tip

The amino acid **Cysteine** contains a sulfur-hydrogen group (-SH). The oxidation of two such groups from nearby Cysteine residues creates a strong covalent **disulfide bond** (-S-S-), stabilizing protein structure.

6. What is the specific function of the BLASTn program?

- (A) Compare protein sequences
- (B) Compare nucleotide sequences
- (C) Predict protein structure
- (D) Align protein structures

Correct Answer: (B) Compare nucleotide sequences

Solution:

This question is about the BLAST (Basic Local Alignment Search Tool) suite, a cornerstone of modern bioinformatics. BLAST is not a single program but a family of fast, heuristic algorithms designed to find regions of local similarity between sequences. Different programs within the suite are specialized for comparing different types of query and database sequences (nucleotide or protein).

Understanding the Question The question asks for the specific purpose of the 'BLASTn' variant.

Key Concepts and Approach The key concept is the naming convention used in the BLAST suite, which indicates the type of sequences being compared. The approach is to decode the name 'BLASTn' to determine its function.

Detailed Solution

1. **Decoding the Name:** The letter 'n' in **BLASTn** stands for **nucleotide**.

2. **Defining the Function:** This tells us that BLASTn is designed to take a nucleotide sequence (like DNA or mRNA) as a query and search it against a database of nucleotide sequences. It identifies sequences in the database that are similar to the query sequence.
3. **Contrasting with Other BLAST Programs:** This is distinct from other members of the suite, such as:
 - **BLASTp:** Compares a **p**rotein query against a **p**rotein database.
 - **BLASTx:** Translates a **n**ucleotide query in all six reading frames and compares the resulting amino acid sequences against a **p**rotein database.
4. **Conclusion:** The specific function of BLASTn is to perform nucleotide versus nucleotide sequence comparisons.

💡 Quick Tip

The naming in the BLAST suite is a guide: BLAST**n** is for **n**ucleotide vs. **n**ucleotide searches, while BLAST**p** is for **p**rotein vs. **p**rotein searches.

7. Which tool is most commonly used for Multiple Sequence Alignment (MSA)?

- (A) BLAST
- (B) ClustalW
- (C) FASTA
- (D) Needleman–Wunsch

Correct Answer: (B) ClustalW

Solution:

This question focuses on Multiple Sequence Alignment (MSA), a process that extends pairwise alignment to three or more sequences. MSA is a powerful tool in molecular biology and bioinformatics used to identify conserved sequence regions, analyze evolutionary relationships (phylogenetics), and predict the secondary or tertiary structure of new proteins.

Understanding the Question The question asks to identify a well-known tool specifically designed for performing MSA, as opposed to other sequence analysis tasks.

Key Concepts and Approach The key is to differentiate between tools for pairwise alignment, database searching, and multiple sequence alignment. The approach involves categorizing the function of each tool listed in the options.

Detailed Solution

1. **Analyzing the Options:**

- **BLAST** and **FASTA** are heuristic tools primarily used for rapid database searching to find sequences similar to a query sequence.

- **Needleman–Wunsch** is an algorithm for optimal *pairwise* (two sequences) global alignment. It is not designed for multiple sequences.
2. **Identifying the MSA Tool: ClustalW** (and its more modern versions like Clustal Omega) is one of the most famous and widely-used programs for creating multiple sequence alignments. It employs a progressive alignment method, where it first aligns the most similar pairs of sequences and then progressively adds more distant sequences to the growing alignment.
 3. **Conclusion:** Among the choices provided, ClustalW is the tool that is specifically and most commonly associated with Multiple Sequence Alignment.

💡 Quick Tip

While BLAST is for database searching and Needleman-Wunsch is for pairwise alignment, **ClustalW** is a classic and widely-used tool specifically for **Multiple Sequence Alignment (MSA)**.

8. The p53 protein is primarily known for which cellular role?

- (A) DNA replication
- (B) Tumor suppression
- (C) Protein synthesis
- (D) Cell membrane transport

Correct Answer: (B) Tumor suppression

Solution:

This question is about a critical protein in cell biology and cancer research. The p53 protein acts as a central hub in a network of signaling pathways that protect the cell from various stresses, particularly those that can damage DNA and lead to cancer. Its role is so fundamental that it is often called the "guardian of the genome."

Understanding the Question The question asks for the main, well-established function of the p53 protein within a cell.

Key Concepts and Approach The key concepts are cell cycle control, DNA repair, and apoptosis (programmed cell death). The approach is to describe how p53 responds to cellular damage and how this response prevents cancer.

Detailed Solution

1. **Response to Cellular Stress:** The p53 protein is activated when a cell experiences stress, most notably DNA damage (e.g., from UV radiation or chemical mutagens).
2. **Key Functions of Activated p53:** Once activated, p53 acts as a transcription factor, initiating one of two main pathways:

- **Cell Cycle Arrest:** It can halt the cell division cycle, providing the cell with time to repair the damaged DNA before it can be replicated.
- **Apoptosis:** If the DNA damage is too severe to be repaired, p53 can trigger apoptosis, or programmed cell death. This eliminates the damaged cell, preventing it from passing on mutations.

3. **The Result: Tumor Suppression:** By preventing cells with damaged DNA from proliferating, p53 effectively stops the development of tumors at an early stage. This is why the gene encoding p53 is the most frequently mutated gene in human cancers.

💡 Quick Tip

The **p53 protein**, known as the "guardian of the genome," acts as a **tumor suppressor** by either halting the cell cycle for DNA repair or initiating cell death (apoptosis) in severely damaged cells.

9. Which molecular technique is used to detect specific mRNA molecules in a sample?

- (A) Southern Blot
- (B) Northern Blot
- (C) Western Blot
- (D) PCR

Correct Answer: (B) Northern Blot

Solution:

This question covers the classic "blotting" techniques in molecular biology. These methods are used to identify a specific molecule of interest (DNA, RNA, or protein) within a complex mixture. They generally involve separating the molecules by size using gel electrophoresis, transferring them to a solid membrane, and then using a labeled probe to detect the target molecule.

Understanding the Question The question asks to identify the specific blotting technique designed for the detection of RNA, particularly messenger RNA (mRNA).

Key Concepts and Approach The key concept is the distinction between the three main types of blotting techniques. A common mnemonic, "SNOW DROP," helps to remember them:

Southern → DNA

Northern → RNA

O - O

Western → Protein

Detailed Solution

1. **Southern Blot:** This technique is used to detect specific sequences of **DNA**. It is named after its inventor, Edwin Southern.

2. **Western Blot:** This method uses antibodies as probes to detect specific **proteins**.
3. **Northern Blot:** This technique was developed as an analogy to the Southern blot and is used to detect and quantify specific sequences of **RNA**, such as mRNA. It is a powerful tool for studying gene expression, as the amount of a specific mRNA in a sample often correlates with the level of activity of that gene.
4. **Conclusion:** Based on this distinction, the Northern blot is the correct technique for detecting mRNA.

💡 Quick Tip

Remember the blotting techniques by the type of molecule they detect: **Southern** for **DNA**, **Northern** for **RNA** (including mRNA), and **Western** for **Protein**.

10. In computer science, what does the Markov Property assume about future states?

- (A) Future states depend on all previous states
- (B) Future states depend only on the current state
- (C) Future states depend only on the initial state
- (D) Future states are completely random

Correct Answer: (B) Future states depend only on the current state

Solution:

This question concerns the Markov Property, a foundational concept in probability theory and stochastic processes that has wide applications in computer science, including bioinformatics (e.g., Hidden Markov Models for gene finding), speech recognition, and machine learning. The property describes a specific kind of "memorylessness" in a system's evolution over time.

Understanding the Question The question asks for the core assumption of the Markov Property regarding the prediction of a system's future behavior.

Key Concepts and Approach The central concept is that of a "memoryless" process. The approach is to explain what this means in terms of state transitions and conditional probabilities.

Detailed Solution

1. **The Core Idea:** The Markov Property states that for a stochastic process, the future is independent of the past, given the present. In other words, to predict where the system will go next, you only need to know where it is right now; you do not need to know the history of how it got there.
2. **The "Memoryless" Nature:** This is often called the "memoryless" property. The process "forgets" all previous states and its future evolution depends solely on its current state.

3. **Formal Definition:** Mathematically, the conditional probability of transitioning to a future state, given the entire history of past states, is the same as the conditional probability of transitioning to that future state given only the current state.
4. **Conclusion:** Therefore, the Markov Property assumes that the probability of moving to any future state depends only on the current state.

💡 Quick Tip

The **Markov Property** is the "memoryless" principle: the future depends only on the **present**, not on the past history that led to the present state.

11. Which database serves as the main archive for 3D macromolecular structures?

- (A) GenBank
- (B) Protein Data Bank (PDB)
- (C) UniProt
- (D) EMBL

Correct Answer: (B) Protein Data Bank (PDB)

Solution:

This question is about the different types of biological databases and their specific roles. While many databases store one-dimensional sequence information (the order of nucleotides or amino acids), structural biology requires a dedicated repository for the three-dimensional atomic coordinates of macromolecules, which describe their complex shapes.

Understanding the Question The question asks to identify the primary, worldwide repository for 3D structural data of biological molecules like proteins and nucleic acids.

Key Concepts and Approach The key is to differentiate between sequence databases and structural databases. The approach involves identifying the function of each database listed in the options.

Detailed Solution

1. **Sequence Databases: GenBank, EMBL, and UniProt** are primarily sequence databases. GenBank and EMBL are archives for DNA/RNA sequences, while UniProt is focused on protein sequences and their functional annotation. They store information as 1D text strings.
2. **The Structural Database:** The **Protein Data Bank (PDB)** is the single global archive for experimentally determined 3D structures of proteins, nucleic acids, and their complex assemblies.
3. **Content of the PDB:** Researchers deposit atomic coordinate files generated from techniques like X-ray crystallography, NMR spectroscopy, and cryo-electron microscopy into the PDB. This data is essential for research in drug design, molecular modeling, and understanding biological function at a molecular level.

4. **Conclusion:** Therefore, the PDB is the main archive specifically for 3D macromolecular structures.

 Quick Tip

While GenBank and UniProt store 1D sequence data, the **Protein Data Bank (PDB)** is the dedicated global archive for **3D structural data** of biological macromolecules.

12. What is the purpose of Emulsion PCR (ePCR) in Next-Generation Sequencing?

- (A) To sequence DNA directly without amplification
- (B) To amplify individual DNA fragments in isolated microreactors
- (C) To detect proteins in sequencing samples
- (D) To separate RNA molecules from DNA

Correct Answer: (B) To amplify individual DNA fragments in isolated microreactors

Solution:

This question deals with a key step in certain Next-Generation Sequencing (NGS) workflows. NGS technologies require a large number of identical copies of each DNA fragment to be sequenced, a process known as clonal amplification. This ensures that the signal generated during the sequencing reaction is strong enough to be detected accurately. Emulsion PCR (ePCR) is one method for achieving this clonal amplification.

Understanding the Question The question asks for the specific role of ePCR in the context of preparing a DNA library for NGS.

Key Concepts and Approach The key concepts are "clonal amplification" and "microreactors." The approach is to describe the physical process of ePCR and explain how it achieves the goal of amplifying single DNA molecules in isolation.

Detailed Solution

1. **The Goal of Clonal Amplification:** The main goal is to take a library of millions of different DNA fragments and produce millions of clusters, where each cluster contains many identical copies of just one original fragment.
2. **Creating Microreactors:** In ePCR, a water-in-oil emulsion is created. This process forms millions of tiny, separate aqueous droplets suspended in oil. Each droplet acts as an independent, microscopic PCR tube or "microreactor."
3. **Isolating DNA Fragments:** The DNA library is diluted such that, ideally, each droplet encapsulates a single DNA template molecule along with a single bead and all the necessary PCR reagents.
4. **Amplification in Isolation:** PCR is then performed on the entire emulsion. Inside each droplet, the single DNA molecule is amplified, and the resulting identical copies

(amplicons) attach to the surface of the bead within that droplet. Because the droplets are separate, amplification from different templates does not mix.

5. **Conclusion:** Thus, the purpose of ePCR is to physically isolate individual DNA fragments in microreactors (droplets) to allow for their massive and parallel clonal amplification.

 Quick Tip

Emulsion PCR (ePCR) uses tiny water-in-oil droplets as isolated microreactors to clonally amplify single DNA fragments onto beads, a key preparation step for some Next-Generation Sequencing platforms.

13. Which component is strictly part of a Microprocessor System Unit, excluding peripherals?

- (A) Keyboard
- (B) Monitor
- (C) CPU
- (D) Printer

Correct Answer: (C) CPU

Solution:

This question tests the fundamental understanding of a computer's architecture. A computer system is broadly divided into two main parts: the central **System Unit** that contains the core processing components, and the **Peripheral Devices** that are connected to the system unit to provide input, output, and extended storage.

Understanding the Question The question asks to identify the component from the list that is an internal part of the system unit, not an external peripheral device.

Key Concepts and Approach The key is the distinction between internal processing hardware and external input/output (I/O) devices. The approach is to classify each option into one of these two categories.

Detailed Solution

1. **Defining Peripherals:** Peripherals are external devices that connect to the system unit to expand its capabilities.
 - A **Keyboard** is an input peripheral.
 - A **Monitor** is an output peripheral.
 - A **Printer** is an output peripheral.
2. **Defining the System Unit:** The system unit is the main enclosure that houses the computer's primary components, such as the motherboard, memory (RAM), and the central processing unit.

3. **Identifying the Core Component:** The **Central Processing Unit (CPU)** is the brain of the computer. It is located on the motherboard inside the system unit and is responsible for executing instructions and performing calculations. It is an essential internal component, not a peripheral.

💡 Quick Tip

The **System Unit** contains the core components like the **CPU** and **RAM**. **Peripherals** are external devices like the keyboard, monitor, and printer that connect to the system unit.

14. **How does a 2D array differ from a 1D array in a programming language like C++?**

- (A) A 2D array stores characters only
- (B) A 2D array has rows and columns
- (C) A 2D array stores only integers
- (D) A 2D array cannot store multiple values

Correct Answer: (B) A 2D array has rows and columns

Solution:

This question explores the concept of arrays, a fundamental data structure in programming. Arrays allow for the storage of multiple elements of the same data type. The "dimensionality" of an array refers to how these elements are organized and accessed.

Understanding the Question The question asks for the primary structural difference between a one-dimensional (1D) array and a two-dimensional (2D) array.

Key Concepts and Approach The key concept is dimensionality in data structures. The approach is to compare the logical structure and element access method for 1D and 2D arrays.

Detailed Solution

1. **One-Dimensional (1D) Array:** A 1D array organizes elements in a single, linear sequence. It can be thought of as a simple list. To access any element, you need only one index. For example, in C++, 'arr[5]' refers to the sixth element in a linear list named 'arr'.
2. **Two-Dimensional (2D) Array:** A 2D array organizes elements in a grid-like or tabular structure. It is composed of multiple **rows and columns**. To access any element, you need two indices: one to specify the row and another to specify the column. For example, 'matrix[2][3]' refers to the element in the third row and fourth column of a grid named 'matrix'.
3. **The Fundamental Difference:** The core distinction is the structure. A 1D array is a line of elements, while a 2D array is a table of elements. Both can store any data type (not just characters or integers) and are designed to hold multiple values.

💡 Quick Tip

A **1D array** is structured like a single list, accessed with one index (e.g., 'list[i]'). A **2D array** is structured like a table with **rows and columns**, accessed with two indices (e.g., 'table[row][col]').

15. Which type of bond characterizes the primary structure of DNA?

- (A) Hydrogen bond
- (B) Ionic bond
- (C) Phosphodiester bond
- (D) Disulfide bond

Correct Answer: (C) Phosphodiester bond

Solution:

This question examines the chemical bonds that define the structure of DNA. The structure of DNA is hierarchical: the primary structure is the linear sequence of nucleotides in a single strand, while the secondary structure is the famous double helix formed by two complementary strands. Different chemical bonds are responsible for each level of structure.

Understanding the Question The question asks to identify the specific type of covalent bond that links nucleotides together to form the backbone of a single DNA strand, which constitutes its primary structure.

Key Concepts and Approach The key is to distinguish between the bonds **within** a single DNA strand (forming the backbone) and the bonds **between** the two strands. The approach is to identify the chemical linkage that creates the sugar-phosphate backbone.

Detailed Solution

1. **Defining Primary Structure:** The primary structure of DNA is the sequence of its nucleotides (A, C, G, T) linked together in a chain. This chain is formed by a repeating sugar-phosphate backbone.
2. **The Backbone Bond:** The nucleotides in this backbone are joined by strong covalent bonds called **phosphodiester bonds**. This bond connects the 5' carbon of one deoxyribose sugar to the 3' carbon of the next via a phosphate group. This creates the directionality (5' to 3') of the DNA strand.
3. **Distinguishing Other Bonds:**
 - **Hydrogen bonds** are weaker bonds that form between the nitrogenous bases of the two opposing strands (A with T, G with C). They are responsible for holding the double helix together (secondary structure), not forming the primary structure.
 - **Disulfide bonds** are found in proteins, not DNA.
4. **Conclusion:** The defining bond of DNA's primary structure is the phosphodiester bond.

💡 Quick Tip

The **primary structure** of DNA (the single strand) is formed by strong **phosphodiester bonds** linking the sugar-phosphate backbone. The **secondary structure** (the double helix) is held together by weaker **hydrogen bonds** between base pairs.
