

# Design and Analysis of Algorithms

|                   |                             |
|-------------------|-----------------------------|
| <b>Developer:</b> | Dr. Katila Charles          |
| <b>Programme:</b> | MSC Artificial Intelligence |



## COURSE CONTENT

---

Introduction to algorithms: The role of algorithms in computing; Designing and Analysis of Algorithms; growth of functions; recurrence; Sorting and Order Statistics: Sorting in linear time; Medians and order statistics; Dynamic Programming: Common subproblems; Optimal substructure; Fibonacci Sequence; Minimum Edit Distance; greedy algorithms; Data structures for set manipulation problems: Queues; Stacks; Search tree; Randomized search tree; Red-black tree; Advanced topics: Greedy algorithms; The Fast Fourier Transform and its applications; Pattern Matching algorithms; Graphs; NP complete problems;



Maximum Modules for each course is set 12 – delete this placeholder after reading this instruction on your copy!

## LIST MODULES

---

|                                             |   |
|---------------------------------------------|---|
| MODULE 0: Course Overview and Expectations. | 7 |
| 0.1 Welcome to the Course                   | 7 |
| 0.2 Purpose of the course                   | 7 |
| 0.3 Course Content                          | 7 |
| 0.4 Expected Learning Outcomes              | 8 |
| 0.5 List of Modules                         | 8 |
| MODULE 1: INTRODUCTION TO ALGORITHMS        | 9 |
| Module Overview                             | 9 |
| Learning Outcomes                           | 9 |
| Learning Activities/Task List               | 9 |

|                                                                       |    |
|-----------------------------------------------------------------------|----|
| Review Reading Material                                               | 10 |
| Introduction to Problems                                              | 10 |
| The Role of Algorithms In Computing                                   | 11 |
| Algorithmic Problem-Solving                                           | 12 |
| Pseudocodes And Flowcharts                                            | 15 |
| Properties of Algorithms                                              | 19 |
| Examples of Algorithms in Daily Computing Tasks                       | 20 |
| Applications of Algorithms in Artificial Intelligence                 | 21 |
| ACTIVITY 2: online discussions:anagram detection algorithm            | 22 |
| Instructions:                                                         | 23 |
| ACTIVITY 3: Design and implementation of average calculator algorithm | 24 |
| REFERENCE LIST AND FURTHER READING                                    | 26 |
| ACTIVITY 4: End of Module Assessment                                  | 26 |
| MODULE 2: designing and analyzing algorithms                          | 29 |
| Module Overview                                                       | 29 |
| Learning Outcomes                                                     | 29 |
| Learning Activities/Task List                                         | 29 |
| Review Reading Material                                               | 30 |
| INTRODUCTION TO DESIGN AND ANALYSIS OF ALGORITHMS                     | 30 |
| Design principles                                                     | 30 |
| Greedy Algorithms                                                     | 31 |
| Dynamic Programming                                                   | 32 |
| Activity 2: Mastery Exercise on Design Principles in Algorithms       | 33 |
| Analysis of algorithms                                                | 35 |
| Time Complexity                                                       | 35 |
| Space Complexity                                                      | 35 |
| Big O Notation and Its Significance                                   | 36 |
| Activity 3: Mastery Exercise on Analysis of Algorithms                | 40 |
| REFERENCE LIST AND FURTHER READING                                    | 41 |
| Activity 4: End of Module Quiz                                        | 41 |
| MODULE 3: Growth of functions and recurrences                         | 43 |
| Module Overview                                                       | 43 |
| Learning Outcomes                                                     | 43 |
| Learning Activities/Task List                                         | 44 |
| Review Module Notes                                                   | 45 |
| Growth of Functions                                                   | 45 |
| Big O Notation ( $O$ )                                                | 45 |
| Big $\Theta$ Notation ( $\Theta$ )                                    | 45 |
| Big $\Omega$ Notation ( $\Omega$ )                                    | 45 |
| Comparing Different Growth Rates                                      | 46 |
| Activity 2: Practical Lab Activity                                    | 48 |
| Activity 3: Mastery Quiz                                              | 48 |
| Explanation:                                                          | 49 |

|                                                                                             |    |
|---------------------------------------------------------------------------------------------|----|
| Recurrences                                                                                 | 50 |
| Solving Recurrences                                                                         | 51 |
| The Master Theorem                                                                          | 52 |
| Activity 3: YouTube Video on Recurrence Relations (youtube link)                            | 53 |
| Activity 4 : Mastery Quiz on recursion and recurrences                                      | 55 |
| Activity 5 : end of module quiz                                                             | 55 |
| REFERENCE LIST AND FURTHER READING                                                          | 56 |
| MODULE 4: sorting and order statistics                                                      | 57 |
| Module Overview                                                                             | 57 |
| Learning Outcomes                                                                           | 57 |
| Learning Activities/Task List                                                               | 57 |
| Review Reading Materials                                                                    | 58 |
| Sorting in Linear Time                                                                      | 58 |
| Linear Time Sorting Algorithms                                                              | 59 |
| Comparison of Linear-Time Sorting Algorithms                                                | 63 |
| Lab Activity: Implementing and Analyzing Radix Sort                                         | 65 |
| Objective:                                                                                  | 65 |
| Scenario: Sorting Phone Numbers                                                             | 65 |
| Instructions:                                                                               | 66 |
| Questions for Reflection:                                                                   | 66 |
| Medians and Order Statistics                                                                | 66 |
| Median                                                                                      | 67 |
| Activity 4 : end of module quiz                                                             | 70 |
| MODULE 5: Dynamic Programming - Part 1                                                      | 71 |
| Module Overview                                                                             | 72 |
| Learning Outcomes                                                                           | 72 |
| Learning Activities/Task List                                                               | 73 |
| Activity 1: Review Reading Materials                                                        | 73 |
| Introduction to Dynamic Programming                                                         | 73 |
| Activity 2: Case Study: Optimizing Fibonacci Sequence Calculation Using Dynamic Programming | 78 |
| Further Reading and References on Dynamic Programming                                       | 80 |
| Activity 3: End of Module quiz                                                              | 80 |
| MODULE 6: Dynamic Programming - Part 2                                                      | 83 |
| Module Overview                                                                             | 83 |
| Learning Outcomes                                                                           | 84 |
| Learning Activities/Task List                                                               | 84 |
| Activity 1: Review Reading Materials                                                        | 85 |
| Minimum Edit Distance Problem                                                               | 85 |
| Longest Common Subsequence (LCS)                                                            | 88 |
| 3. Knapsack Problem                                                                         | 89 |
| Summary of problem Solving Approach                                                         | 91 |
| Activity 2: Lab Activity on Longest Common Subsequence (LCS)                                | 92 |

|                                                                                       |     |
|---------------------------------------------------------------------------------------|-----|
| Further Reading and References on Dynamic Programming                                 | 93  |
| Activity 3: End of Module Quiz                                                        | 93  |
| MODULE 7: Greedy algorithms                                                           | 96  |
| Module Overview                                                                       | 96  |
| Learning Outcomes                                                                     | 97  |
| Learning Activities/Task List                                                         | 97  |
| Review Module Reading Material                                                        | 98  |
| Introduction to Greedy Algorithms                                                     | 98  |
| Characteristics of Greedy Algorithms                                                  | 99  |
| Huffman Coding                                                                        | 99  |
| How Huffman Coding Works                                                              | 99  |
| Example of Huffman Coding                                                             | 100 |
| Encoding Example                                                                      | 101 |
| Activity 2: Lab Activity                                                              | 102 |
| Steps of Prim's Algorithm:                                                            | 102 |
| Activity 2: Understanding Minimum Spanning Tree and Prim's Algorithm                  | 103 |
| Activity 3: Mastery Quiz on Prim's Algorithm                                          | 104 |
| Kruskal's Algorithm                                                                   | 105 |
| Activity 3: Watch a YouTube Video on Kruskal's Algorithm                              | 108 |
| Dijkstra's Algorithm                                                                  | 110 |
| Activity 3 Discussion Forum Activity: Understanding Minimum Spanning Trees            | 110 |
| Further Reading and References on Greedy Algorithms                                   | 111 |
| Activity 4: End of Module Module Quiz                                                 | 112 |
| MODULE 8: Data Structures for Set Manipulation                                        | 113 |
| Module Overview                                                                       | 113 |
| Learning Outcomes                                                                     | 114 |
| Learning Activities/Task List                                                         | 114 |
| Review reading Material                                                               | 115 |
| Fundamental Data Structures:                                                          | 115 |
| Stacks:                                                                               | 116 |
| Binary Search Trees (BSTs):                                                           | 116 |
| Key Properties:                                                                       | 116 |
| Randomized Search Trees:                                                              | 119 |
| Performance:                                                                          | 120 |
| Advanced Data Structures:                                                             | 121 |
| AVL Trees:                                                                            | 125 |
| Activity 2: Laboratory Activity ( Implementing and Analyzing AVL Trees )              | 128 |
| Activity 3: Comparing Data Structures for Set Manipulation(Discussion Forum Activity) | 131 |
| Further Reading and References                                                        | 133 |
| Activity 4: End of the module Quiz                                                    | 133 |
| MODULE 9: Advanced Topics in Algorithms                                               | 136 |
| Module Overview                                                                       | 137 |
| Learning Outcomes                                                                     | 137 |

|                                                                             |     |
|-----------------------------------------------------------------------------|-----|
| Learning Activities/Task List                                               | 137 |
| Review Reading Material                                                     | 138 |
| Fast Fourier Transform (FFT)                                                | 138 |
| Pattern Matching Algorithms                                                 | 140 |
| 1. Naive String Matching Algorithm                                          | 140 |
| 2. Knuth-Morris-Pratt (KMP) Algorithm                                       | 141 |
| 3. Boyer-Moore Algorithm                                                    | 142 |
| Activity 2: Laboratory on Advanced Topics in Algorithms                     | 143 |
| Further Reading and References on advanced algorithms                       | 145 |
| Activity 3: End of Module Quiz                                              | 145 |
| MODULE 10: NP-Complete Problems                                             | 148 |
| Module Overview                                                             | 148 |
| Learning Outcomes                                                           | 148 |
| Learning Activities/Task List                                               | 149 |
| Review Module Material                                                      | 149 |
| Introduction to NP-Completeness                                             | 149 |
| The Cook-Levin Theorem                                                      | 150 |
| NP-Completeness and Reducibility                                            | 150 |
| Examples of NP-Complete Problems                                            | 151 |
| NP-Completeness Proofs                                                      | 152 |
| 3. Approximation Algorithms for NP-Hard Problems                            | 153 |
| Types of Approximation Algorithms                                           | 154 |
| Activity 2: Laboratory Activity on NP-Complete Problems(Group Lab Activity) | 155 |
| Further Reading and References on NP-Completeness                           | 157 |
| Activity 3: End of Module Quiz                                              | 158 |

# MODULE 0: COURSE OVERVIEW AND EXPECTATIONS.

## 0.1 Welcome to the Course

Welcome to **Design and Analysis of Algorithms**, a fundamental course in the MSc in Artificial Intelligence program. This course will provide a deep understanding of the development and analysis of algorithms, a cornerstone of AI. Mastering these concepts is essential for creating intelligent systems capable of efficiently processing data, learning from it, and making informed decisions.

Throughout this course, you'll delve into the core principles of algorithm design, study the growth of functions, and explore the critical role of data structures in problem-solving. We'll also cover advanced topics like the Fast Fourier Transform, pattern matching algorithms, and NP-complete problems, all of which are vital for tackling complex challenges in AI.

## 0.2 Purpose of the course

The purpose of this course is to introduce learners to various algorithmic strategies and techniques and also ways to analyse algorithm efficiency and to understand the theoretical foundations that underpin them

## 0.3 Course Content

Introduction to algorithms: The role of algorithms in computing; Designing and Analysis of Algorithms; growth of functions; recurrence; Sorting and Order Statistics: Sorting in linear time; Medians and order statistics; Dynamic Programming: Common subproblems; Optimal substructure; Fibonacci Sequence; Minimum Edit Distance; greedy algorithms; Data structures for set manipulation problems: Queues; Stacks; Search tree; Randomized search tree; Red-black tree; Advanced topics: Greedy algorithms; The Fast Fourier Transform and its applications; Pattern Matching algorithms; Graphs; NP complete problems;

## **0.4 Expected Learning Outcomes**

By the end of the course, the learner should be able to:

Describe time and space complexities of algorithms

Use various algorithmic techniques to solve problems efficiently

Analyse different types of algorithms for space and time efficiency

Choose appropriate algorithms for specific real world problems

## **0.5 List of Modules**

The course is divided into 10 modules:

1. Introduction to Algorithms
2. Designing and Analyzing Algorithms
3. Growth of Functions and Recurrence
4. Sorting and Order Statistics
5. Dynamic Programming Part 1
6. Dynamic Programming Part 2
7. Greedy Algorithms
8. Data Structures for Set Manipulation
9. Advanced Sorting and Pattern Matching
10. NP-Complete Problems and Approximation Algorithms



# MODULE 1: INTRODUCTION TO ALGORITHMS

---

INSTRUCTIONAL HOURS: 4, 5 or 6

## MODULE OVERVIEW

In this module, you will be introduced to the fundamental concepts of algorithms and their critical role in computing. The module is designed to lay the foundation for understanding how algorithms are developed, analyzed, and applied across various domains, especially in artificial intelligence.

## LEARNING OUTCOMES

At the end of this module, you will be able to:

1. Explain the role of algorithms in computing, including their significance and impact on various applications.
2. Describe fundamental algorithm design techniques such as divide and conquer, dynamic programming, and greedy algorithms.
3. Analyze the efficiency of algorithms using methods such as time and space complexity, and apply Big O notation.
4. Evaluate the use of algorithms in real-world computing scenarios and their effects on system performance.

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Discussion Forum
3. Lab Activity on Design of Algorithms
4. End of Module Quiz

## Introduction to Problems

### What is a problem in Computing?

A problem is a task that needs to be solved. It is typically described by a set of inputs and the expected outputs or results

### Properties of Problems

**Deterministic vs. Non-deterministic:** A problem is deterministic if a given input will always produce the same output. Non-deterministic problems may have different outputs for the same input.

**Well-defined vs. Ill-defined:** A problem is well-defined if it has clear, specific inputs and outputs, and a clear criterion for correctness. An ill-defined problem lacks these specifications.

**Computability:** This refers to whether a problem can be solved using an algorithm within finite time.

**Complexity:** This involves the resources required to solve the problem, typically time and space (memory).

## Types of Problems

**Decision Problems:** Problems that require a yes/no answer.

**Optimization Problems:** Problems that seek the best solution from a set of feasible solutions.

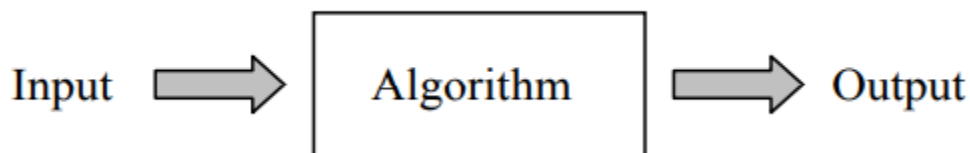
**Search Problems:** Problems that involve finding a specific item or set of items in a collection.

**Sorting Problems:** Problems that involve arranging data in a particular order.

## The Role of Algorithms In Computing

### What is an Algorithm?

An algorithm is a finite set of well-defined instructions for solving a problem or performing a task. This means that the instructions must be unambiguous, executable, and produce a result within a finite amount of time.



### Characteristics of Algorithms

**Finiteness:** An algorithm must always terminate after a finite number of steps.

**Definiteness:** Each step of the algorithm must be precisely defined; the actions to be performed should be clear.

**Input:** An algorithm should have zero or more inputs, which are provided before the algorithm starts.

**Output:** An algorithm should produce one or more outputs, which are the solution to the problem.

**Effectiveness:** The steps of the algorithm should be basic enough to be carried out by a human using a pen and paper.

### Example: Sorting a List of Numbers

Consider the task of sorting a list of numbers. An algorithm for this task might involve comparing each pair of numbers and swapping them if necessary until the list is in order. For instance, the Bubble Sort algorithm repeatedly traverses the list, compares adjacent elements, and swaps them if they are in the wrong order. This process is repeated until the list is sorted.

## Algorithmic Problem-Solving

Algorithmic problem-solving is the process of formulating a problem in a way that allows it to be solved by an algorithm. This process is foundational in computer science, as it enables us to automate the resolution of complex tasks.

### Steps in Algorithmic Problem-Solving

#### Problem Understanding:

Before designing an algorithm, it is essential to thoroughly understand the problem at hand. This involves:

**Clarifying the Problem Statement:** Ensuring that the problem is clearly defined, including the inputs, expected outputs, and any constraints.

**Identifying the Objectives:** Understanding what constitutes a successful solution and any specific goals (e.g., minimizing time or memory usage).

#### 2. Problem Formulation:

Once the problem is understood, it needs to be formulated in a way that can be addressed algorithmically. This includes:

**Defining Inputs and Outputs:** Clearly specify what data will be provided to the algorithm (inputs) and what it should produce (outputs).

**Specifying Constraints:** Any limits or conditions that the solution must adhere to, such as time limits or maximum memory usage.

#### 3. Designing the Algorithm:

This is the core step where the actual algorithm is developed. The key considerations include:

**Choosing the Right Approach:** Based on the problem type, decide on an approach (e.g., divide and conquer, dynamic programming, greedy algorithms).

**Step-by-Step Breakdown:** Break down the problem into smaller, manageable tasks and outline the steps needed to solve each.

**Pseudocode:** Write the algorithm in pseudocode, which is a high-level description of the steps in a format that is independent of any specific programming language. Alternatively you can use a flowchart to represent the algorithm

#### 4. **Analyzing the Algorithm:**

After designing the algorithm, it's crucial to analyze its efficiency:

**Time Complexity:** Determine how the runtime of the algorithm scales with the size of the input. Common notations include  $O(n)$ ,  $O(\log n)$ ,  $O(n^2)$ , etc.

**Space Complexity:** Assess the amount of memory the algorithm uses as the input size grows.

**Correctness:** Prove that the algorithm works correctly for all possible inputs.

#### 5. **Implementation:**

**After analysis,** the algorithm is implemented in a specific programming language. During implementation:

**Code Translation:** Convert the pseudocode into actual code, ensuring that all edge cases are handled.

**Debugging:** Test the algorithm with various inputs to ensure it works as expected. Debug any issues that arise.

#### 6. **Testing and Verification:**

Rigorous testing is done to ensure the algorithm works correctly. During testing the following is done:

**Unit Testing:** Test individual components or functions of the algorithm.

**Integration Testing:** Test how different parts of the algorithm work together.

**Performance Testing:** Test the algorithm under different conditions to evaluate its efficiency and scalability.

**Verification:** Ensure that the algorithm meets all requirements and constraints defined during the problem formulation stage.

## 7. Optimization:

After the initial implementation, look for ways to improve the algorithm's performance:

**Optimize Time Complexity:** Explore alternative approaches or data structures that could reduce the time complexity.

**Optimize Space Complexity:** Reduce memory usage by optimizing data storage or eliminating unnecessary variables.

## 9. Documentation:

Proper documentation is crucial for future maintenance and understanding:

## 10. Evaluation:

Finally, evaluate the algorithm in a broader context:

**Comparative Analysis:** Compare the designed algorithm with existing solutions to assess its relative strengths and weaknesses.

**Real-World Applicability:** Consider how the algorithm would perform in real-world scenarios, taking into account factors like data distribution and hardware constraints.

## Pseudocodes And Flowcharts

Algorithms are often presented using pseudo-code or flowcharts. Pseudocode is a way of describing an algorithm using a mix of plain English, mathematical symbols, and keywords from programming languages. It's a way to express the steps of an algorithm without being tied to the specific syntax of any programming language.

There isn't a strict format for writing pseudo-code, so different authors may use different styles. However, the main goal is to ensure that the pseudo-code is clear and easy to understand.

### Example of Pseudocode

Consider the pseudocode below for an algorithm used to calculate the Factorial of number n

```
Begin
  If n is 0 Then
    Return 1
  Else
    Set result to 1
    For i from 1 to n Do
      result = result * i
    EndFor
    Return result
  EndIf
End
```

## Explanation for the Pseudocode

The algorithm checks if  $n$  is 0 because the factorial of 0 is defined as 1.

If  $n$  is not 0, it initializes a variable `result` to 1.

It then loops from 1 to  $n$ , multiplying `result` by each number in this range.

After the loop, `result` holds the factorial of  $n$ , which is then returned.

Algorithms can also be represented using diagrams. One commonly used method is the flowchart. **Flowcharts** visually depict algorithms through a series of standardized symbols:

**Terminator Boxes:** Represent the start and end points of the algorithm.

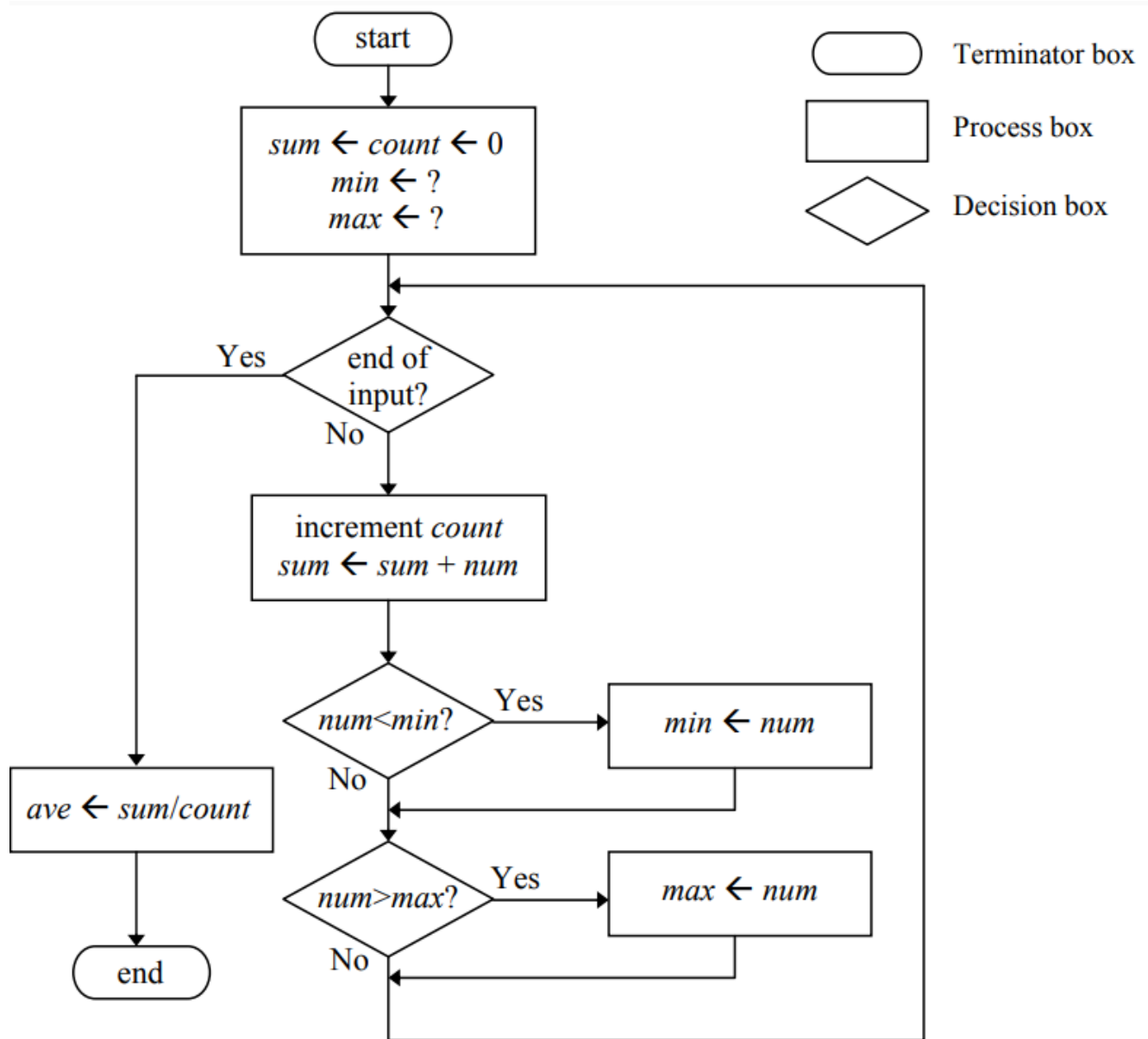
**Process Boxes:** Show the steps or actions that need to be performed.

**Decision Boxes:** Indicate points where a decision is required, leading to different branches based on conditions.

**Arrows:** Represent the flow of control, showing the direction in which the steps are executed



**Example:** An algorithm to find the minimum, maximum, and average of a list of number



Regardless of whether you use pseudo-code or flowcharts to represent your algorithm, it's essential to validate it by testing it with different sets of data. Carefully walk through the algorithm with various inputs to ensure it functions correctly and produces the desired outcomes in all scenarios.

## Properties of Algorithms

When evaluating and designing algorithms, several critical properties are considered to ensure their effectiveness and suitability for specific tasks. These properties help in assessing whether an algorithm will perform well in practice and how it will handle different scenarios. The table below summarizes essential properties explaining their significance and how they impact the performance and reliability of computational solutions

| Property            | Description                                                                                                                                                                                                                |
|---------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Correctness         | An algorithm is considered correct if it produces the right output for every possible input and terminates as expected. It must solve the problem it was designed to address.                                              |
| Partial Correctness | An algorithm is partially correct if, whenever it terminates, it provides the correct result for the given input. This means that if the algorithm completes, the result is accurate, but it may not always terminate.     |
| Total Correctness   | An algorithm is totally correct if it is both partially correct and guaranteed to terminate for all possible inputs. It ensures that the algorithm will always complete its execution and produce the correct output.      |
| Efficiency          | Efficiency measures how effectively an algorithm uses computational resources, such as time and memory. An efficient algorithm performs tasks using minimal resources, leading to faster execution and lower memory usage. |
| Big-O Notation      | Big-O Notation is a mathematical tool used to describe the upper bound of an algorithm's time complexity. It helps to understand how                                                                                       |

|             |                                                                                                                                                                                                                                |
|-------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|             | the runtime or space requirements of an algorithm grow relative to the size of the input.                                                                                                                                      |
| Scalability | Scalability refers to an algorithm's ability to handle increasing amounts of input data without a significant rise in computational resources. A scalable algorithm remains effective and efficient as the problem size grows. |
| Robustness  | Robustness is the ability of an algorithm to handle errors or unexpected inputs gracefully without failing. Robust algorithms are designed to be resilient and adaptable to various input scenarios and conditions.            |

Examples of Algorithms in Daily Computing Tasks

Algorithms play a critical role in various aspects of daily computing tasks, optimizing and facilitating a wide range of operations. From sorting and searching data to securing information and enhancing user experiences, algorithms are foundational to modern technology. Below is a table showcasing examples of algorithms used in different computing tasks, illustrating their functions and applications in everyday scenarios.

| Algorithm Category     | Algorithm                          | Description                                                                                                |
|------------------------|------------------------------------|------------------------------------------------------------------------------------------------------------|
| Sorting Algorithms     | QuickSort                          | Efficiently sorts large datasets by dividing and conquering.                                               |
|                        | MergeSort                          | A stable sort that divides the list into smaller segments, sorts them, and then merges them back together. |
| Search Algorithms      | Binary Search                      | Quickly finds items in a sorted list by repeatedly dividing the search interval in half.                   |
|                        | Linear Search                      | Finds an item by checking each element sequentially in an unsorted list.                                   |
| Compression Algorithms | ZIP                                | A widely used file compression format that reduces file size without loss of data.                         |
|                        | JPEG Compression                   | Reduces the file size of images by compressing image data with minimal loss of quality.                    |
| Encryption Algorithms  | AES (Advanced Encryption Standard) | A symmetric key encryption algorithm widely used to secure sensitive data.                                 |

|                        |                                |                                                                                                                                           |
|------------------------|--------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------|
|                        | RSA<br>(Rivest-Shamir-Adleman) | An asymmetric encryption algorithm used for secure data transmission.                                                                     |
| Search Engines         | PageRank Algorithm             | Developed by Google, it ranks web pages based on their link structure, helping users find relevant pages.                                 |
|                        | BM25                           | A probabilistic-based ranking function used to retrieve documents from a corpus based on query relevance.                                 |
| Recommendation Systems | Collaborative Filtering        | Makes recommendations based on user interactions and preferences, such as suggesting movies or products based on similar users' behavior. |
|                        | Content-Based Filtering        | Recommends items similar to those the user has liked before, based on the content of the items.                                           |

## Applications of Algorithms in Artificial Intelligence

Algorithms are fundamental to the field of Artificial Intelligence (AI), serving as the backbone of intelligent systems and applications. They are step-by-step procedures or formulas for solving problems, and their implementation in AI enables machines to perform tasks that typically require human intelligence. These tasks span various domains, from understanding and generating human language to interpreting visual data and making predictions based on historical patterns. Common Applications are summarized in the table below.

| Application                       | Key Algorithms                                                      | Impact                                                                                                             |
|-----------------------------------|---------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------|
| Natural Language Processing (NLP) | Transformer Models (e.g., BERT, GPT)                                | Powers chatbots, voice assistants, and translation services, enabling natural communication between humans and AI. |
| Computer Vision                   | Convolutional Neural Networks (CNNs), Object Detection (e.g., YOLO) | Used in facial recognition, autonomous vehicles, and medical imaging to improve safety and healthcare.             |
| Recommendation Systems            | Collaborative Filtering, Deep Learning-Based Recommendations        | Personalized user experiences on platforms like Netflix and Amazon, increasing engagement and satisfaction.        |
| Predictive Analytics              | Linear Regression, Time Series Forecasting                          | Helps predict trends and risks in finance and healthcare, enabling proactive decision-making.                      |
| Fraud Detection                   | Anomaly Detection, Logistic Regression                              | Protects against fraud in banking and online platforms by detecting suspicious activities and patterns.            |

## ACTIVITY 2: ONLINE DISCUSSIONS:ANAGRAM DETECTION ALGORITHM

Welcome to this week's **graded discussion forum**! This activity is an opportunity for you to engage with your peers, explore the concept of anagram detection, and develop critical thinking and problem-solving skills in a collaborative setting.

**Discussion Topic: Designing an Algorithm to Detect Anagrams**

**Objective:**

The purpose of this discussion is to design an algorithm that determines if two given words are anagrams. An anagram is when two words contain the exact same characters with the same frequency, but arranged differently. For example:

**Anagrams:** "LISTEN" and "SILENT", "NOW" and "WON"

**Not Anagrams:** "SELL" and "LESS"

---

**Instructions:****Initial Post:**

**Task:** Post an original response proposing your approach to solving the problem of detecting anagrams.

**Pseudocode:** Outline a step-by-step pseudocode for your solution.

**Algorithm Design:** Discuss your algorithm design, considering how you would compare two words.

**Time Complexity:** Analyze the time and space complexity of your proposed solution.

**Discussion Points:**

Should your algorithm handle **case sensitivity** (i.e., uppercase vs. lowercase)?

Should **special characters** (spaces, punctuation) be ignored, or should they be considered as part of the solution?

How will your algorithm handle **edge cases** like empty strings or very long words?

**Peer Interaction:**

**Task:** Provide constructive feedback on at least **two peers' posts**. Critically evaluate their pseudocode and algorithm design, suggesting possible improvements or alternative approaches.

**Alternative Methods:** If you see room for optimization (e.g., using frequency tables instead of sorting), share your ideas.

**Engagement:** Engage in meaningful discussions on topics such as case sensitivity, handling special characters, and optimizing the algorithm.

### **Final Reflection:**

**Task:** Reflect on the feedback you received from peers and make revisions to your initial post if needed. Briefly summarize any changes you made and why.

---

## **ACTIVITY 3: DESIGN AND IMPLEMENTATION OF AVERAGE CALCULATOR ALGORITHM**

### **Objective:**

In this graded lab activity, you are expected to create a flowchart, write pseudocode, and implement the algorithm in Python. This activity aims to assess your ability to develop basic algorithms and translate them into code.

### **Task: Implementing the Average Calculator Algorithm**

The Average Calculator Algorithm is a basic algorithm that calculates the average of a set of numbers. It involves summing the numbers and then dividing by the count of the numbers.

### **Instructions:**

#### **Flowchart Creation:**

**Tool:** Use a flowcharting tool (e.g., Lucidchart, Microsoft Visio) or draw by hand.



**Content:** Design a flowchart to represent the Average Calculator Algorithm. Include key steps such as inputting numbers, calculating the sum, dividing by the count, and outputting the average.

### **Pseudocode Development:**

**Format:** Write pseudocode in plain text.

**Content:** Develop pseudocode for the Average Calculator Algorithm, outlining the steps to input numbers, calculate the sum, compute the average, and display the result.

### **Python Implementation:**

**Code:** Write and test the Python code to implement the Average Calculator Algorithm.

**Content:** Ensure your code includes comments explaining each step of the algorithm.

### **Submission:**

You are required to submit the flowchart, pseudocode and python code on the LMS

---

Deliverables:

**Flowchart:** Submit a digital or scanned copy of your flowchart.

**Pseudocode:** Provide a text file or document containing your pseudocode.

**Python Code:** Submit your Python script file (.py) with clear comments.

Grading Criteria:

**Flowchart Accuracy (30%):** Correct representation of the Average Calculator Algorithm steps.

**Pseudocode Clarity (30%):** Clear, concise, and correct pseudocode.

**Python Implementation (40%):** Correct implementation, code quality, and comments.

---

## REFERENCE LIST AND FURTHER READING

Dimri, S. C., Malik, P., Ram, M. (2021). Algorithms: Design and Analysis.  
Germany: De Gruyter.

---

## ACTIVITY 4: END OF MODULE ASSESSMENT

**What defines a deterministic problem in computing?**

- A. A problem that may have different outputs for the same input.
- B. A problem where the output is not always predictable.
- C. A problem where a given input will always produce the same output.
- D. A problem with multiple correct answers.

**Which property ensures that an algorithm always produces the correct output if it terminates?**

- A. Total Correctness
- B. Partial Correctness
- C. Efficiency
- D. Scalability

**What does Big-O Notation describe in an algorithm?**

- A. The total correctness of the algorithm.

B. The time complexity and how it grows with input size.

C. The efficiency of an algorithm.

D. The robustness of an algorithm.

**Which of the following is a characteristic of a robust algorithm?**

A. Always terminates within finite steps.

B. Can handle errors or unexpected inputs gracefully.

C. Uses minimal computational resources.

D. Requires clear, well-defined steps.

**In the context of algorithmic problem-solving, what is 'Problem Formulation'?**

A. The process of implementing the algorithm in code.

B. Understanding the problem's requirements and constraints.

C. Designing a flowchart to visualize the algorithm.

D. Testing the algorithm with various inputs.

**What is the main goal of pseudocode?**

A. To describe an algorithm using specific programming language syntax.

B. To represent an algorithm in a high-level, language-independent format.

C. To visualize the algorithm using diagrams.

D. To write executable code directly.

**Which algorithmic property measures the resources required to solve a problem?**

A. Correctness

B. Efficiency

C. Scalability

D. Robustness

**In the Average Calculator Algorithm, what is the purpose of dividing the sum by the count of numbers?**

- A. To determine the maximum value in the list.
- B. To calculate the average value of the numbers.
- C. To find the median of the numbers.
- D. To sort the numbers in ascending order.

**What does 'Scalability' refer to in the context of an algorithm?**

- A. The algorithm's ability to handle errors.
- B. The algorithm's capacity to handle increasing input sizes efficiently.
- C. The algorithm's correctness for all inputs.
- D. The algorithm's effectiveness in solving problems.

**Which step in algorithmic problem-solving involves converting pseudocode into actual code?**

- A. Problem Understanding
- B. Problem Formulation
- C. Designing the Algorithm
- D. Implementation

# MODULE 2: DESIGNING AND ANALYZING ALGORITHMS

---

INSTRUCTIONAL HOURS: 4, 5 or 6

## MODULE OVERVIEW

In this module, students will explore the fundamental principles of designing and analyzing algorithms. They will learn various problem-solving strategies such as divide and conquer, greedy approaches, and dynamic programming, equipping them with the tools to tackle complex computational challenges. Additionally, the module will cover the analysis of algorithms, focusing on time complexity, space complexity, and the significance of Big O notation. Through a combination of theoretical concepts and practical applications, students will gain a comprehensive understanding of algorithm design and analysis, preparing them for advanced topics in computer science.

## LEARNING OUTCOMES

At the end of this module, you will be able to:

1. Understand and apply various problem-solving strategies, including divide and conquer, greedy algorithms, and dynamic programming.
2. Analyze algorithms for their time complexity and space complexity to evaluate their efficiency and resource usage.
3. Interpret and utilize Big O notation to describe the performance and scalability of algorithms.
4. Compare and contrast different algorithm design principles to select the most appropriate strategy for a given problem.

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Lab Activity on Design
3. Mastery Exercises
4. End of Module Quiz

---

## REVIEW READING MATERIAL

### INTRODUCTION TO DESIGN AND ANALYSIS OF ALGORITHMS

#### What is the difference between algorithm Design and Analysis ?

**Definition and focus of Algorithm Design:** Algorithm design is the process of defining a systematic solution to a problem. It involves creating a step-by-step procedure that can be followed to solve a problem or accomplish a task. The focus is on how the algorithm works, what data structures are used, and what approach (e.g., brute force, divide and conquer, dynamic programming) will be used.

**Definition and focus of Algorithm Analysis:** Algorithm analysis involves evaluating the performance and efficiency of an algorithm. It helps to understand how the algorithm scales with the input size in terms of time and space resources. The focus is on the computational complexity of the algorithm, which includes time complexity (how fast the algorithm runs) and space complexity (how much memory the algorithm uses).

#### Design principles

On Design Principles , you will learn strategies to help you design effective algorithms. Think of these strategies as tools you can use to solve complex problems by breaking them down into smaller, manageable parts. The Three common problem-solving strategies include: Divide and Conquer, Greedy Algorithms, and Dynamic Programming

#### Divide and Conquer

Imagine you have a huge task at hand—like cleaning your entire house. Instead of trying to tackle it all at once, you might clean room by room, or even break each room into smaller tasks like vacuuming, dusting, or organizing. This is the essence of **divide and conquer**. You break a problem into smaller sub-problems, solve each one individually, and then combine those solutions.

### **How does the strategy work?**

Divide the problem into smaller, similar sub-problems.

Conquer each sub-problem recursively (solving it in the same way you broke it down).

Combine the solutions of the sub-problems to get the final solution.

### **Real-World Example:**

Imagine organizing a set of student report cards. You divide the cards into smaller groups by class, sort each group by student name, and then combine all the sorted groups to have a complete list of report cards sorted by name.

### **When to Use:**

When the problem can naturally be broken into similar smaller problems, or when working with large datasets where solving smaller pieces is more efficient.

## **Greedy Algorithms**

Greedy algorithms focus on making the best possible choice at each step, with the aim that these local solutions will lead to the best overall solution. Imagine you're shopping on a tight budget. At each shop, you pick the cheapest option, hoping that this will save you the most money in the end.

### **How it Works:**

**Choose:** Select the best option available at the current step.

**Move Forward:** Make decisions without looking back or considering future choices.

**Hope:** Trust that these local choices will collectively lead to the optimal overall outcome.

### **When to Use:**

Greedy algorithms work well when making a series of local optimal choices results in a global optimum. They are suitable for problems like:

Scheduling: Assigning tasks to time slots in a way that maximizes efficiency.

Minimum Spanning Trees: Connecting all points in a graph with the minimal total edge weight.

Routing: Finding the shortest path in a graph or Network. In this case Dijkstra's Algorithm can be applied

## Dynamic Programming

Imagine you're trying to climb a staircase, and at each step, you can choose to take one or two steps. If you want to know how many different ways there are to reach the top, one approach is to solve this problem for every possible step combination. However, dynamic programming (DP) allows you to store the results of each calculation so you don't repeat the same work multiple times.

### How it Works:

Break the problem into overlapping sub-problems.

Solve each sub-problem only once and store the result.

Use these stored results to avoid recalculating the same problem multiple times.

### Real-World Example:

The **Fibonacci sequence** can be solved using dynamic programming by storing the result of each Fibonacci number, so you don't have to recompute them.

### When to Use:

Use DP when a problem has overlapping sub-problems and optimal substructure (i.e., the solution to a problem can be built from the solutions of smaller problems). It's commonly used for problems like the knapsack problem, sequence alignment, or shortest path problems.



What is the main idea behind the Divide and Conquer strategy?

- A) Solving the problem in a single step.
- B) **Breaking the problem into smaller sub-problems, solving each one, and combining their solutions.**
- C) Making the best choice at each step without considering future steps.
- D) Directly combining all possible solutions to find the best one.

2. Which of the following is a step in the Divide and Conquer approach?

- A) Only solve the problem without dividing it.
- B) Solve all sub-problems at once.
- C) **Divide the problem into smaller sub-problems, solve each sub-problem recursively, and combine the solutions.**
- D) Focus on the most complex part of the problem first.

3. When is the Greedy Algorithm approach most appropriate?

- A) When a problem can be solved by breaking it into smaller sub-problems.
- B) When local optimal choices do not necessarily lead to a global optimum.
- C) **When making the best choice at each step will lead to the best overall solution.**
- D) When the problem requires dividing and merging solutions.

4. What is an example of a Greedy Algorithm?

A) Merge Sort

B) Quick Sort

**C) Dijkstra's Algorithm for finding the shortest path in a graph**

D) Binary Search

5. In the context of a Greedy Algorithm, what does "Move Forward" mean?

A) Revisit previous decisions to adjust them.

B) Consider future choices before making a decision.

**C) Make decisions based solely on the current best option without looking back.**

D) Combine all possible choices to determine the best option.

6. Which principle does the following scenario illustrate? You are tasked with organizing a large dataset by breaking it into smaller, more manageable chunks, processing each chunk, and then combining the results.

A) Greedy Algorithms

**B) Divide and Conquer**

C) Dynamic Programming

D) Brute Force

7. In a Divide and Conquer algorithm, what is the purpose of the "Combine" step?

A) To solve each sub-problem independently.

**B) To merge the solutions of all sub-problems into the final solution.**

C) To divide the problem into smaller sub-problems.

D) To make the best choice at each step.

## Analysis of algorithms

Understanding how to analyze algorithms is crucial for determining their efficiency and performance. In the analysis of algorithms you are expected to know key concepts like time complexity, space complexity, and Big O notation.

### Time Complexity

Time complexity measures how the running time of an algorithm changes with the size of the input. It helps to estimate the amount of time an algorithm takes to complete based on the input size. The following are the key concepts in time complexity analysis:

**Growth Rate of an Algorithm:** Time complexity is often expressed as a function of the input size  $n$ . For example, if an algorithm takes  $2n+3$  steps for input size  $n$ , its time complexity is  $O(n)$ .

#### Common Time Complexities:

| Time Complexity | Description      | Growth Rate                                                       |
|-----------------|------------------|-------------------------------------------------------------------|
| $O(1)$          | Constant Time    | The running time remains constant regardless of input size.       |
| $O(n)$          | Linear Time      | The running time increases linearly with the size of the input.   |
| $O(n^2)$        | Quadratic Time   | The running time increases quadratically as the input size grows. |
| $O(\log n)$     | Logarithmic Time | The running time increases logarithmically with the input size.   |

**Example:** If you have an algorithm that searches for an item in a list by checking each item one by one, its time complexity is  $O(n)$  because the time it takes grows linearly with the size of the list.

### Space Complexity

Space complexity measures how the amount of memory an algorithm needs changes with the size of the input. It helps to estimate the amount of memory required for an algorithm to run.

Key concepts in space complexity include:

**Memory Usage:** space complexity measures how much memory an algorithm uses. This includes both the memory required for the algorithm's variables (like temporary data) and any extra space needed for input and output data.

### Common Space Complexities:

| Space Complexity | Description                                                         | Memory Usage Example                                       |
|------------------|---------------------------------------------------------------------|------------------------------------------------------------|
| $O(1)$           | Constant Space – Memory usage does not change with input size.      | Using a fixed number of variables regardless of input size |
| $O(n)$           | Linear Space – Memory usage grows linearly with input size.         | Storing a list of size $n$                                 |
| $O(n^2)$         | Quadratic Space – Memory usage grows quadratically with input size. | Storing a 2D matrix of size $n \times n$                   |

**Example:** If an algorithm needs to store a list of size  $n$ , its space complexity is  $O(n)$  because the memory required grows linearly with the size of the input.

### Big O Notation and Its Significance

Big **O** notation is a mathematical notation used to describe the **upper bound of an algorithm's time and space complexity**. It provides a way to express the **worst-case scenario** of an algorithm's performance, helping to compare the efficiency of different algorithms.

Big O notation expresses how the **time or space complexity** of an algorithm grows relative to the input size  $n$ . For example,  $O(n)$  indicates that the complexity grows linearly with input size.

The Significance of Big **O** notation is that:

It allows for comparison between different algorithms based on their efficiency,

It helps to determine how well an algorithm will scale with larger inputs.

Provides insight into the maximum amount of time or space an algorithm will require.

**Example:** If you compare two sorting algorithms, one with  **$O(n \log n)$**  time complexity and another with  **$O(n^2)$** , the first one is generally more efficient for large input sizes because its running time grows slower relative to the size of the input.

### Example 1: Time Complexity Analysis

**Algorithm:** Finding the maximum value in an unsorted array.

```
def find_max(arr):  
    max_value = arr[0]  
    for i in range(1, len(arr)):  
        if arr[i] > max_value:  
            max_value = arr[i]  
    return max_value
```

## Time Complexity Analysis:

**Identify Basic Operations:** The main operation is the comparison inside the loop.

**Count Basic Operations:** The loop runs  $n-1$  times (where  $n$  is the number of elements in the array). Each iteration performs a constant-time comparison and assignment.

**Big O Notation:** The time complexity is  $O(n)$  because the running time grows linearly with the size of the input array.

**Explanation:** The function iterates through the array once, making the time complexity linear,  $O(n)$

## Example 2: Time Complexity Analysis

**Algorithm:** Printing all pairs of elements from an array.

```
def print_all_pairs(arr):  
    for i in range(len(arr)):  
        for j in range(len(arr)):  
            print(arr[i], arr[j])
```

## Time Complexity Analysis:

**Identify Basic Operations:** The basic operation is printing each pair of elements.

**Count Basic Operations:** There are two nested loops, each running  $n$  times. Thus, there are  $n \times n$  (or  $n^2$ ) print operations.

**Big O Notation:** The time complexity is  $O(n^2)$  because the number of operations grows quadratically with the size of the input array.

**Explanation:** Each element is paired with every other element, resulting in a quadratic growth in the number of operations.

### Example 3: Space Complexity Analysis

**Algorithm:** Creating a list of size  $n$  and populating it with integers from  $0$  to  $n-1$

```
def create_list(n):  
    lst = []  
    for i in range(n):  
        lst.append(i)  
    return lst
```

### Space Complexity Analysis:

**Identify Memory Usage:** The algorithm uses additional space to store the list `lst`.

**Count Memory Usage:** The size of `lst` grows linearly with  $n$ . No additional significant space is used beyond `lst`.

**Big O Notation:** The space complexity is  $O(n)$  because the amount of memory required grows linearly with the size of the input.

**Explanation:** The space used by the list grows directly with the number of elements  $n$ , resulting in linear space complexity,  $O(n)$ .

## Activity 3: Mastery Exercise on Analysis of Algorithms

### Objective

To demonstrate understanding of time complexity, space complexity, and Big O notation by analyzing different algorithms and their efficiency.

### Instructions

Complete the following tasks and questions to test your mastery of algorithm analysis concepts. Provide explanations and justifications for your answer.

#### Task 1: Analyzing Time Complexity

Analyze the following code snippet and determine its time complexity. Explain your reasoning.

```
def find_max(arr):  
  
    max_value = arr[0]  
  
    for i in range(1, len(arr)):  
  
        if arr[i] > max_value:  
  
            max_value = arr[i]  
  
    return max_value
```

**Task 2:** Analyze the following python code snippet and determine its space complexity. Explain your reasoning.

```
def create_list(n):  
  
    lst = []
```



```

for i in range(n):

    lst.append(i)

return lst

```

**Task 3:** Match the following algorithms with their time complexities. Explain why each algorithm has the given time complexity.

| Algorithm             | Time Complexity | Explanation |
|-----------------------|-----------------|-------------|
| <b>Binary Search</b>  |                 |             |
| <b>Bubble Sort</b>    |                 |             |
| <b>Merge Sort</b>     |                 |             |
| <b>Insertion Sort</b> |                 |             |

#### REFERENCE LIST AND FURTHER READING

1. Dimri, S. C., Malik, P., Ram, M. (2021). Algorithms: Design and Analysis. Germany: De Gruyter.
2. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>

#### Activity 4: End of Module Quiz

##### Question 1

Which of the following scenarios is the best fit for using the Divide and Conquer strategy?

A) Finding the shortest path between two points in a weighted graph.

**B) Sorting a large array by dividing it into smaller sub-arrays, sorting each sub-array, and then merging them.**

C) Selecting the maximum value in an unsorted array.

D) Scheduling tasks to maximize resource usage.

#### Question 2

Consider an algorithm that performs a binary search on a sorted array of size  **$n$** . What is its time complexity in Big O notation?

A)  $O(n)$  B)  $O(n \log n)$  C)  **$O(\log n)$**  D)  $O(n^2)$

#### Question 3:

For a recursive algorithm that uses  **$O(n)$**  space to store function call frames and  **$O(1)$**  space for auxiliary variables, what is the overall space complexity?

A)  $O(1)$  B)  **$O(n)$**  C)  $O(n \log n)$  (D)  $O(n^2)$

#### Question 4

Which of the following statements correctly describes the significance of Big O notation?

A) It provides an exact measurement of an algorithm's running time or memory usage.

**B) It estimates the worst-case upper bound of an algorithm's running time or space requirement.**

C) It compares the performance of different algorithms in terms of their average case efficiency.

D) It measures the average running time of an algorithm for all possible input sizes.

#### Question 5

In which of the following situations is a Greedy algorithm most likely to fail in providing the optimal solution?

- A) Finding the minimum spanning tree of a graph.
- B) Solving the Knapsack problem where items have fractional values.
- C) Scheduling jobs with deadlines to minimize total completion time.
- D) Determining the optimal way to allocate resources with constraints.**

## MODULE 3: GROWTH OF FUNCTIONS AND RECURRENCES

INSTRUCTIONAL HOURS: 4, 5 or 6

---

### MODULE OVERVIEW

In this module, students will delve into the concepts of growth of functions and recurrences, essential for understanding algorithm efficiency and performance analysis. The module will begin with mathematical notations for growth rates, allowing students to compare different growth rates effectively. This foundational knowledge will facilitate discussions on how various algorithms behave as input sizes increase. The module will also cover the techniques for solving recurrences, with a particular focus on the Master Theorem. Through practical examples and theoretical insights, students will develop the skills to analyze and predict the behavior of algorithms based on their growth functions and recurrences.

---

## LEARNING OUTCOMES

At the end of this module, you will be able to:

1. Understand and utilize mathematical notations to represent and compare growth rates of functions.
2. Analyze and categorize different growth rates to evaluate their impact on algorithm performance.
3. Solve recurrences using various methods, including iteration and substitution, to determine algorithm behavior.
4. Apply the Master Theorem to efficiently solve common classes of recurrences encountered in algorithm analysis.

---

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Watch YOUTUBE Video
3. Lab Activity
4. Mastery Exercises
5. End of Module Quiz

## Growth of Functions

In algorithm analysis, we use several mathematical notations to describe how the running time or space requirement of an algorithm grows as the input size increases. The most common notations are Big O, Big  $\Theta$  (Theta), and Big  $\Omega$  (Omega). These notations help us classify algorithms based on their efficiency and scalability.

### Big O Notation (O)

The **Big O** notation describes the upper bound of the growth rate of a function. It provides a worst-case scenario for the growth of the function. It is denoted as  **$O(f(n))$** .

**Example:** If an algorithm's running time is  **$O(n^2)$** , it means that the running time grows no faster than a quadratic function of the input size  $n$ .

### Big $\Theta$ Notation ( $\Theta$ )

Big  $\Theta$  notation describes both the upper and lower bounds of the growth rate of a function. It provides a tight bound on the function's growth. It is denoted as  **$\Theta(f(n))$** .

**Example:** If an algorithm's running time is  **$\Theta(n \log n)$** , it means that the running time grows asymptotically at the same rate as  **$n \log n$** .

### Big $\Omega$ Notation ( $\Omega$ )

Big  $\Omega$  notation describes the lower bound of the growth rate of a function. It provides a best-case scenario for the growth of the function. It is denoted as  **$\Omega(f(n))$** .

**Example:** If an algorithm's running time is  **$\Omega(n)$** , it means that the running time grows at least linearly with the input size  $n$ .

## Comparing Different Growth Rates

To understand the efficiency of algorithms, it's essential to compare their growth rates. Here are common growth rates and how they compare:

| Growth Rate  | Big O Notation | Description                                                | Example                                               |
|--------------|----------------|------------------------------------------------------------|-------------------------------------------------------|
| Constant     | $O(1)$         | The running time or space does not change with input size. | Accessing an element in an array.                     |
| Logarithmic  | $O(\log n)$    | The running time or space grows logarithmically.           | Binary search in a sorted array.                      |
| Linear       | $O(n)$         | The running time or space grows linearly with input size.  | Finding the maximum in an array.                      |
| Linearithmic | $O(n \log n)$  | The running time or space grows as $n \log n$ .            | Merge sort and quicksort algorithms.                  |
| Quadratic    | $O(n^2)$       | The running time or space grows quadratically.             | Bubble sort and insertion sort.                       |
| Cubic        | $O(n^3)$       | The running time or space grows cubically.                 | Some naive 3D matrix algorithms.                      |
| Exponential  | $O(2^n)$       | The running time or space grows exponentially.             | Recursive solution to the traveling salesman problem. |
| Factorial    | $O(n!)$        | The running time or space grows factorially.               | Generating all permutations of a set.                 |

## Comparing Growth Rates

**Constant vs. Logarithmic:** Constant time is faster than logarithmic time for large input sizes.

**Logarithmic vs. Linear:** Logarithmic time is more efficient than linear time.

**Linear vs. Linearithmic:** Linearithmic time (e.g.,  $O(n \log n)$ ) grows faster than linear time but is more efficient than quadratic time.

**Quadratic vs. Cubic:** Quadratic time grows slower than cubic time, which means algorithms with  $O(n^2)$  complexity are more scalable than those with  $O(n^3)$ .

**Exponential vs. Factorial:** Exponential time grows faster than factorial time. However, both are impractical for large input sizes due to their rapid growth.

**Example:** Suppose we have the following values of  $n$ : 1, 10, 100, 1000, and 10,000. The following table shows the corresponding values for each function.

| $n$    | $O(1)$ | $O(\log n)$ | $O(n)$ | $O(n \log n)$ | $O(n^2)$    | $O(n^3)$  | $O(2^n)$ | $O(n!)$   |
|--------|--------|-------------|--------|---------------|-------------|-----------|----------|-----------|
| 1      | 1      | 0           | 1      | 0             | 1           | 1         | 2        | 1         |
| 10     | 1      | 2           | 10     | 20            | 100         | 1,000     | 1,024    | 3,628,800 |
| 100    | 1      | 6           | 100    | 600           | 10,000      | 1,000,000 | ~        | ~         |
| 1,000  | 1      | 9           | 1,000  | 9,000         | 1,000,000   | 1 billion | ~        | ~         |
| 10,000 | 1      | 13          | 10,000 | 130,000       | 100 million | ~         | ~        | ~         |

#### Observations:

**Factorial time  $O(n!)$ :** This is the most extreme case. Even for relatively small inputs like  $n=10$ , the function value reaches over 3.6 million. For larger values like  $n=100$  or  $n=1000$ , the function becomes astronomically large and impractical for computation.

## ACTIVITY 2: PRACTICAL LAB ACTIVITY

### Objective:

In this lab, you will:

- Implement two algorithms: Linear Search (with linear time complexity) and Bubble Sort (with quadratic time complexity) using python Programming Language
- Compare their performance on datasets of different sizes.
- Analyze the relationship between input size and algorithm running time.

### Task Overview:

You are required to Implement two algorithms, one with linear time complexity (e.g., Linear Search) and one with quadratic time complexity (e.g., Bubble Sort), and compare their performance on datasets of increasing sizes.

### Steps

Step 1: Write pseudocode for each algorithm.

Step 2: Implement the algorithms in Python.

Step 3: Run the algorithms on input sizes of 10, 100, 1000, and 10,000.

Step 4: Measure and record the time taken by each algorithm.

Step 5: Analyze how the running time increases as the input size grows.



Step6: Upload the documentation on LMS

### ACTIVITY 3: MASTERY QUIZ

Consider the algorithm whose pseudocode is provided below, and analyze the time complexity using **Big O**, **Big  $\Theta$** , and **Big  $\Omega$**  notations.

```
for k = 2 to B.length ----- 1
    temp = B[k]-----2
    // Insert B[k] into sorted array B[1.....k-1]----- 3
    l = k - 1 -----4
    while l > 0 and B[l] > temp-----5
        B[l + 1] = B[l]-----6
        l = l - 1 -----7
    B[l + 1] = temp-----8
```

What is the time complexity of the Insertion Sort algorithm in terms of **Big O**, **Big  $\Theta$** , and **Big  $\Omega$**  notations?

- a) **Big O:**  $O(n)$ , **Theta:**  $\Theta(n)$ , **Omega:**  $\Omega(n)$
- b) **Big O:**  $O(n^2)$ , **Theta:**  $\Theta(n^2)$ , **Omega:**  $\Omega(n)$
- c) **Big O:**  $O(n^2)$ , **Theta:**  $\Theta(n^2)$ , **Omega:**  $\Omega(n \log n)$
- d) **Big O:**  $O(n \log n)$ , **Theta:**  $\Theta(n \log n)$ , **Omega:**  $\Omega(n \log n)$

**Explanation:**

**Big O Notation ( $O$ ):** Represents the worst-case time complexity. Since the pseudocode describes the Insertion Sort algorithm, the worst case is when the array is sorted in reverse, leading to  $O(n^2)$ .

**Theta Notation ( $\Theta$ ):** Represents the average-case time complexity. Insertion Sort typically has a quadratic average-case complexity, so  $\Theta(n^2)$  is correct for average cases.

**Omega Notation ( $\Omega$ ):** Represents the best-case time complexity. In the best case, when the array is already sorted, Insertion Sort performs  $O(n)$  operations. Thus,  $\Omega(n)$  is appropriate for the best case.

For reference and further reading on asymptotic notations, you can explore the following link: [Asymptotic Notations](#). The page also includes additional relevant examples

## Recurrences

### What is a recurrence Relation?

A recurrence relation is a mathematical formula that recursively defines a sequence. In the context of algorithm analysis, it often expresses the time complexity of recursive algorithms. For example:

- Linear Recurrence Relation:

$$T(n) = T(n - 1) + c$$

- Divide and Conquer Recurrence:

$$T(n) = aT\left(\frac{n}{b}\right) + f(n)$$

Where:

- $T(n)$  is the function describing the time complexity.
- $a$  is the number of subproblems in the recursion.
- $\frac{n}{b}$  is the size of each subproblem.
- $f(n)$  is the cost outside the recursive calls.

## Solving Recurrences

There are several methods to solve recurrences, including:

### Substitution Method:

Guess the form of the solution and use mathematical induction to prove it.

**Example:** Suppose

$$T(n) = 2T(n/2) + n$$

Guess that  $T(n) = O(n \log n)$ . Use induction to prove this guess.

### Iteration Method:

Expand the recurrence until a pattern is observed, then sum up the results.

**Example:** To solve

$$T(n) = 2T\left(\frac{n}{2}\right) + n$$

expand it:

$$T(n) = 2\left(2T\left(\frac{n}{4}\right) + \frac{n}{2}\right) + n = 4T\left(\frac{n}{4}\right) + 2n + n$$

Repeat this process to observe the pattern.

### Characteristic Equation Method:

Used for solving linear recurrences with constant coefficients.

**Example:** For

$$T(n) = 2T(n - 1) + 1$$

Find the characteristic equation and solve for its roots.

### The Master Theorem

The Master Theorem provides a straightforward way to analyze the time complexity of divide-and-conquer algorithms. It applies to recurrences of the form:

$$T(n) = aT(n/b) + f(n)$$

Where:

**a** is the number of subproblems.

**b** is the factor by which the problem size is divided.

**f(n)** is the cost outside the recursive calls.

#### Master Theorem Cases:

1. **Case 1:** If  $f(n) = O(n^c)$  where  $c < \log_b a$ , then

$$T(n) = \Theta(n^{\log_b a})$$

- **Example:**  $T(n) = 2T\left(\frac{n}{2}\right) + O(1)$ . Here,  $a = 2$ ,  $b = 2$ , and  $f(n) = O(1)$ . Since  $f(n)$  is polynomially smaller than  $n^{\log_b a}$ ,  $T(n) = \Theta(n^{\log_2 2}) = \Theta(n)$ .

2. **Case 2:** If  $f(n) = \Theta(n^c)$  where  $c = \log_b a$ , then

$$T(n) = \Theta(n^c \log^k n)$$

where  $k \geq 0$  depends on the specific form of  $f(n)$ .

- **Example:**  $T(n) = 2T\left(\frac{n}{2}\right) + O(n)$ . Here,  $a = 2$ ,  $b = 2$ , and  $f(n) = O(n)$ . Since  $f(n)$  matches  $n^{\log_b a}$ ,  $T(n) = \Theta(n \log n)$ .

3. **Case 3:** If  $f(n) = \Omega(n^c)$  where  $c > \log_b a$  and regularity condition holds (i.e.,  $af\left(\frac{n}{b}\right) \leq kf(n)$  for some  $k < 1$  and sufficiently large  $n$ ), then

$$T(n) = \Theta(f(n))$$

- **Example:**  $T(n) = 2T\left(\frac{n}{2}\right) + O(n^2)$ . Here,  $a = 2$ ,  $b = 2$ , and  $f(n) = O(n^2)$ . Since  $f(n)$  grows faster than  $n^{\log_b a}$ ,  $T(n) = \Theta(n^2)$ .

For further reading on Recursion and Recurrences refer to Unit 1 on this link [Recursion](#).

---

#### ACTIVITY 3: YouTube Video on Recurrence Relations [\(YOUTUBE LINK\)](#)

##### Objective:

To understand how to form recurrence relations by watching a video and then applying this knowledge to create your own recurrence relations based on different scenarios.

##### Instructions:

### Watch the Video:

Watch the YouTube video titled "Recurrence Relations" by **Satpal Singh** . The video explains how to form recurrence relations.

**Video Link:** [Recurrence Relations](#)

### Complete the Following Activities:

#### a. Summarize Key Points:

Write a brief summary (150-200 words) of the key points discussed in the video. Focus on the process of forming recurrence relations and any examples provided.

**b. Form Recurrence Relations:** Based on the video, create recurrence relations for the following scenarios:

Scenario 1: An algorithm that divides an input of size  $n$  into 2 subproblems, each of size  $n/3$ , and performs an additional  $O(n)$  operation. Write the recurrence relation for the time complexity.

Scenario 2: A recursive function that solves a problem of size  $n$  by breaking it into 3 subproblems of size  $n/2$  each, with an additional constant time operation. Write the recurrence relation for the time complexity.

Scenario 3: A problem where each step of the recursion reduces the problem size by 1, with a cost of  $O(n)$  at each step. Write the recurrence relation for this problem.

### Submit Your Work:

Submit your summary and formulated recurrence relations to the provided email or learning management system (LMS) by **[Submission Deadline]**.

#### ACTIVITY 4 : MASTERY QUIZ ON RECURSION AND RECURRENCES

Which method is best suited for solving the recurrence  $T(n) = T(n - 1) + n$ ? Explain Why? **Answer Iteration method**

Consider the recurrence  $T(n) = 3T(n/3) + n$ . According to the Master Theorem, what is the time complexity? **Answer  $O(n \log n)$**

What is the general solution of the recurrence relation  $T(n) = 2T(n - 1) + 1$  using the characteristic equation method? **Answer  $T(n) = O(2^n)$**

**Using the substitution method, solve the recurrence  $T(n) = 4T(n/2) + n$ . What is the time complexity? Answer  $O(n^2)$**

#### ACTIVITY 5 : END OF MODULE QUIZ

##### Question 1

Which of the following Big O notations correctly describes the growth of the function  $f(n) = 3n^2 + 5n + 2$ ?

A)  $O(n^3)$

**B)  $O(n^2)$**

C)  $O(n \log n)$

D)  $O(n)$

##### Question 2

The recurrence relation  $T(n) = 2T(n/2) + n$  can be solved using the Master Theorem. What is the time complexity of this recurrence?

A)  $O(n^3)$

B)  $O(n^2)$

C)  $O(n \log n)$

D)  $O(n)$

### Question 3

Using the substitution method, solve the recurrence  $T(n) = 3T(n/3) + n$ . What is the time complexity?

A)  $O(n^3)$

B)  $O(n^2)$

C)  $O(n \log n)$

D)  $O(n)$

### Question 4

Consider the recurrence  $T(n) = T(n - 1) + n$ . By expanding the recurrence using the iteration method, what is the time complexity?

A)  $O(n^3)$

B)  $O(n^2)$

C)  $O(n \log n)$

D)  $O(n)$

---

### REFERENCE LIST AND FURTHER READING

1. Dimri, S. C., Malik, P., Ram, M. (2021). Algorithms: Design and Analysis. Germany: De Gruyter.

2. Khan Academy. (n.d.). Algorithms.

<https://www.khanacademy.org/computing/computer-science/algorithms>



3. Thomas, P., & Hunt, S. (2020). *An introduction to algorithms and data structures*. University of London.

## MODULE 4: SORTING AND ORDER STATISTICS

---

INSTRUCTIONAL HOURS: 4, 5 or 6

### MODULE OVERVIEW

In this module, you will explore efficient sorting algorithms that operate in linear time and learn how to quickly find the median and other ranked elements in data using order statistics techniques.

### LEARNING OUTCOMES

By the end of this module, you will be able to:

1. Implement linear-time sorting algorithms like **Counting Sort**, **Radix Sort**, and **Bucket Sort**.
2. Use the **Quickselect** algorithm to efficiently find the k-th smallest or largest element.
3. Compute the **median** without fully sorting the data.
4. Analyze and compare the efficiency of linear-time sorting methods against comparison-based algorithms.

---

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Lab Activity
3. Complete End of Module Quiz

### Sorting in Linear Time

Why is sorting essential in the design and analysis of algorithms?

Sorting is crucial in algorithm design and analysis for several reasons:

**Improves Efficiency:** Sorting helps solve problems faster. For example, binary search works only on sorted data, reducing the search time from  $O(n)$  in an unsorted array to  $O(\log n)$  in a sorted one. Sorting is often used as a pre-processing step to make algorithms more efficient.

**Simplifies Problem Solving:** Problems like finding the median, mode, or duplicates become easier when data is sorted. Sorting allows us to quickly access values like minimum, maximum, or closest numbers.

**Organizes Data:** Sorting arranges data in a structured way, making it easier to use in applications such as database indexing, file organization, and data retrieval.

**Supports Other Algorithms:** Many advanced algorithms, like those for searching, clustering, or scheduling, rely on sorting. Algorithms like merge sort and quicksort are key examples in divide-and-conquer strategies.

**Helps Analyze Efficiency:** Sorting algorithms are often used to study time and space complexity. Comparing algorithms like bubble sort, merge sort, and quicksort helps us understand different techniques and trade-offs.

**Enables Efficient Comparisons:** Once data is sorted, tasks like finding duplicates or comparing datasets can be done more efficiently, often in  $O(n)$  time, improving overall algorithm performance.

## Linear Time Sorting Algorithms

### What is a linear-time sorting algorithm, and why is it important?

A **linear-time sorting algorithm** sorts data in  **$O(n)$**  time, meaning the time it takes grows directly with the number of elements ( $n$ ). These algorithms are important because they can handle large datasets efficiently, unlike comparison-based algorithms that take longer ( $O(n \log n)$  or  $O(n^2)$ ).

Linear-time sorting works by using special features of the data, like knowing the range of values or having a specific structure, to sort faster.

Examples include **Counting Sort, Radix Sort, and Bucket Sort**.

### Counting Sort

#### How does **Counting Sort** work?

Counting Sort is an algorithm that works by **counting the occurrences** of each unique element in the input array. It is ideal for sorting integers or objects that can be mapped to integers within a fixed range. After counting, it uses the count to position the elements in the sorted output array.

#### Steps:

Create an auxiliary array to store the count of each element.

Modify the count array so that each index holds the position of the element in the output array.

Traverse the input array and place each element in its correct sorted position.

Counting Sort runs in  **$O(n + k)$**  time, where  $n$  is the number of elements and  $k$  is the range of the input (difference between the max and min values).

#### In what scenarios is Counting Sort most effective?

Counting Sort is most effective when:

The range of input values is relatively small compared to the number of elements (i.e., when  $k$  is smaller than  $n$ ).

it is inefficient for data with a large range, as the auxiliary array size increases with the range of input values.

## Radix Sort

How does Radix Sort differ from Counting Sort?

Radix Sort processes numbers digit by digit, starting from the least significant digit (LSD) or the most significant digit (MSD), depending on the variant used. It uses **Counting Sort** (or another stable sorting algorithm) as a subroutine to sort the numbers at each digit level.

### Steps for LSD Radix Sort:

Start from the least significant digit.

Use Counting Sort to sort the numbers based on the current digit.

Move to the next digit and repeat until all digits are processed.

Radix Sort runs in  $O(d * (n + k))$  time, where  $d$  is the number of digits in the largest number,  $n$  is the number of elements, and  $k$  is the range of each digit (typically 0-9 for decimal numbers). It is efficient for sorting large sets of integers or strings.

What are the strengths of Radix Sort?

Radix Sort is advantageous because:

It can handle large datasets efficiently, especially for numbers with a limited number of digits (e.g., phone numbers, zip codes).

It avoids the overhead of comparison-based sorting algorithms.

It is stable, maintaining the relative order of elements with the same key.

However, it requires extra space for counting arrays and may not perform well for very large digit ranges.

## Bucket Sort

### What is Bucket Sort?

Bucket Sort works by **distributing elements into buckets** based on their value. Each bucket is then sorted individually, either using a different sorting algorithm or recursively applying Bucket Sort.

#### Steps:

Divide the input range into several buckets.

Distribute the input elements into these buckets.

Sort the elements within each bucket using a stable sorting algorithm (such as Insertion Sort).

Concatenate the sorted buckets to form the final sorted array.

Bucket Sort runs in  **$O(n + k)$**  time, where  $n$  is the number of elements and  $k$  is the number of buckets. It works best when the input is uniformly distributed over a known range.

### When is Bucket Sort most efficient?

Bucket Sort is most efficient when:

The input data is uniformly distributed.

The number of buckets is similar to the number of elements, ensuring balanced distribution.

It is less effective for skewed data distributions, where most elements fall into a small number of buckets, causing some buckets to be overfilled and requiring additional sorting.

## Comparison of Linear-Time Sorting Algorithms

How do Counting Sort, Radix Sort, and Bucket Sort compare?

| Algorithm     | Time Complexity  | Space Complexity | Best Use Case                                    |
|---------------|------------------|------------------|--------------------------------------------------|
| Counting Sort | $O(n + k)$       | $O(k)$           | Small integer ranges                             |
| Radix Sort    | $O(d * (n + k))$ | $O(n + k)$       | Large numbers or strings with a fixed digit size |
| Bucket Sort   | $O(n + k)$       | $O(n)$           | Uniformly distributed data                       |

## Practical Examples For Linear Time Algorithms

### Counting Sort

**Scenario:** Sorting exam scores

**Example:** You have a list of exam scores for 10 students:

[4, 2, 2, 8, 3, 3, 1, 5, 6, 7]

**Steps:**

**Count occurrences:** Create a count array to track how many times each score appears.

**Build the output array:** Use the count array to determine the position of each score in the sorted array.

**Sorted Output:**

[1, 2, 2, 3, 3, 4, 5, 6, 7, 8]

## 2. Radix Sort

**Scenario:** Sorting phone numbers

**Example:** You have a list of phone numbers:

[415, 212, 303, 123, 556]

**Steps:**

**Sort by least significant digit (LSD):** Use Counting Sort for the last digit, resulting in:

Sorted by last digit: [212, 303, 123, 415, 556]

**Sort by next digit:** Repeat the process for the second digit. After sorting by the second digit, you might get:

Sorted by second digit: [123, 212, 303, 415, 556]

**Continue sorting by more significant digits** until all digits are processed.

**Sorted Output:**

[123, 212, 303, 415, 556]

## 3. Bucket Sort

**Scenario:** Sorting floating-point numbers

**Example:** You have a list of decimal numbers:

[0.42, 0.32, 0.23, 0.75, 0.15]



## Steps:

**Create buckets:** Divide the range  $[0, 1]$  into buckets. For example, 5 buckets:

Bucket 1:  $[0.0, 0.2)$

Bucket 2:  $[0.2, 0.4)$

Bucket 3:  $[0.4, 0.6)$

Bucket 4:  $[0.6, 0.8)$

Bucket 5:  $[0.8, 1.0)$

**Distribute elements into buckets:**

Bucket 1:  $[0.15]$

Bucket 2:  $[0.23, 0.32]$

Bucket 3:  $[0.42]$

Bucket 4:  $[0.75]$

Bucket 5:  $[]$

**Sort each bucket:** Sort the numbers within each bucket (using a simple algorithm like Insertion Sort).

**Concatenate sorted buckets:** Combine the sorted buckets to produce the final sorted list.

**Sorted Output:**

$[0.15, 0.23, 0.32, 0.42, 0.75]$

---

## LAB ACTIVITY: IMPLEMENTING AND ANALYZING RADIX SORT

**Objective:**

The objective of this lab is to help you understand the **Radix Sort** algorithm by implementing it in Python. You will sort a list of phone numbers and analyze the sorting process at each stage, specifically by examining how the digits are sorted from the least significant to the most significant using **Counting Sort** as a subroutine.

### Scenario: Sorting Phone Numbers

You are provided with a list of unsorted phone numbers:

[564, 213, 987, 432, 123, 765, 321, 654, 876]

Your task is to implement the **Radix Sort** algorithm to arrange these phone numbers in ascending order. **Counting Sort** will be used at each digit level (starting from the least significant digit to the most significant).

### Instructions:

**Implement Radix Sort** in Python.

Use **Counting Sort** as a helper function to sort each digit.

Sort the provided list of phone numbers and observe the output after sorting by each digit.

Analyze the intermediate and final results.

### Questions for Reflection:

What is the order of sorting for the digits in Radix Sort?

After sorting by the least significant digit, what does the intermediate array look like?

After sorting by the most significant digit, what is the final sorted output?

## Medians and Order Statistics

### Introduction to Order Statistics

Order statistics refer to the statistics obtained by arranging a dataset in ascending or descending order and then selecting particular elements based on their position. The most common order statistics are:

**Minimum:** The first element in a sorted list.

**Maximum:** The last element in a sorted list.

**Median:** The middle value in a sorted list.

**k-th Order Statistic:** The element that would appear in the k-th position if the data is sorted.

### Median

The median is a measure of central tendency that represents the middle value in a dataset when the data points are arranged in order. It divides the dataset into two equal halves.

**Odd Number of Elements:** If the dataset has an odd number of elements, the median is the middle element.

Median=Element at position  $(n + 1)/2$

Example: For the dataset 3,7,9,12,15, 7, 9, 12, 15,3,7,9,12,15, the median is 9.

**Even Number of Elements:** If the dataset has an even number of elements, the median is the average of the two middle elements.

$$\text{Median} = \frac{\text{Elements at Position } \frac{n}{2} + \text{Elements at position } \frac{n+1}{2}}{2}$$

Example: For the dataset 3,7,9,12,3, 7, 9, 12,3,7,9,12, the median is 8

## Order Statistics Definitions

**Minimum (1st Order Statistic):** The smallest element in the dataset.

**Maximum (n-th Order Statistic):** The largest element in the dataset.

**k-th Order Statistic:** The k-th smallest element in the dataset.

## Finding the k-th Smallest Element: Quickselect Algorithm

The **Quickselect algorithm** is an efficient way to find the k-th smallest element in an unsorted array. It works similarly to the QuickSort algorithm by using the partitioning method, but only recursively processes one side of the partition based on the desired k-th position.

### Algorithm Steps:

Choose a pivot element.

Partition the array such that elements smaller than the pivot come before it, and elements greater than the pivot come after it.

Recursively apply the Quickselect algorithm to the side of the array that contains the k-th element until the k-th element is found.

**Time Complexity:** On average, Quickselect runs in  $O(n)$ , though in the worst case (poor pivot selection), it can run in  $O(n^2)$ .

## Finding the Median Using the Quickselect Algorithm

The median of a dataset can be considered as the  $\frac{n}{2}$ -th order statistic (for an odd number of elements) or the average of the two middle order statistics (for an even number of elements).

Using Quickselect:

For **odd-sized datasets**, the median is the  $(\frac{n+1}{2})$ -th smallest element.

For **even-sized datasets**, the median is the average of the  $(\frac{n}{2})$ -th and  $(\frac{n+1}{2})$ -th smallest elements.

## Applications of Order Statistics

**Finding Percentiles:** Percentiles are specific order statistics. For example, the 90th percentile is the element that is larger than 90% of the dataset.

**Selection Problems:** Many practical applications, such as selection of top-k elements, median filtering in image processing, and statistical analysis, require order statistics.

**Database Queries:** Median queries or percentile queries in large datasets can be optimized using algorithms like Quickselect or Median of Medians.

## Median of Medians Algorithm

The **Median of Medians** algorithm is a deterministic algorithm to find the median of a dataset. It ensures a linear time complexity  $O(n)$  even in the worst case.

### Steps:

Divide the array into groups of 5 elements each.

Find the median of each group (which can be done using a simple sorting algorithm for small groups).

Recursively apply the Median of Medians algorithm to find the median of the medians of these groups.

Use the median of medians as the pivot to perform a partition similar to Quickselect.

This method is often used when a guaranteed linear time complexity is required, such as in certain critical systems or large-scale applications.

### **Practical Applications of Median and Order Statistics**

**Data Analysis:** Finding the median or k-th smallest element is essential for understanding the distribution of data.

**Computer Graphics:** Median filtering is used to reduce noise in images.

**Machine Learning:** Order statistics are used in algorithms such as k-nearest neighbors, which relies on selecting k smallest distances.

**Statistics:** Percentiles and quartiles are examples of order statistics used for data summarization.

---

## ACTIVITY 4 : END OF MODULE QUIZ

### **1. Which of the following sorting algorithms runs in linear time?**

- a) Quick Sort
- b) Merge Sort
- c) Radix Sort**
- d) Heap Sort

### **2. In which scenario is Counting Sort most efficient?**

- a) When sorting large numbers with many digits
- b) When sorting floating-point numbers
- c) When sorting integers with a small range**
- d) When the input is uniformly distributed

### **3. What is the primary drawback of Counting Sort?**

- a) It is not stable
- b) It requires a large amount of extra memory for the count array**
- c) It only works for strings
- d) It is slower than comparison-based algorithms

### **4. What is the main difference between Counting Sort and Radix Sort?**

- a) Radix Sort uses comparison-based sorting, Counting Sort does not

**b) Counting Sort processes all elements at once, while Radix Sort processes digits**

c) Radix Sort uses divide-and-conquer, Counting Sort does not

d) Counting Sort is recursive, while Radix Sort is not

**5. Which of the following is NOT a step in the Radix Sort algorithm?**

a) Sorting elements by each digit using Counting Sort

**b) Starting from the most significant digit (MSD)**

c) Sorting digits from least significant to most significant

d) Using comparison-based sorting to arrange numbers

**6. In Bucket Sort, why is it important for the input data to be uniformly distributed?**

a) To ensure all elements fall into a single bucket

**b) To prevent overfilling some buckets and underfilling others**

c) To reduce the time complexity to  $O(n^2)$

d) To avoid the need for auxiliary space

**7. What is the time complexity of Radix Sort when sorting numbers with  $d$  digits?**

a)  $O(n \log n)$

**b)  $O(d * (n + k))$**

c)  $O(n^2)$

d)  $O(k^2 + n)$

**8. Which sorting algorithm is typically used as a subroutine within Radix Sort?**

a) Merge Sort

b) Quick Sort

**c) Counting Sort**

d) Insertion Sort

**9. What is the time complexity of Counting Sort?**

a)  $O(n^2)$

b)  $O(n \log n)$

**c)  $O(n + k)$**

d)  $O(d * n)$

**10. What is the role of the auxiliary array in Counting Sort?**

a) To store the sorted elements directly

**b) To count the occurrences of each element**

c) To store intermediate results for bucket merging

d) To track the position of the smallest element in the array

---

## REFERENCE LIST AND FURTHER READING

1. Dimri, S. C., Malik, P., Ram, M. (2021). Algorithms: Design and Analysis. Germany: De Gruyter.
2. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>
3. Thomas, P., & Hunt, S. (2020). *An introduction to algorithms and data structures*. University of London.

# MODULE 5: DYNAMIC PROGRAMMING - PART 1

---

INSTRUCTIONAL HOURS: 4, 5 or 6

---

## MODULE OVERVIEW

This module introduces the foundational concepts of **Dynamic Programming (DP)**, a critical problem-solving paradigm in computer science and algorithm design. It focuses on breaking down complex problems into smaller, overlapping subproblems and solving each subproblem just once. By doing so, DP optimizes the computational process and reduces redundancy.

Through the exploration of dynamic programming, you will learn about the two primary approaches used to implement DP—**memoization** (top-down) and **tabulation** (bottom-up)—and understand how to apply them in practical scenarios. The module also includes a **case study on the Fibonacci sequence**, which serves as a classic example to demonstrate these techniques in action.

---

## LEARNING OUTCOMES

By the end of this module, you will be able to:



**Understand Dynamic Programming Principles:** Identify problems that exhibit optimal substructure and overlapping subproblems and learn how to break them down for efficient solutions.

**Top-Down vs. Bottom-Up Approaches:** Compare the memoization (top-down) and tabulation (bottom-up) techniques and understand their differences in solving dynamic programming problems.

**Implement Dynamic Programming Solutions:** Develop Python solutions for dynamic programming problems using both memoization and tabulation methods.

**Analyze DP Algorithm Efficiency:** Evaluate time and space complexities of dynamic programming solutions and understand their advantages over recursive approaches.

---

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Review Case Study with integrated Lab Activity
3. Complete End of Module Quiz

## Introduction to Dynamic Programming

### What is Dynamic Programming (DP)?

Dynamic Programming is an algorithmic technique for solving problems that can be broken down into smaller, overlapping subproblems. It solves each subproblem only once and stores the result for future reference (either using memoization or tabulation), thereby improving efficiency by avoiding redundant computations.

### What are the main principles of Dynamic Programming?

The main principles of **Dynamic Programming (DP)** are as follows:

#### 1. Optimal Substructure

A problem exhibits **optimal substructure** if an optimal solution to the problem can be constructed efficiently from optimal solutions of its smaller subproblems.

**Example:** In the shortest path problem, the shortest path from a source to a destination can be derived by combining the shortest paths from the source to an intermediate node and from that node to the destination.

#### 2. Overlapping Subproblems

In dynamic programming, overlapping subproblems refer to situations where a problem can be broken down into smaller subproblems that are reused multiple times. This characteristic often occurs in recursive algorithms, where the same subproblem is solved repeatedly without any optimization, leading to inefficiency.

**Key Idea:**

Dynamic programming addresses this inefficiency by solving each subproblem only once and storing its result. This way, when the same subproblem arises again, the algorithm can simply retrieve the stored solution instead of recomputing it.

### **Example:**

Consider the computation of Fibonacci numbers:

The Fibonacci sequence is defined as  $F(n) = F(n-1) + F(n-2)$ .

If we calculate Fibonacci numbers recursively without optimization, the same subproblems, like  $F(n-1)$  and  $F(n-2)$ , will be solved multiple times as the recursion progresses.

For instance, to compute  $F(5)$ , you will compute  $F(4)$  and  $F(3)$ , but to compute  $F(4)$ , you will again need  $F(3)$  and  $F(2)$ , leading to unnecessary repetition.

## **3. Memoization**

Memoization is a technique used in dynamic programming to enhance the efficiency of recursive algorithms. It follows a top-down approach where the solution to each subproblem is stored (or "memorized") to avoid redundant calculations when the same subproblem recurs. Instead of recomputing the result every time, the algorithm simply retrieves the stored value, saving both time and processing power.

It works as follows:

First, the problem is typically broken down recursively into smaller subproblems.

Then, Whenever a subproblem is solved for the first time, its result is saved.

If the same subproblem arises again, the algorithm fetches the saved result instead of solving it again.

### **Example:**

Consider the Fibonacci sequence, where each number is the sum of the two preceding ones:

Without memoization, calculating `fib(5)` would involve multiple redundant calculations for `fib(3)` and `fib(2)`. With memoization, these values are stored after the first calculation and reused, drastically improving the performance of the algorithm.

#### **4. Tabulation**

Tabulation is a dynamic programming technique that employs a bottom-up approach to solve problems. In this method, subproblems are solved in a sequential order, starting from the smallest and progressing to the largest. The results of these subproblems are stored in a table (often an array), allowing for efficient retrieval.

##### **How Tabulation Works:**

The algorithm begins by solving the base cases of the problem and stores their results in the table.

Subsequent subproblems are solved iteratively by referencing previously computed values, filling up the table until the desired solution is reached.

##### **Key Features:**

**No Recursion:** Unlike top-down approaches (like memoization), tabulation avoids recursion entirely. This reduces the overhead associated with function calls, making the algorithm more efficient.

**Iterative Implementation:** Tabulation is often easier to implement as it follows a straightforward iterative pattern, which can enhance clarity and reduce the potential for stack overflow errors in languages with limited recursion depth.

##### **Example:**

Consider the Fibonacci sequence again:

In tabulation, you would create an array to store Fibonacci values.

You start by initializing the base cases:  $F(0) = 0$  and  $F(1) = 1$ .

Then, you iteratively compute and store values for  $F(2)$ ,  $F(3)$ , and so on, up to  $F(n)$ , by referencing previously stored values.

## 5. Trade-off Between Time and Space

In dynamic programming, a crucial principle is the **trade-off between time and space**. This trade-off arises from the practice of storing the results of subproblems (typically in arrays or tables) to avoid redundant computations. By saving these results, dynamic programming significantly reduces the time complexity of algorithms.

Key Concepts:

**Time Complexity:** Refers to the amount of time an algorithm takes to complete based on the size of the input. By storing results, dynamic programming minimizes the need to recompute values, thus speeding up the process.

**Space Complexity:** Refers to the amount of memory an algorithm uses during its execution. Storing subproblem results requires additional memory, which can become a limiting factor, especially for problems with large input sizes.

### Storage Strategies:

The approach you choose for storing subproblem solutions can significantly influence both time and space complexity:

#### Tabulation:

In this bottom-up approach, all subproblem solutions are computed and stored in a table, leading to predictable memory usage. While this method is efficient in terms of time, it can consume a significant amount of space, especially for large problems.

#### Memoization:

This top-down approach computes subproblem solutions only as needed, storing them to avoid future recomputation. While it can save space by not storing unnecessary results, it may involve higher time overhead due to recursive function calls.

## Example:

Consider a scenario where you are calculating the  $n$ th Fibonacci number:

**Using Tabulation:** You would create an array of size  $n+1$  to store all Fibonacci values up to  $F(n)$ . This method uses  $O(n)$  space and reduces time complexity to  $O(n)$ .

**Using Memoization:** You would store only the Fibonacci values needed during the recursion, potentially saving space if not all values are used. However, it might require more time due to the overhead of recursive calls, but the total complexity remains  $O(n)$ .

---

### ACTIVITY 2: CASE STUDY: OPTIMIZING FIBONACCI SEQUENCE CALCULATION USING DYNAMIC PROGRAMMING

#### BACKGROUND

The Fibonacci sequence is a classic example in computer science and mathematics, defined as follows:

$$F(0)=0$$

$$F(1)=1$$

$$F(n)=F(n-1)+F(n-2) \text{ for } n \geq 2$$

The naive recursive implementation to compute Fibonacci numbers can lead to significant inefficiencies due to repeated calculations of the same subproblems.

---

#### Objective

To understand how dynamic programming can optimize the computation of Fibonacci numbers by applying the principles of **memoization** and **tabulation**.

---

#### Part 1: Problem Definition

##### Identify the Problem:

Write the recursive function for calculating the Fibonacci number and analyze its time complexity.

Highlight how this approach results in recalculating the same values multiple times.

### **Explore the Impact of Overlapping Subproblems:**

Discuss which subproblems are repeated during the recursive calculation.

---

## Part 2: Memoization Approach

### **Implementing Memoization:**

Modify the recursive function to include a memoization technique.

Create a structure (e.g., a dictionary or list) to store previously computed Fibonacci numbers.

### **Performance Analysis:**

Compare the time complexity of the original recursive approach with the memoized version.

Discuss how memoization reduces redundant calculations and speeds up the computation.

---

## Part 3: Tabulation Approach

### **Implementing Tabulation:**

Develop a non-recursive function that calculates Fibonacci numbers using tabulation.

Fill an array with Fibonacci values from the base cases up to the desired index.

### **Performance Analysis:**

Analyze the time and space complexity of the tabulation approach.

Discuss the advantages and disadvantages of using tabulation compared to memoization.

---

## Part 4: Practical Implementation in Python Programming Language

### Coding Activity:

Implement both the memoization and tabulation methods in Python.

Ensure that your code is modular, with functions for each approach.

### Testing the Functions:

Write test cases to validate your implementations. Consider edge cases such as  $F(0)$ ,  $F(1)$ , and larger values like  $F(10)$ ,  $F(20)$ , etc.

## Part 5: Reflection Questions

### Understanding Optimal Substructure:

How does the Fibonacci problem demonstrate the concept of optimal substructure?

### Comparative Analysis:

In what scenarios might you prefer one approach over the other (memoization vs. tabulation)?

### Real-World Applications:

Discuss how dynamic programming techniques can be applied to other computational problems, such as coin change or knapsack problems.

---

## FURTHER READING AND REFERENCES ON DYNAMIC PROGRAMMING

For further information and references on the topic refer to:



1. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2009). Dynamic programming (Chapter 14, pp. 362-400). In *Introduction to algorithms* (3rd ed.). MIT Press
2. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>
3. Dimri, S. C., Malik, P., Ram, M. (2021). *Algorithms: Design and Analysis*. Germany: De Gruyter.

---

### ACTIVITY 3: END OF MODULE QUIZ

Question 1: What is the primary characteristic of a problem that exhibits optimal substructure?

- A) The problem can be solved in polynomial time.
- B) An optimal solution can be constructed from optimal solutions of its subproblems.**
- C) The problem has no overlapping subproblems.
- D) The solution requires recursion.

Question 2: Which of the following is a feature of overlapping subproblems?

- A) Each subproblem can be solved independently without affecting others.
- B) The same subproblem is solved multiple times during the computation.**
- C) Subproblems do not recur.
- D) The problem can only be solved using brute force.

Question 3: In dynamic programming, what is memoization?

- A) A technique that uses a bottom-up approach to solve problems.
- B) A method of storing solutions to subproblems to avoid recomputation.**
- C) A process of merging overlapping subproblems.
- D) An algorithm that only works with recursive functions.

Question 4: Which statement is true regarding tabulation?

- A) It is a top-down approach that requires recursion.
- B) It stores intermediate results in a table and fills it in a bottom-up manner.**
- C) It is less space-efficient than memoization.
- D) It is used for problems that do not have overlapping subproblems.

Question 5: Which of the following problems can be solved using dynamic programming?

- A) Finding the shortest path in a graph.
- B) Calculating Fibonacci numbers.
- C) The Knapsack problem.
- D) All of the above.**

Question 6: What is the time complexity of calculating the n-th Fibonacci number using the naive recursive approach?

- A)  $O(n)$
- B)  $O(n^2)$
- C)  $O(2^n)$**
- D)  $O(n \log n)$

Question 7: In the context of the Fibonacci sequence, which of the following statements is true?

- A) Both memoization and tabulation can be used to optimize its computation.**
- B) The Fibonacci sequence does not have optimal substructure.
- C) The recursive approach is always faster than dynamic programming methods.
- D) The Fibonacci numbers are defined only for positive integers.

Question 8: Which of the following best describes the difference between memoization and tabulation?

- A) Memoization is iterative, while tabulation is recursive.
- B) Memoization is a bottom-up approach, while tabulation is a top-down approach.

**C) Memoization stores results of recursive calls, while tabulation builds a table from the ground up.**

D) There is no difference; they are the same technique.

# MODULE 6: DYNAMIC PROGRAMMING - PART 2

---

INSTRUCTIONAL HOURS: 4, 5 or 6

## MODULE OVERVIEW

This module delves deeper into dynamic programming techniques, focusing on three classic problems that demonstrate the power of this approach: the **Minimum Edit Distance problem**, the **Longest Common Subsequence (LCS)**, and the **Knapsack problem**. Each problem will be explained with detailed examples, showcasing how dynamic programming optimally solves these complex challenges by breaking them down into overlapping subproblems. By the end of this module, students will have a firm understanding of these problems and will be able to implement dynamic programming solutions to similar optimization problems.

## LEARNING OUTCOMES

By the end of this module, you will be able to:

1. Apply principles of dynamic Programming to solve advanced computational problems such as the Minimum Edit Distance, LCS, and Knapsack problem.
2. Develop dynamic programming solutions for the Minimum Edit Distance problem, analyzing different types of operations (insertion, deletion, substitution).
3. Implement algorithms for the Longest Common Subsequence (LCS) and understand their application in sequence comparison and data alignment.
4. Solve optimization problems like the Knapsack problem, applying memoization and tabulation techniques to efficiently allocate resources.

---

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Perform Lab activity
3. Complete End of Module Quiz

---

## ACTIVITY 1: REVIEW READING MATERIALS

### Minimum Edit Distance Problem

The **Minimum Edit Distance** is the minimum number of operations (insertions, deletions, or substitutions) required to transform one string into another. This is a common problem in text processing, including spell-checking, DNA sequence alignment, and Natural Language Processing (NLP).

#### Problem Definition

Given two strings  $s_1$  and  $s_2$ , find the minimum number of operations to convert  $s_1$  into  $s_2$ .

#### Allowed Operations:

**Insertion:** Insert a character in  $s_1$ .

**Deletion:** Delete a character from  $s_1$ .

**Substitution:** Replace a character in  $s_1$  with one from  $s_2$ .

#### Example

For the words "kitten" and "sitting", the transformation can be done in 3 steps:

Replace 'k' with 's'

Replace 'e' with 'i'

Insert 'g' at the end

#### Dynamic Programming Approach

We use a 2D table  $dp[i][j]$  where:

$i$  is the length of the prefix of  $s_1$ .

j is the length of the prefix of s2.

The entry  $dp[i][j]$  will store the minimum number of operations required to convert the first i characters of s1 into the first j characters of s2.

## Recurrence Relation

If  $s1[i] == s2[j]$ : No operation is needed, so  $dp[i][j] = dp[i-1][j-1]$ .

If  $s1[i] != s2[j]$ : We take the minimum of the three possible operations:

Insert:  $dp[i][j-1] + 1$

Delete:  $dp[i-1][j] + 1$

Substitute:  $dp[i-1][j-1] + 1$

## Implementation in Python

```
def min_edit_distance(s1, s2):

    m, n = len(s1), len(s2)

    dp = [[0] * (n + 1) for _ in range(m + 1)]

    for i in range(m + 1):

        for j in range(n + 1):

            if i == 0:

                dp[i][j] = j # Insert all characters from s2

            elif j == 0:

                dp[i][j] = i # Delete all characters from s1

            elif s1[i-1] == s2[j-1]:
```

```
    dp[i][j] = dp[i-1][j-1]

    else:

        dp[i][j] = 1 + min(dp[i-1][j], dp[i][j-1], dp[i-1][j-1])

    return dp[m][n]
```

### Time and Space Complexities

**Time Complexity:**  $O(m * n)$ , where  $m$  is the length of  $s1$  and  $n$  is the length of  $s2$ .

**Space Complexity:**  $O(m * n)$ , as a 2D table is used.

## Longest Common Subsequence (LCS)

The **Longest Common Subsequence** (LCS) problem is about finding the longest subsequence common to two sequences (strings) without changing the order of elements. It has applications in bioinformatics (DNA sequence alignment) and version control systems (detecting differences between file versions).

### Problem Definition

Given two sequences  $s1$  and  $s2$ , find the length of their longest common subsequence.

### Example

For the sequences "ABCB DAB" and "BDCAB", the LCS is "BDAB", with a length of 4.

### Dynamic Programming Approach



We construct a 2D table  $dp[i][j]$  where:

$i$  is the length of the prefix of  $s_1$ .

$j$  is the length of the prefix of  $s_2$ .

The entry  $dp[i][j]$  represents the length of the LCS of the first  $i$  characters of  $s_1$  and the first  $j$  characters of  $s_2$ .

## Recurrence Relation

If  $s_1[i] == s_2[j]$ : LCS can be extended, so  $dp[i][j] = dp[i-1][j-1] + 1$ .

If  $s_1[i] != s_2[j]$ : We take the maximum of ignoring one character from either  $s_1$  or  $s_2$ :  $dp[i][j] = \max(dp[i-1][j], dp[i][j-1])$ .

## Implementation in Python

```
def lcs(s1, s2):  
  
    m, n = len(s1), len(s2)  
  
    dp = [[0] * (n + 1) for _ in range(m + 1)]  
  
    for i in range(1, m + 1):  
        for j in range(1, n + 1):  
            if s1[i-1] == s2[j-1]:  
                dp[i][j] = dp[i-1][j-1] + 1  
            else:  
                dp[i][j] = max(dp[i-1][j], dp[i][j-1])  
  
    return dp[m][n]
```

## Time and space Complexities

**Time Complexity:**  $O(m * n)$ , where  $m$  is the length of  $s_1$  and  $n$  is the length of  $s_2$ .

**Space Complexity:**  $O(m * n)$ .

## 3. Knapsack Problem

The **0/1 Knapsack Problem** is a classic optimization problem where you have a set of items, each with a weight and a value, and you must select a subset of items to maximize the total value without exceeding a given weight limit.

### Problem Definition

You are given:

A set of items, each with a weight and value.

A knapsack with a fixed capacity.

Find the maximum value you can achieve by selecting a subset of items such that their combined weight does not exceed the knapsack's capacity.

### Example

Consider 4 items with weights [2, 3, 4, 5] and values [3, 4, 5, 6], and a knapsack capacity of 5. The maximum value is 7 (items with weights 2 and 3).

### Dynamic Programming Approach

We use a 2D table  $dp[i][w]$  where:

$i$  is the number of items considered.

$w$  is the current capacity of the knapsack.

The entry  $dp[i][w]$  represents the maximum value achievable with the first  $i$  items and capacity  $w$ .

## Recurrence Relation

If  $weight[i] > w$ : The current item cannot fit in the knapsack, so  $dp[i][w] = dp[i-1][w]$ .

Otherwise: We have two options—either include the item or exclude it:

Exclude the item:  $dp[i][w] = dp[i-1][w]$ .

Include the item:  $dp[i][w] = dp[i-1][w - weight[i]] + value[i]$ .

```
def knapsack(weights, values, capacity):  
    n = len(values)  
    dp = [[0] * (capacity + 1) for _ in range(n + 1)]  
  
    for i in range(1, n + 1):  
        for w in range(1, capacity + 1):  
            if weights[i-1] <= w:  
                dp[i][w] = max(dp[i-1][w], dp[i-1][w - weights[i-1]] + values[i-1])  
            else:  
                dp[i][w] = dp[i-1][w]  
  
    return dp[n][capacity]
```

## Complexities

**Time Complexity:**  $O(n * \text{capacity})$ , where  $n$  is the number of items.

**Space Complexity:**  $O(n * \text{capacity})$ .

### Summary of problem Solving Approach

For each problem:

**Understand the problem:** Define the variables and determine which subproblems repeat.

**Formulate a recurrence relation:** Develop a formula that expresses the problem in terms of smaller subproblems.

**Set up the dynamic programming table:** Initialize the table based on the recurrence relation and iteratively fill it in.

**Implement the solution:** Translate the logic into a Python implementation.

## ACTIVITY 2: LAB ACTIVITY ON LONGEST COMMON SUBSEQUENCE (LCS)

### Problem Statement:

You are required to implement a dynamic programming solution to find the **Longest Common Subsequence (LCS)** between two strings  $s1$  and  $s2$ . The LCS problem involves finding the longest subsequence common to both strings without rearranging the order of characters.

### Task:

Write a Python function `lcs(s1, s2)` that calculates the length of the longest common subsequence between the two input strings  $s1$  and  $s2$  using dynamic programming.

Implement the following steps in your function:

Create a 2D table `dp` where `dp[i][j]` represents the length of the LCS of the first  $i$  characters of  $s1$  and the first  $j$  characters of  $s2$ .

Fill the table using the recurrence relation:

If  $s1[i-1] == s2[j-1]$ : The LCS can be extended, so  $dp[i][j] = dp[i-1][j-1] + 1$

If  $s1[i-1] != s2[j-1]$ : Take the maximum of ignoring one character from either  $s1$  or  $s2$ , i.e.,  $dp[i][j] = \max(dp[i-1][j], dp[i][j-1])$

After filling in the table, the length of the longest common subsequence will be stored in  $dp[\text{len}(s1)][\text{len}(s2)]$ .

Provide a brief explanation of the dynamic programming approach used.

### Instructions:

Use the Colab or Jupyter notebook platform for your implementation.

Name your notebook as **Last\_Firstname\_LCS.ipynb**.

Submit your notebook for evaluation on Learning Management System

---

### FURTHER READING AND REFERENCES ON DYNAMIC PROGRAMMING

For further information and references on the topic refer to:

1. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2009). Dynamic programming (Chapter 14, pp. 362-400). In *Introduction to algorithms* (3rd ed.). MIT Press
2. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>
3. Dimri, S. C., Malik, P., Ram, M. (2021). *Algorithms: Design and Analysis*. Germany: De Gruyter.

4. Thomas, P., & Hunt, S. (2020). *An introduction to algorithms and data structures*. University of London.

---

ACTIVITY 3: END OF MODULE QUIZ

What is the minimum number of operations required to transform one string into another in the **Minimum Edit Distance** problem?

**A. Insertion, Deletion, and Substitution**

B. Insertion, Substitution, and Reversal

C. Deletion, Insertion, and Transposition

D. Insertion, Substitution, and Swapping

2. Which dynamic programming table entry represents the final result in the **Minimum Edit Distance** problem?

A.  $dp[0][0]$

**B.  $dp[m][n]$**

C.  $dp[1][n]$

D.  $dp[m][0]$

In the **LCS** problem, which of the following statements is true about the recurrence relation when  $s1[i-1] == s2[j-1]$ ?

A.  $dp[i][j] = dp[i-1][j] + 1$

**B.  $dp[i][j] = dp[i-1][j-1] + 1$**

C.  $dp[i][j] = \max(dp[i-1][j], dp[i][j-1])$

D.  $dp[i][j] = dp[i][j-1] + 1$

What is the time complexity of finding the LCS using dynamic programming for two strings of lengths  $m$  and  $n$ ?

**A.  $O(m * n)$**

B.  $O(m + n)$

C.  $O(m^2)$

D.  $O(n^2)$

In the **0/1 Knapsack problem**, what does the entry  $dp[i][w]$  represent in the dynamic programming table?

A. The total number of items selected.

B. The total weight of the items selected.

**C. The maximum value achieved with the first  $i$  items and a capacity  $w$ .**

D. The minimum weight that can be achieved with  $i$  items.

If the weight of the current item exceeds the knapsack's remaining capacity, what does the recurrence relation in the **Knapsack problem** do?

- A. It excludes the item and sets  $dp[i][w] = dp[i-1][w]$ .
- B. It includes the item and reduces the capacity by the item's weight.
- C. It skips the entire table calculation.
- D. It sets the value to zero.



## MODULE 7: GREEDY ALGORITHMS

---

INSTRUCTIONAL HOURS: 4, 5 or 6

### MODULE OVERVIEW

This module introduces greedy algorithms, a problem-solving approach used in optimization problems. The key characteristic of a greedy algorithm is to make a series of choices, each of which looks the best at that moment (locally optimal choices), with the hope that the result will lead to a globally optimal solution. Greedy algorithms are widely used due to their simplicity and efficiency in many practical problems such as networking, scheduling, and graph theory.

### LEARNING OUTCOMES

By the end of this module, students will be able to:

Understand the core characteristics of greedy algorithms and how they differ from other algorithmic techniques.

Apply greedy algorithms to specific problems such as Huffman coding, Prim's, and Kruskal's algorithms.

Analyze case studies including interval scheduling maximization and Dijkstra's shortest path algorithm.

Evaluate the efficiency and limitations of greedy algorithms in solving optimization problems.

---

### LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Perform Lab activity
3. Review youtube Video
4. Complete End of Module Quiz

## Introduction to Greedy Algorithms

### What is a greedy algorithm?

Greedy algorithms are a problem-solving paradigm used in optimization problems. These algorithms build a solution step by step, choosing the best available option at each step (locally optimal choice), with the hope that the final solution will be globally optimal. Unlike dynamic programming, which solves problems by breaking them into overlapping subproblems, greedy algorithms focus on the current step without revisiting or reconsidering choices made in previous steps.

### Characteristics of Greedy Algorithms

Greedy algorithms make a series of choices in a problem, each time picking the locally optimal solution with the hope that it leads to a globally optimal solution. The characteristics of greedy algorithms include:

**Greedy Choice Property:** The algorithm makes choices that seem best at the moment.

**Optimal Substructure:** The solution to a problem contains within it solutions to smaller subproblems.

**Local Optimality:** Each step is taken to ensure the locally optimal solution is selected, even if it is not clear that this will yield the globally optimal solution.

**Irrevocable Decisions:** Once a choice is made, it cannot be changed, leading to efficient computation but also potential suboptimality in certain cases.

Examples of greedy algorithms include Huffman Coding, Prim's Algorithm, Kruskal's Algorithm, Dijkstra's Shortest Path Algorithm, etc

## Huffman Coding

Huffman coding is an algorithm used for lossless data compression. It is based on the frequency of occurrence of characters in a data set. The algorithm assigns shorter binary codes to more frequent characters and longer codes to less frequent ones, thus minimizing the overall length of the encoded data.

### How Huffman Coding Works

#### Frequency Count

First, count the frequency of each character in the data to be encoded.

#### Building a Min-Heap

Create a priority queue (min-heap) where each node is a character and its frequency. Characters with lower frequencies have higher priority.

#### Tree Construction

The next step is to construct a binary tree. In each iteration, two nodes with the lowest frequencies are extracted from the heap, and a new node is created by merging these two nodes. This new node's frequency is the sum of the frequencies of the two nodes. This process continues until there is only one node left in the heap, which becomes the root of the Huffman Tree.

#### Assigning Codes

Once the Huffman Tree is built, traverse the tree to assign binary codes to each character. A left edge represents a "0," and a right edge represents a "1."

## Encoding the Data

Replace each character in the original data with its corresponding binary code.

### Example of Huffman Coding

Consider the following character frequencies:

| Character | Frequency |
|-----------|-----------|
| A         | 5         |
| B         | 9         |
| C         | 12        |
| D         | 13        |
| E         | 16        |
| F         | 45        |

#### Step 1: Initial Heap

Create a min-heap with each character and its frequency:

$[(A, 5), (B, 9), (C, 12), (D, 13), (E, 16), (F, 45)]$

#### Step 2: Build the Tree

Merge A and B ( $5+9=14$ ), creating a new node with frequency 14. New heap:  $[(C, 12), (D, 13), (E, 16), (F, 45), (AB, 14)]$

Merge C and D ( $12+13=25$ ). New heap:  $[(E, 16), (F, 45), (AB, 14), (CD, 25)]$

Merge AB and E ( $14+16=30$ ). New heap:  $[(F, 45), (CD, 25), (ABE, 30)]$

Merge CD and ABE ( $25+30=55$ ). New heap:  $[(F, 45), (ABCDE, 55)]$

Merge F and ABCDE ( $45+55=100$ ), creating the final tree.

### Step 3: Assign Binary Codes

Traverse the tree to assign binary codes to each character:

A: 1100

B: 1101

C: 100

D: 101

E: 111

F: 0

### Encoding Example

If the data to be encoded is "ABCDEF", it would be represented as:

11001101100101110

### Why Huffman Coding is Greedy

Huffman coding uses the greedy strategy by always merging the two nodes with the smallest frequencies, aiming to reduce the overall cost (in terms of total number of bits). This locally optimal choice at each step ensures that the final solution is the most compressed representation of the data, demonstrating both the greedy choice property and optimal substructure.

---

## ACTIVITY 2: LAB ACTIVITY

**Write a Python program that implements Huffman coding for a given set of characters and their frequencies.**

### Prim's Algorithm

**Prim's Algorithm** is a classical greedy algorithm used to find a **minimum spanning tree (MST)** for a connected, undirected graph with weighted edges. The minimum spanning tree ensures that all the vertices are connected with the minimum possible total edge weight, without forming any cycles. The algorithm was first discovered in 1930 by Vojtech Jarnik and later independently rediscovered by Robert Clay Prim in 1957.

#### Steps of Prim's Algorithm:

##### Initialization:

Start from an arbitrary vertex (source).

Mark this vertex as part of the MST.

Add all edges connected to this vertex to a priority queue (or min heap), sorted by edge weight.

##### Choose the Minimum Edge:

Select the edge with the smallest weight from the priority queue, provided it connects to a vertex not yet in the MST.

##### Add to MST:

Add the selected edge and the connected vertex to the MST.

Mark the newly connected vertex as part of the MST.

##### Repeat:

Add all edges from the newly connected vertex to the priority queue, ensuring that no edges are selected that lead to a vertex already in the MST.

Continue this process until all vertices are included in the MST.

### **Termination:**

The algorithm terminates when all vertices are included in the MST.

## **Activity 2: Understanding Minimum Spanning Tree and Prim's Algorithm**

### **Objective:**

To watch the YouTube video on Minimum Spanning Tree (MST) and Prim's Algorithm, and reflect on the concepts presented in the video.

### **Instructions:**

Watch the YouTube video titled "*Minimum Spanning Tree (MST) and Prim's Algorithm*" at this link: [Watch Video](#) (video is by WIT Solapur)

Pay attention to how the MST is derived and the step-by-step explanation of Prim's Algorithm.

After watching, answer the reflection questions below.

### **Reflection Questions**

#### **Explain the concept of a Minimum Spanning Tree (MST).**

What is the main objective of constructing an MST in a graph?

#### **Describe how Prim's Algorithm works to derive the Minimum Spanning Tree.**

What is the importance of selecting the minimum edge at each step?



## Why is it important to avoid cycles when constructing a Minimum Spanning Tree using Prim's Algorithm?

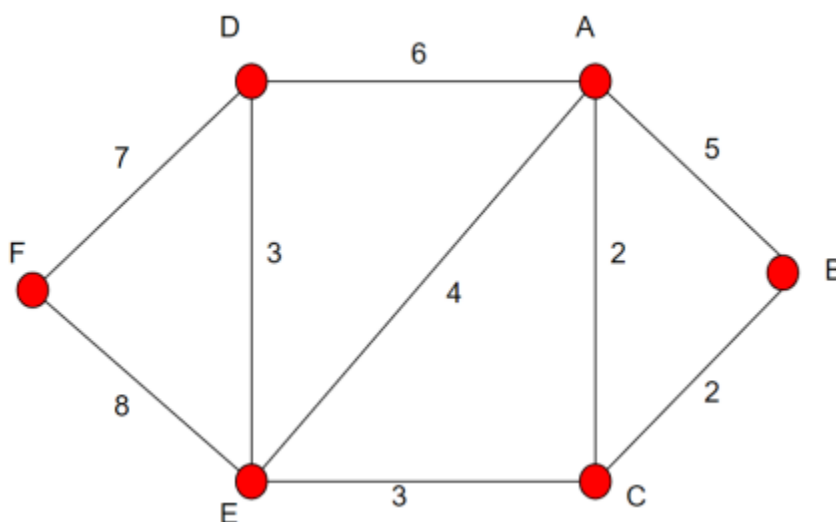
Provide an example from the video where a cycle might have been formed and how it was avoided.

**Reflect on the computational efficiency of Prim's Algorithm.**

**Discuss a real-world application of Minimum Spanning Trees.**

### Activity 3: Mastery Quiz on Prim's Algorithm

Maria is an irrigation specialist who has installed a system of water distribution points across her farm, as illustrated in the diagram below. Use Prim's algorithm to determine the network that will connect all of the water distribution points with the least amount of piping. Additionally, calculate the total length of piping needed. In this scenario, each vertex represents a water distribution point, and the weight of each edge signifies the distance in meters.



## Kruskal's Algorithm

Kruskal's Algorithm is a popular algorithm for finding the Minimum Spanning Tree (MST) of a connected, weighted graph. It operates by selecting the edges of the graph in order of increasing weight while ensuring that no cycles are formed.

### Steps of Kruskal's Algorithm

#### Sort the Edges:

Begin by sorting all the edges of the graph in non-decreasing order based on their weights.

#### Initialize the MST:

Create an empty set to store the edges that will be included in the MST.

#### Union-Find Data Structure:

Use a union-find (disjoint-set) data structure to manage the vertices and to efficiently check for cycles. This structure supports two main operations:

**Find:** Determines which component a particular element belongs to.

**Union:** Merges two components into a single component.

#### Add Edges:

Iterate through the sorted edges, and for each edge:

Use the **Find** operation to check if the endpoints of the edge belong to different components.

If they do not form a cycle (i.e., belong to different components), add the edge to the MST and perform a **Union** operation to merge the components.

## Repeat Until MST is Complete:

Continue this process until you have included  $V-1$  edges in the MST, where  $V$  is the number of vertices in the graph.

## Pseudocode

Here's a simple pseudocode representation of Kruskal's Algorithm:

```
Kruskal(G):  
  
    MST = []  
  
    sort edges of G by weight  
  
    create a union-find structure for each vertex in G  
  
    for each edge (u, v) in sorted edges:  
        if Find(u) ≠ Find(v): // Check if u and v are in different components  
            MST.add((u, v)) // Add edge to the MST  
            Union(u, v) // Merge the components  
  
    return MST
```

## Time Complexity

Sorting the edges takes  $O(E \log E)$ , where  $E$  is the number of edges.

Each union and find operation takes  $O(\alpha(V))$ , where  $\alpha$  is the inverse Ackermann function, which grows very slowly.

The overall time complexity is  $O(E \log E + E \alpha(V))$ , which simplifies to  $O(E \log E)$  for dense graphs.

### Example

Consider a graph with the following edges and weights:

(A, B, 1)

(A, C, 3)

(B, C, 2)

(B, D, 4)

(C, D, 5)

### Sorted Edges:

(A, B, 1)

(B, C, 2)

(A, C, 3)

(B, D, 4)

(C, D, 5)

### Process Edges:

Add (A, B)  $\rightarrow$  MST: {(A, B)}

Add (B, C)  $\rightarrow$  MST: {(A, B), (B, C)}

Skip (A, C)  $\rightarrow$  forms a cycle.

Add (B, D)  $\rightarrow$  MST: {(A, B), (B, C), (B, D)}

Skip (C, D) → forms a cycle.

**Final MST:**

MST: {(A, B), (B, C), (B, D)} with a total weight of  $1 + 2 + 4 = 7$ .

**Applications of Kruskal Algorithm**

1. Network design (e.g., designing the layout of cables, roads, etc.)
2. Clustering algorithms
3. Creating efficient communication networks

## ACTIVITY 3: WATCH A YOUTUBE VIDEO ON KRUSKAL'S ALGORITHM

**Video Link:** [Kruskal's Algorithm\(Video is by WIT Solapur\)](#)

**Instructions:**

- Watch the YouTube video on Kruskal's Algorithm.
- Take notes on the key concepts, steps of the algorithm, and any examples provided in the video.

**Reflection Questions based on the video**

**Understanding the Basics:**

What is the primary purpose of Kruskal's Algorithm? How does it differ from other algorithms used for finding a Minimum Spanning Tree?

**Algorithm Steps:**

List the main steps involved in executing Kruskal's Algorithm. Why is it important to sort the edges before processing them?

### **Cycle Detection:**

How does the union-find data structure assist in detecting cycles during the execution of Kruskal's Algorithm? Explain the Find and Union operations.

### **Practical Applications:**

In what real-world scenarios can Kruskal's Algorithm be effectively applied? Provide at least two examples.

### **Algorithm Efficiency:**

Discuss the time complexity of Kruskal's Algorithm. Under what circumstances might this algorithm be more efficient compared to other MST algorithms like Prim's Algorithm?

### **Personal Insight:**

Reflect on your understanding of Kruskal's Algorithm after watching the video. What concepts or steps did you find most challenging, and how would you approach them differently in the future?

### **Submission**

Write your answers to the reflection questions and any additional notes from the video.

Submit your work in the LMS

### **Dijkstra's Algorithm**

Dijkstra's algorithm finds the shortest path from a source vertex to all other vertices in a graph with non-negative edge weights. It operates by always selecting the vertex with the shortest known distance to the source and updating the distances to its neighbors.

To understand how the algorithm works watch youtube video in this link

Video Link:[Dijkstra's Algorithm](#)( Video is created by [Mr ARUL SUJU D](#))

---

## ACTIVITY 3 DISCUSSION FORUM ACTIVITY: UNDERSTANDING MINIMUM SPANNING TREES

**Topic:** Edge Weights in Minimum Spanning Trees

Consider a graph  $G$ , with  $T$  as its minimum spanning tree (MST) and  $L$  representing the sorted list of edge weights of  $T$ . It is claimed that for any other minimum spanning tree  $T'$  of  $G$ , the list  $L$  remains the same, i.e., the sorted list of edge weights in  $T'$  will also be  $L$ .

### Discussion Questions:

**Do you agree with the claim?** Why or why not? Use examples of graphs and their minimum spanning trees to support your answer.

**What is the significance of this property?** How does it help in understanding the uniqueness or variability of minimum spanning trees in a graph?

**Explore different algorithms** for constructing minimum spanning trees (such as Kruskal's and Prim's algorithms). Do they always yield the same sorted list of edge weights?

### Instructions:

Post your thoughts and reasoning on the forum, ensuring that you explain your ideas with clarity and include examples where possible.

Respond to at least two peers' posts, agreeing or disagreeing with their points and providing your reasoning or additional insights.

---

## FURTHER READING AND REFERENCES ON GREEDY ALGORITHMS

For further information and references on the topic refer to:

1. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2009). Prim and Kruskal algorithms (pp. 591-597) & Dijkstra's algorithm (pp. 620-624). In *Introduction to algorithms* (3rd ed.). MIT Press.
2. Dimri, S. C., Malik, P., Ram, M. (2021). *Algorithms: Design and Analysis*. Germany: De Gruyter.

---

## ACTIVITY 4: END OF MODULE MODULE QUIZ

1. **What is the primary function of Dijkstra's Algorithm?**
  - a) To find the Minimum Spanning Tree of a graph
  - b) To find the shortest path from a source vertex to all other vertices in a weighted graph**
  - c) To sort the edges of a graph
  - d) To detect cycles in a graph
2. **. Which of the following statements is true regarding Dijkstra's Algorithm?**
  - a) It can be used on graphs with negative edge weights.
  - b) It guarantees the shortest path from the source to all vertices.**



- c) It requires the graph to be connected.
  - d) It uses depth-first search (DFS) for traversal.
3. **In Kruskal's Algorithm, what data structure is commonly used to keep track of the connected components?**
- a) Stack
  - b) List
  - c) Union-Find (Disjoint Set)**
  - d) Queue
4. **What is the main difference between Prim's and Kruskal's algorithms when constructing a Minimum Spanning Tree (MST)?**
- a) Prim's algorithm starts from any vertex, while Kruskal's algorithm starts from a specific edge.
  - b) Prim's algorithm builds the MST one vertex at a time, while Kruskal's algorithm adds edges to the MST.**
  - c) Prim's algorithm requires sorting of edges, while Kruskal's does not.
  - d) Prim's algorithm can handle negative weights, while Kruskal's cannot.
5. **In which of the following scenarios would you prefer to use Kruskal's Algorithm over Prim's Algorithm?**
- a) When the graph is dense (many edges)
  - b) When the graph is sparse (few edges)**
  - c) When you need to find the shortest path from one vertex to another
  - d) When the graph is unweighted

# MODULE 8: DATA STRUCTURES FOR SET MANIPULATION

---

INSTRUCTIONAL HOURS: 4, 5 or 6

## MODULE OVERVIEW

In this module, you will explore both fundamental and advanced data structures used for efficient set manipulation. The module starts with basic structures like queues, stacks, and binary search trees and progresses to more sophisticated structures such as Red-Black Trees and AVL Trees. These data structures are essential for managing ordered data and optimizing operations like insertion, deletion, and searching. By the end of the module, you will understand how different data structures impact the efficiency of algorithms and how to choose the right structure for various computational tasks.

## LEARNING OUTCOMES

By the end of this module, students will be able to:

1. **Understand and implement fundamental data structures** such as queues, stacks, and binary search trees for set manipulation tasks.
2. **Analyze the properties and performance** of binary search trees and randomized search trees in set operations.
3. **Compare advanced data structures**, such as Red-Black Trees and AVL Trees, in terms of their balancing mechanisms and efficiency in dynamic set operations.
4. **Evaluate and apply appropriate data structures** to solve problems that require efficient insertion, deletion, and searching in sets.
5. **Design algorithms** that leverage the unique strengths of advanced tree structures to improve performance in complex computational tasks.

---

## LEARNING ACTIVITIES/TASK LIST



Review Reading Material

Perform Lab activity

Attend to Discussion Forum

Complete End of Module Quiz

---

## REVIEW READING MATERIAL

### Fundamental Data Structures:

#### Queues and Stacks:

##### Queues:

A **queue** is a linear data structure that follows the **FIFO (First-In, First-Out)** principle. This means that the element added first will be the one removed first. Queues are widely used in scenarios where order matters and tasks must be processed in the sequence they arrive.

##### Key Operations:

**Enqueue:** Add an element to the back (rear) of the queue.

**Dequeue:** Remove an element from the front of the queue.

**Peek/Front:** Retrieve the element at the front of the queue without removing it.

**IsEmpty:** Check if the queue is empty.

### **Use Cases:**

**Task scheduling:** Queues are used in operating systems for managing processes and tasks.

**Buffering:** In situations like streaming or handling asynchronous data, queues help in managing data packets in a sequential order.

**Breadth-First Search (BFS):** This graph traversal algorithm uses a queue to explore nodes level by level.

### **Stacks:**

A **stack** is a linear data structure that follows the **LIFO (Last-In, First-Out)** principle, meaning that the last element added to the stack is the first to be removed. Stacks are particularly useful in scenarios involving nested or recursive operations.

### **Key Operations:**

**Push:** Add an element to the top of the stack.

**Pop:** Remove the element from the top of the stack.

**Peek/Top:** Retrieve the top element without removing it.

**IsEmpty:** Check if the stack is empty.

### **Use Cases:**

**Function calls:** Stacks are used to keep track of function calls in programming. The most recent function call is placed on top of the stack and removed after completion.

**Expression evaluation:** Stacks are used to evaluate and convert infix expressions to postfix or prefix (used in compilers).

**Backtracking:** In algorithms like Depth-First Search (DFS) or maze-solving problems, stacks help in remembering previous states when backtracking is needed.

## Binary Search Trees (BSTs)

A **Binary Search Tree (BST)** is a hierarchical data structure where each node has at most two children, referred to as the left and right child. The BST property ensures that the left subtree of a node contains only nodes with keys less than the node's key, and the right subtree contains only nodes with keys greater than the node's key.

### Key Properties:

**Binary Tree Structure:** Each node in a BST has up to two children.

**Ordered Structure:** For any node, all elements in the left subtree are smaller, and all elements in the right subtree are larger.

**No Duplicates:** Typically, BSTs do not allow duplicate keys, but this can be modified based on specific implementations.

### Key Operations:

**Insertion:** Inserting a new key into a BST involves finding the correct position by following the BST property. The time complexity is  $O(\log n)$  for balanced trees, but can degrade to  $O(n)$  in the worst case (when the tree is skewed).

**Search:** Searching for a key follows the same principle of comparing the target key with the current node and moving left or right accordingly. The time complexity is  $O(\log n)$  on average but can degrade to  $O(n)$  in the worst case.

**Deletion:** Deleting a node involves three cases:

**Leaf node deletion:** Simply remove the node.

**Node with one child:** Replace the node with its child.

**Node with two children:** Replace the node with its in-order predecessor (maximum value in the left subtree) or in-order successor (minimum value in the right subtree).

**Traversal:**

**In-order Traversal:** Visits nodes in ascending order (left subtree → root → right subtree).

**Pre-order Traversal:** Visits the root before its subtrees (root → left subtree → right subtree).

**Post-order Traversal:** Visits subtrees before the root (left subtree → right subtree → root).

**Example**

```
    15
   /  \
  10   20
 / \  / \
8 12 17 25
```

**Search for 12:** Start at 15 (root), move left to 10, and then right to 12.

**Insert 13:** Start at 15, move left to 10, right to 12, and insert 13 as the right child of 12.

**Performance:**

**Average-case time complexity:**  $O(\log n)$  for search, insert, and delete operations when the tree is balanced.

**Worst-case time complexity:**  $O(n)$  occurs when the tree becomes skewed (like a linked list).

### Use Cases:

**Search applications:** Efficient searching and retrieval in databases.

**Sorted data management:** Maintaining a dynamically changing set of sorted data.

**Implementing associative arrays:** Used in the implementation of map or dictionary data types.

## Randomized Search Trees

A **Randomized Search Tree** is a type of binary search tree (BST) that incorporates randomization to maintain balance and avoid worst-case scenarios. By ensuring that the structure is balanced with high probability, randomized search trees provide consistently good performance for operations like insertion, deletion, and searching.

One common example of a randomized search tree is the **Treap**, a hybrid data structure that combines properties of both a binary search tree and a heap. Each node in a treap has a key (following the BST property) and a priority (following the heap property), where the priority is chosen randomly.

### Key Concepts:

**Binary Search Tree Property:** The keys in a randomized search tree follow the standard BST property, where for each node, all keys in the left subtree are smaller, and all keys in the right subtree are larger.

**Randomized Priority:** Each node is assigned a random priority. The heap property ensures that the priority of any node is greater than the priorities of its children, giving the tree a random structure.

**Balanced with High Probability:** The random priorities ensure that the tree stays balanced with high probability, leading to expected logarithmic depth and performance.

### **Operations:**

#### **Insertion:**

Insert the new node using the BST insertion method based on the key.

Assign the new node a random priority.

Restore the heap property by rotating the tree if necessary (rotations are performed if the priority of a node is greater than its parent's priority).

**Time complexity:**  $O(\log n)$  expected, as the tree is balanced with high probability.

#### **Deletion:**

Find the node to be deleted using the BST property.

To delete the node, rotate it down to a leaf position (maintaining the heap property) and then remove it.

**Time complexity:**  $O(\log n)$  expected.

**Search:** The search operation is similar to a standard BST, where you traverse the tree based on key comparison.

**Time complexity:**  $O(\log n)$  expected.

### **Example (Treap):**



Consider inserting nodes with keys and randomly assigned priorities:

Insert key 15 with priority 20.

Insert key 10 with priority 35 (heap property requires rotations to maintain priorities).

Insert key 20 with priority 15.

The resulting tree would satisfy both the BST property (based on keys) and the heap property (based on priorities), ensuring it remains balanced with high probability.

### **Performance:**

**Expected time complexity for insertion, deletion, and search:**  $O(\log n)$

**Worst-case time complexity:**  $O(n)$ , but this is rare due to the randomized nature of the structure.

### **Advantages:**

**Balancing without additional bookkeeping:** Unlike AVL trees or Red-Black trees, which require careful rebalancing rules, randomized search trees rely on random priorities to keep the tree balanced.

**Simple balancing mechanism:** Rotations are based on priority comparisons, making the balancing process straightforward.

**Probabilistic guarantees:** The tree stays balanced with high probability, providing consistent logarithmic performance for dynamic operations.

### **Use Cases:**

**Dynamic set operations:** Randomized search trees are useful when you need to maintain a dynamic set of sorted data, where the order of data is constantly changing.

**Probabilistic data structures:** Ideal for applications where probabilistic balancing is preferred over deterministic balancing (such as in databases and caches).

# Advanced Data Structures:

## Red-Black Trees

A **Red-Black Tree** is a type of self-balancing binary search tree that ensures the tree remains approximately balanced, which guarantees  $O(\log n)$  time complexity for insertion, deletion, and search operations. This is achieved through a set of properties that maintain the height of the tree close to the minimum possible for the number of nodes.

### Key Properties:

**Each node is either red or black:** The color of each node helps maintain balance without needing complex height comparisons.

**The root is always black:** This simplifies the balancing rules.

**Every leaf (NULL) is black:** These are the "sentinel" leaves that help with maintaining the structure.

**Red nodes cannot have red children:** This property ensures that the tree does not become unbalanced by limiting consecutive red nodes.

**Every path from a given node to its descendant leaves has the same number of black nodes:** This is called the **black-height property**, which guarantees that no path in the tree is more than twice as long as any other.

### Operations:

#### Insertion:

A new node is always inserted as a red node to avoid violating the black-height property.

After insertion, if the parent of the new node is also red, this creates a violation of the red-black tree properties, and **rotations** and **recoloring** are performed to restore balance.

**Time complexity:**  $O(\log n)$ , since the tree's height is at most  **$2\log n$** , and only a few rotations are required to restore balance.

### **Deletion:**

The node to be deleted is removed, and the tree is rebalanced by adjusting the color and performing rotations if necessary.

Deletion is more complex than insertion, as it might involve several recoloring and rotation steps to maintain the red-black properties.

**Time complexity:**  $O(\log n)$ .

**Search:** The search operation works similarly to a standard binary search tree, where comparisons are made to traverse the tree from the root down to the leaves.

**Time complexity:**  $O(\log n)$ .

### **Rotations:**

Rotations are used in both insertion and deletion operations to restore the red-black tree properties.

**Left Rotation:** Rotates a subtree to the left, moving the right child up and the current node down to the left.

**Right Rotation:** Rotates a subtree to the right, moving the left child up and the current node down to the right.

These rotations ensure that the tree remains balanced after modifications.

### **Example:**

Consider the following sequence of operations on a Red-Black Tree:

Insert nodes with keys 7, 3, 18, 10, 22, 8, 11, and 26.

As you insert nodes, rotations and recoloring occur when necessary to maintain the red-black properties.

After all the insertions, the tree maintains a balanced structure, where the black-height property holds and no path is more than twice as long as any other path.

### **Performance:**

**Time complexity** for search, insertion, and deletion:  $O(\log n)$  in all cases.

The maximum height of a red-black tree is  $2\log n$ , ensuring efficient operations even in the worst case.

### **Advantages:**

**Efficient balancing:** Red-Black Trees maintain balance with relatively simple rules, compared to AVL Trees, which require stricter balance and more frequent rotations.

**Consistent performance:** The tree remains approximately balanced, providing predictable  $O(\log n)$  performance for all key operations.

**Low overhead for rebalancing:** Insertions and deletions require only a few rotations and recoloring operations, which are efficient.

### **Use Cases:**

**Balanced searching:** Used in databases, file systems, and language libraries where efficient searching and modifications are critical.

**Associative arrays:** Red-Black Trees are often used to implement map and set data structures (like in C++ STL or Java's `TreeMap` and `TreeSet`).

# AVL Trees:

An **AVL Tree** is a type of self-balancing binary search tree (BST) where the heights of the two child subtrees of any node differ by at most one. Named after its inventors **Adelson-Velsky and Landis**, the AVL tree maintains this balance through rotations during insertion and deletion, ensuring logarithmic time complexity for search, insertion, and deletion operations.

## Key Properties:

**Binary Search Tree Property:** Like all binary search trees, the left subtree contains nodes with keys less than the parent node, and the right subtree contains nodes with keys greater than the parent.

**Height-Balanced Property:** For any node in the tree, the difference in height between the left and right subtrees (called the **balance factor**) is at most 1. This ensures that the tree remains balanced and avoids degenerating into a linear structure like a linked list.

## Operations:

**Balance Factor:** The balance factor of a node is the difference between the height of the left subtree and the right subtree:

$$\text{Balance Factor} = \text{Height of Left Subtree} - \text{Height of Right Subtree}$$

If the balance factor is 1, 0, or -1, the node is balanced.

If the balance factor is outside this range (greater than 1 or less than -1), the tree needs to be rebalanced using rotations.

## Rotations:

To restore balance when the balance factor goes beyond the allowed range, AVL Trees use four types of rotations:

**Left Rotation** (LL Rotation): Applied when the right subtree of a node is taller, causing the tree to skew to the right.

**Right Rotation** (RR Rotation): Applied when the left subtree of a node is taller, causing the tree to skew to the left.

**Left-Right Rotation** (LR Rotation): A combination of left and right rotations, applied when the tree is first skewed left, then right.

**Right-Left Rotation** (RL Rotation): A combination of right and left rotations, applied when the tree is first skewed right, then left.

### **Key Operations:**

#### **Insertion:**

Insert the node as in a regular BST.

After insertion, check the balance factor of the affected nodes.

If the balance factor is violated, apply the appropriate rotation to rebalance the tree.

**Time complexity:**  $O(\log n)$  as the tree is kept balanced, and at most one rotation per insertion is required.

#### **Deletion:**

Delete the node as in a regular BST.

After deletion, check the balance factor of the affected nodes.

If the balance factor is violated, apply the appropriate rotation to restore balance.

**Time complexity:**  $O(\log n)$  since multiple rotations may be required, but the height of the tree remains logarithmic.

**Search:** Searching follows the same process as a regular binary search tree, where comparisons are made to decide whether to traverse the left or right subtree.

**Time complexity:**  $O(\log n)$ , since the tree's height is kept logarithmic due to the balancing property.

### **Example:**

Consider inserting the following keys into an AVL tree: 10, 20, 30, 40, 50, 25.

After each insertion, the AVL tree rebalances itself to maintain the height-balance property using rotations. For example:

After inserting 50, the tree might become skewed to the right, so a left rotation is applied to restore balance.

After inserting 25, a right-left (RL) rotation might be necessary to keep the tree balanced.

### **Performance:**

**Time complexity** for search, insertion, and deletion:  $O(\log n)$ .

The height of the AVL tree is strictly maintained as  $O(\log n)$ , ensuring fast access times even in the worst case.

### **Advantages:**

**Strict balancing:** AVL Trees are more strictly balanced than Red-Black Trees, ensuring better search performance in the worst case.

**Faster search performance:** The strict balance property results in slightly better search times compared to other balanced trees like Red-Black Trees.

**Height guarantee:** The maximum height of an AVL tree with  $n$  nodes is approximately  $1.44 \log n$ , which ensures logarithmic performance.

### **Disadvantages:**

**Higher overhead for insertion and deletion:** Maintaining strict balance means that AVL Trees often require more rotations than Red-Black Trees, leading to a slight overhead for insertions and deletions.

**More complex to implement:** The balancing logic and multiple types of rotations make AVL Trees more complex to implement compared to simpler balanced trees.

### **Use Cases:**

**Memory-constrained environments:** AVL Trees are preferred in memory-limited environments because they guarantee  $O(\log n)$  performance for all operations, minimizing the depth of the tree.

**Applications requiring frequent searches:** Ideal for situations where search performance is critical, such as in databases or in-memory data structures.

---

## ACTIVITY 2: LABORATORY ACTIVITY ( IMPLEMENTING AND ANALYZING AVL TREES )

**Objective :**



The purpose of this lab activity is to help students understand how AVL trees work, including the insertion and rotation processes to maintain balance. By the end of the activity, students should be able to:

### **Implement an AVL tree in Python.**

Understand and apply tree rotations (left, right, left-right, right-left).

Analyze the performance of AVL trees in terms of balance and efficiency.

Instructions:

**Environment:** Use either Jupyter Notebook or Colab to complete this lab activity.

Rename the notebook as `Lastname_Firstname_AVL_Lab.ipynb`.

**Setup:** Install the necessary Python packages if needed (e.g., matplotlib for visualizing trees).

**Submission:** Ensure you include your code, visualizations, and written answers for the analysis questions below. Submit your notebook as a `.ipynb` file through the eLearning platform (LMS).

#### Task 1: Implementing AVL Tree

**Create an AVL Tree Class:** Write a Python class `AVLTree` that includes the following methods:

`insert(key)`: Insert a new key into the AVL tree while maintaining balance using rotations.

`delete(key)`: Delete a key from the tree while maintaining the AVL balance.

`preorder()`, `inorder()`, and `postorder()`: Methods to print the tree in different traversal orders.

`height()`: Calculate the height of the AVL tree.

**Rotation Methods:** Your `AVLTree` class should include the following private methods to perform rotations:

`_left_rotate(node)`

`_right_rotate(node)`

`_left_right_rotate(node)`

`_right_left_rotate(node)`

Make sure each rotation method updates the tree structure correctly and keeps it balanced.

Task 2: Testing the AVL Tree

**Insert and Balance:**

Insert the following sequence of keys into your AVL Tree: 30, 40, 50, 20, 10, 5, 35, 25.

After each insertion, output the tree's **preorder traversal** and the **balance factors** for each node.

Visualize the tree after all insertions are complete (you can use text-based or graphical visualization).

**Analyze Rotations:**

Document and describe any rotations that occurred during the insertion process (e.g., Left Rotation at node 30).

Explain why the rotation was necessary and how it restored balance.

Task 3: AVL Tree Deletion

### Delete Nodes:

Delete the following keys from the AVL Tree: 20, 30, 35.

Output the tree's **preorder traversal** and **balance factors** after each deletion.

Document any rotations that occur during the deletion process and explain how they restored balance.

### Task 4: Analyzing AVL Tree Performance

#### Height Analysis:

Calculate the height of the AVL Tree after the insertions and deletions are complete.

Compare the height of the AVL Tree to the height of an unbalanced binary search tree (e.g., a tree formed by inserting the same keys in ascending order without balancing).

#### Efficiency:

Estimate the time complexity for search, insertion, and deletion in your AVL Tree and compare it to an unbalanced binary search tree.

Discuss the advantages of using an AVL Tree in applications where frequent insertions and deletions occur.

---

### ACTIVITY 3: COMPARING DATA STRUCTURES FOR SET MANIPULATION[DISCUSSION FORUM ACTIVITY]

#### Choosing the Right Data Structure: Queues, Stacks, Binary Search Trees, Red-Black Trees, and AVL Trees

#### Objective:

In this discussion, students will critically compare and contrast different data structures used for set manipulation. The focus will be on understanding the appropriate use cases for **Queues**, **Stacks**, **Binary Search Trees (BSTs)**, **Red-Black Trees**, and **AVL Trees** in terms of performance, complexity, and application in real-world scenarios.

### Instructions:

#### Initial Post:

Research and explain the key differences between two of the data structures mentioned (you may choose any pair: Queues vs Stacks, BSTs vs AVL Trees, Red-Black Trees vs AVL Trees, etc.).

Discuss which data structure is better suited for specific types of problems (e.g., searching, insertion, deletion, balancing).

Provide examples of real-world applications where one data structure might outperform the other.

Your initial post should be at least **250 words** and include references to any sources you use.

#### Follow-up Discussion :

Respond to **two** of your peers' posts. In your response:

Ask questions to further the discussion or challenge their perspective.

Provide an alternative viewpoint or real-world example where a different data structure could be more efficient.

Your responses should be at least **100 words** each.

#### Example of Discussion Points:

Why might **AVL Trees** be preferred over **Binary Search Trees (BSTs)** for applications requiring frequent insertions and deletions?

In what scenarios would a **Queue** be more efficient than a **Stack**?

How do the self-balancing properties of **Red-Black Trees** compare to those of **AVL Trees** in terms of practical performance?

---

#### FURTHER READING AND REFERENCES

For further information and references on the topic refer to:

1. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2022). *Introduction to algorithms* (4th ed.). The MIT Press.
2. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>
3. Dimri, S. C., Malik, P., Ram, M. (2021). *Algorithms: Design and Analysis*. Germany: De Gruyter.

---

#### ACTIVITY 4: End of the module Quiz

1. Which of the following data structures follows the First-In-First-Out (FIFO) principle?

- a) Stack
- b) Queue**
- c) Binary Search Tree
- d) AVL Tree

**2. What is the time complexity of searching for an element in a balanced Binary Search Tree (BST)?**

- a)  $O(n)$
- b)  $O(\log n)$**
- c)  $O(n \log n)$
- d)  $O(1)$

**3. In a stack, which operation removes an element from the data structure?**

- a) Enqueue
- b) Dequeue
- c) Push
- d) Pop**

**4. What is the primary advantage of an AVL tree compared to a standard Binary Search Tree (BST)?**

- a) Simplicity of insertion
- b) Ensures balance for improved search performance**
- c) Lower memory usage
- d) Faster deletion

**5. Which of the following statements is true about Red-Black Trees?**

- a) Every path from the root to a leaf has the same number of black nodes.
- b) Every path from the root to a leaf has the same number of red nodes.
- c) All nodes are either red or black, and leaf nodes are always black.**
- d) Red-Black Trees are always perfectly balanced.

**6. What is the key difference between a Binary Search Tree (BST) and a Randomized Search Tree (RST)?**

a) BSTs are unbalanced, while RSTs are always balanced.

**b) In RSTs, node positions are determined randomly.**

c) BSTs do not require rotations, but RSTs do.

d) RSTs are always faster for insertions and deletions.

**7. In an AVL tree, what is the maximum difference allowed in the heights of the left and right subtrees of any node?**

**a) 1**

b) 2

c) 3

d) No limit

**8. Which of the following operations on a queue is responsible for adding an element?**

a) Push

**b) Enqueue**

c) Pop

d) Peek

**9. In a Red-Black Tree, which property helps maintain the tree's balance?**

a) Left-leaning property

**b) No consecutive red nodes**

c) Each node has at most two children

d) No black nodes can have two black children

**10. Which data structure would be most suitable for implementing an undo operation in a text editor?**

a) Queue

**b) Stack**

c) Binary Search Tree

d) AVL Tree

MODULE 9: ADVANCED TOPICS IN ALGORITHMS

---



---

## MODULE OVERVIEW

This module explores **Advanced Topics in Algorithms**, focusing on critical algorithms used in signal processing and pattern matching. The first section introduces the **Fast Fourier Transform (FFT)**, a powerful tool in both theoretical and practical applications like signal processing. The second part of the module deals with **Pattern Matching Algorithms**, which are fundamental in string searching. These algorithms are crucial for a wide range of applications, from text editors to bioinformatics.

---

## LEARNING OUTCOMES

By the end of this module, students will be able to:

1. Understand the principles of the Fast Fourier Transform (FFT) and its applications in signal processing.
2. Explain how FFT improves efficiency compared to the Discrete Fourier Transform (DFT).
3. Analyze and implement different pattern matching algorithms including Naive string matching, Knuth-Morris-Pratt (KMP), and Boyer-Moore.
4. Evaluate the time complexity and efficiency of these algorithms in various contexts, such as large-scale text searches.

---

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Perform Lab activity

### 3. Complete End of Module Quiz

---

## REVIEW READING MATERIAL

### Fast Fourier Transform (FFT)

The **Fast Fourier Transform (FFT)** is an efficient algorithm for computing the **Discrete Fourier Transform (DFT)** and its inverse. The Fourier Transform is fundamental in signal processing, as it transforms a signal from its time domain into its frequency domain. The FFT significantly reduces the computational complexity of performing a DFT from  $O(n^2)$  to  $O(n \log n)$ , where  $n$  is the number of data points in the signal.

#### 1. Introduction to FFT

The **Discrete Fourier Transform (DFT)** is used to analyze the frequency components of a discrete signal. It expresses a sequence of  $n$  complex numbers  $x_0, x_1, \dots, x_{n-1}$  as a sum of sinusoids, breaking the signal down into its frequency components. Mathematically, the

DFT of a sequence is defined as:

$$X_k = \sum_{n=0}^{N-1} x_n e^{-2\pi i k n / N}$$

where:

- $x_n$  represents the input signal in the time domain,
- $X_k$  is the transformed signal in the frequency domain,
- $N$  is the total number of points,
- $i$  is the imaginary unit, and
- $k$  corresponds to the index of the frequency component.

FFT is a faster algorithm to compute this DFT, and it works by recursively breaking down a DFT of any composite size  $n = N_1 \times N_2$  into smaller DFTs of sizes  $N_1$  and  $N_2$ .

## 2. Applications in Signal Processing

FFT is widely used in digital signal processing (DSP) and has numerous practical applications:

**Audio and Music Processing:** FFT helps analyze and process sound signals by breaking them down into frequency components. It is used for equalizers, sound compression algorithms (like MP3), and identifying musical notes in a signal.

**Image Processing:** In image processing, FFT is used to process images by converting spatial data into frequency data. This allows for image filtering, noise reduction, and edge detection.

**Communication Systems:** In modern communication, FFT is essential for converting signals into their frequency components for modulation and demodulation in systems such as Orthogonal Frequency Division Multiplexing (OFDM), used in Wi-Fi and LTE.

**Seismic Data Analysis:** FFT is used to analyze earthquake data and other geophysical signals by transforming the data into the frequency domain to identify hidden patterns.

**Radar and Sonar:** FFT is applied in radar and sonar systems to detect objects by analyzing the frequency shift of reflected signals.

## Pattern Matching Algorithms

### 1. Naive String Matching Algorithm

The **Naive String Matching Algorithm** is the simplest string matching approach. It checks for the presence of a substring (pattern) in a larger string (text) by sliding the pattern one character at a time and comparing characters.

#### Steps:

Slide the pattern from the start to the end of the text.

At each position, compare the pattern with the substring of the text.

If all characters match, return the position. If not, move to the next position.

#### Time Complexity:

Worst-case:  $O((n - m + 1) * m)$  where  $n$  is the length of the text and  $m$  is the length of the pattern.

Example:

Text: "ABCABCD"

Pattern: "ABC"

Matches found at index: 0 and 4

## 2. Knuth-Morris-Pratt (KMP) Algorithm

The KMP Algorithm improves on the naive approach by reducing redundant comparisons. It preprocesses the pattern to determine the "longest proper prefix which is also a suffix" (LPS) array, which helps avoid unnecessary shifts during matching.

### Steps:

Preprocess the pattern to compute the LPS array.

Use the LPS array to skip some comparisons while sliding the pattern over the text.

### Key Idea:

If a mismatch occurs after matching  $k$  characters, the KMP algorithm uses the LPS array to determine where the next comparison should begin.

### Time Complexity:

Both preprocessing and matching have linear time complexity:  $O(n + m)$ .

### Example:

Text: "ABABABCA"

Pattern: "ABAB"

LPS Array: [0, 0, 1, 2]

Matches found at index: 0 and 2

### 3. Boyer-Moore Algorithm

The Boyer-Moore Algorithm is one of the most efficient string matching algorithms in practice. It preprocesses the pattern and uses two heuristics to skip sections of the text: the Bad Character Heuristic and the Good Suffix Heuristic.

#### Steps:

Use the Bad Character Heuristic to skip ahead when a mismatch occurs by aligning the last occurrence of the mismatched character in the pattern with the mismatched text character.

Use the Good Suffix Heuristic when the last characters match, skipping sections based on matched suffixes.

#### Time Complexity:

Best-case:  $O(n/m)$

Worst-case:  $O(n * m)$  (less common due to the heuristics).

#### Example:

Text: "HERE IS A SIMPLE EXAMPLE"

Pattern: "EXAMPLE"

Bad Character Heuristic helps skip several unnecessary comparisons.

Match found at index: 17

---

## ACTIVITY 2: LABORATORY ON ADVANCED TOPICS IN ALGORITHMS

### Objective:

To implement the Fast Fourier Transform (FFT) and various Pattern Matching Algorithms (Naive, KMP, and Boyer-Moore) in Python, and understand their applications and efficiencies.

### Activity Structure:

#### Fast Fourier Transform (FFT)

Task: Implement the Fast Fourier Transform algorithm from scratch. Use a simple example to demonstrate its application in signal processing.

Instructions:

Write a Python function that implements the FFT algorithm.

Create a test signal (e.g., a sine wave) and use FFT to analyze its frequency components.

Plot the original signal and its frequency spectrum using Matplotlib.

Output: A Python script (fft\_last\_firstname.py) and a plot showing the original signal and its FFT result.

Expected Summary: Explain the FFT algorithm and its importance in signal processing.

#### Pattern Matching Algorithms

Task: Implement the Naive string matching algorithm, Knuth-Morris-Pratt (KMP) algorithm, and Boyer-Moore algorithm in Python.

Instructions:

Create three separate functions for each algorithm.

Test these functions with a sample text and pattern, and display the indices where matches are found.

Discuss the time complexity of each algorithm and under what scenarios each would be preferred.

Output: A Python script (**pattern\_matching\_last\_firstname.py**) that contains all three algorithms and test cases for each.

Expected Summary: Compare the efficiencies of the three algorithms and provide insights into their use cases in text processing.

### **Deliverables:**

Two Python scripts:

**fft\_last\_firstname.py** (for FFT)

**pattern\_matching\_last\_firstname.py** (for pattern matching)

A brief report (1-2 pages) summarizing your findings and experiences with the implementations.

### **Evaluation Criteria:**

Implementation: Correctness and efficiency of the algorithms.

Documentation: Clarity of comments and structure of the code.

Analysis: Depth of understanding demonstrated in the summaries.

### **Submission Instructions:**

Save your files as **fft\_last\_firstname.py** and **pattern\_matching\_last\_firstname.py**.



**Submit the scripts and the report on LMS by the due date.**

---

## FURTHER READING AND REFERENCES ON ADVANCED ALGORITHMS

For further information and references on the topic refer to:

1. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2022). *Introduction to algorithms* (4th ed.). The MIT Press.
  - Fast Fourier Transform (pp. 889-892)
  - String Matching (p. 957)
2. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>
3. Dimri, S. C., Malik, P., Ram, M. (2021). *Algorithms: Design and Analysis*. Germany: De Gruyter.

---

## ACTIVITY 3: END OF MODULE QUIZ

**What is the primary advantage of the Knuth-Morris-Pratt (KMP) algorithm over the naive string matching algorithm?**

- A) It uses more memory.
- B) It has a lower worst-case time complexity.**
- C) It is easier to implement.
- D) It requires a longer pattern.

**In the Boyer-Moore algorithm, which heuristic is used to skip sections of the text when a mismatch occurs?**

- A) Longest Common Prefix
- B) Bad Character Heuristic**

- C) Good Suffix Heuristic
- D) Prefix Function

**What is the time complexity of the naive string matching algorithm in the worst case?**

- A)  $O(n)$
- B)  $O(m)$
- C)  $O(n * m)$**
- D)  $O(n + m)$

**In the Fast Fourier Transform (FFT), what is the primary use of the algorithm?**

- A) Sorting strings
- B) Analyzing frequency components of signals**
- C) Searching for patterns in texts
- D) Encrypting data

**7. Which of the following algorithms is NOT a pattern matching algorithm?**

- A) Naive String Matching
- B) Knuth-Morris-Pratt (KMP)
- C) Boyer-Moore
- D) Dijkstra's Algorithm**

**6. The LPS array in the KMP algorithm stands for:**

- A) Longest Proper Suffix
- B) Longest Prefix Suffix**
- C) Longest Pattern Shift

D) Last Pattern Suffix

**7. In which scenario would the Boyer-Moore algorithm be most efficient?**

A) Searching in a very short text.

**B) Searching for patterns in large texts with many unique characters.**

C) Searching in texts with repeated patterns only.

D) Searching in binary data.

**8. The Fast Fourier Transform is most commonly used in:**

A) Text mining

B) Image processing

**C) Signal processing**

D) Database querying

**9. What is the time complexity of the Boyer-Moore algorithm in the average case?**

A)  $O(n)$

B)  $O(m)$

**C)  $O(n/m)$**

D)  $O(n + m)$

**10. Which of the following statements is true regarding the KMP algorithm?**

A) It is slower than the naive algorithm.

**B) It avoids unnecessary comparisons by using the LPS array.**

C) It can only be used for fixed-length patterns.

D) It requires preprocessing of the text, not the pattern.

# MODULE 10: NP-COMPLETE PROBLEMS

---

INSTRUCTIONAL HOURS: 4, 5 or 6

## MODULE OVERVIEW

In this module, students will explore the concept of NP-completeness, a fundamental topic in computational theory. The module will cover essential definitions, the significance of NP-complete problems, and the importance of polynomial-time reductions. Students will delve into various NP-complete problems, learning to identify them and apply appropriate algorithms to tackle these challenges. Through practical examples and case studies, students will gain insights into heuristic and approximation techniques used in solving NP-complete problems.

## LEARNING OUTCOMES

By the end of this module, students will be able to:

1. Explain the concept of NP-completeness and recognize NP-complete problems within the context of computational theory.
2. Classify problems as NP-complete by understanding their characteristics and implications.
3. Perform polynomial-time reductions to demonstrate the NP-completeness of specific problems.
4. Analyze real-world case studies that illustrate the applications of NP-complete problems.

---

## LEARNING ACTIVITIES/TASK LIST



1. Review Reading Material
2. Perform Lab activity
3. Complete End of Module Quiz

---

## REVIEW READING MATERIAL

### Introduction to NP-Completeness

#### Definition and Background

NP-completeness is a classification used in computational theory to describe certain problems that are, informally speaking, the hardest problems in the class NP (nondeterministic polynomial time).

A problem is in NP if a solution can be verified in **polynomial time**. A problem is NP-complete if:

It is in NP.

Every problem in NP can be reduced to it in polynomial time.

**Recall that:**

**A** problem is said to be solvable in polynomial time if its solution can be computed in time proportional to a polynomial expression in the size of the input.

Characteristics of Polynomial-Time Algorithms

Efficiency: Polynomial-time algorithms are generally considered efficient, as their running time grows at a manageable rate.

Examples: Algorithms like sorting (e.g., quicksort, mergesort) and basic arithmetic operations (addition, multiplication) run in polynomial time.

**Note that:** The concept of verification is central to understanding NP-completeness because it establishes that while a problem may be difficult to solve, checking a solution can be efficient.

**Significance** NP-completeness

Understanding NP-completeness is crucial for recognizing the limitations of algorithmic problem-solving. If a polynomial-time solution exists for any NP-complete problem, it implies that all problems in NP can be solved in polynomial time ( $P = NP$ ).

**The Cook-Levin Theorem**

Proposed by Stephen Cook in 1971, this theorem established the first NP-complete problem: the Boolean satisfiability problem (SAT). It showed that if SAT can be solved in polynomial time, all NP problems can be solved in polynomial time.

**NP-Completeness and Reducibility**

**Definition of Reductions**

A reduction is a way of converting one problem into another, showing that if one problem can be solved efficiently, so can the other.

If problem A can be reduced to problem B in polynomial time, solving problem B would also solve problem A.

## Types of Reductions

**Many-One Reduction:** Transforming instances of one problem to instances of another, maintaining yes/no answers.

**Turing Reduction:** More general, allowing the use of algorithms for the second problem during the solution of the first.

## Significance

Reductions are used to demonstrate that a problem is NP-complete by showing that all problems in NP can be reduced to it.

## Examples of NP-Complete Problems

NP-complete problems share the property that if any NP-complete problem can be solved in polynomial time, then all problems in NP can be solved in polynomial time.

Here are some well-known NP-complete problems:

### Boolean Satisfiability Problem (SAT)

Given a Boolean formula, the task is to determine if there exists an assignment of variables that makes the formula true.

### 3-SAT

A specific case of SAT where the formula is in conjunctive normal form with exactly three literals per clause.

## Vertex Cover

Given a graph and an integer  $k$ , the problem asks whether there is a subset of vertices of size  $k$  such that each edge in the graph is incident to at least one vertex in the subset.

## Clique Problem

Given a graph and an integer  $k$ , the problem is to determine whether there is a clique (a subset of vertices all adjacent to each other) of size  $k$ .

## Traveling Salesman Problem (Decision Version)

Given a set of cities and distances between them, and a number  $C$ , the question is whether there is a tour that visits each city exactly once and has a total distance less than or equal to  $C$ .

## Hamiltonian Cycle

Determine whether a cycle exists in a given graph that visits each vertex exactly once.

## Subset Sum Problem

Given a set of integers and a target sum, the problem is to determine if any subset of the integers sums to the target.

These examples illustrate the diversity of NP-complete problems, spanning various domains such as graph theory, optimization, and decision-making.

## NP-Completeness Proofs

### Structure of Proofs

Proving that a problem is NP-complete typically involves two main steps:



**Proving the Problem is in NP:** Show that a solution can be verified in polynomial time.

**Proving it is NP-hard:** Show that at least one known NP-complete problem can be reduced to this problem in polynomial time.

### Example of a Proof

A common example is the proof of NP-completeness for the 3-SAT problem. The proof shows that:

3-SAT is in NP because a satisfying assignment can be verified quickly.

It is NP-hard by demonstrating a polynomial-time reduction from another NP-complete problem, such as SAT.

### Implications of NP-Completeness

Once a problem is proven NP-complete, it is widely accepted that no polynomial-time algorithm exists for it unless  $P = NP$ .

### Applications

NP-complete problems appear in various fields such as scheduling, resource allocation, and network design, making their study significant for practical applications in computer science and operations research.

## 3. Approximation Algorithms for NP-Hard Problems

### Definition of NP-Hard

NP-hard problems are at least as hard as NP-complete problems, meaning they may not even be in NP, as they do not have polynomial-time verifiable solutions.

### Importance of Approximation Algorithms

Since many NP-complete and NP-hard problems cannot be solved in polynomial time, approximation algorithms provide a practical means of obtaining solutions that are "good enough" within reasonable time frames.

## Types of Approximation Algorithms

### Polynomial-Time Approximation Schemes (PTAS)

A family of algorithms that can find solutions within a factor of  $(1+\epsilon)$  of the optimal solution for any small constant  $\epsilon > 0$ , but may take exponential time in the worst case.

### Fully Polynomial-Time Approximation Schemes (FPTAS)

These schemes guarantee that the running time is polynomial in both the input size and  $1/\epsilon$ , making them more efficient than PTAS.

### Constant Factor Approximation Algorithms

Provide solutions guaranteed to be within a constant factor of the optimal solution.

### Examples of Approximation Algorithms

**Greedy Algorithms** for problems like the **Knapsack Problem**, which can provide good solutions in a fraction of the time.

**Local Search Algorithms** for problems such as the **Traveling Salesman Problem**, where a solution is iteratively improved by making small changes.

### Challenges in Designing Approximation Algorithms

One of the main challenges is to establish bounds on the performance of the algorithm compared to the optimal solution.

Determining whether an approximation algorithm is feasible often involves exploring trade-offs between accuracy and computational efficiency.

---

## ACTIVITY 2: LABORATORY ACTIVITY ON NP-COMPLETE PROBLEMS (GROUP LAB ACTIVITY)

### **Activity Title: Exploring NP-Complete Problems and Reductions**

#### **Objective:**

To understand NP-completeness through hands-on experience with specific NP-complete problems, perform reductions, and implement basic algorithms to solve instances of these problems.

#### **Materials Needed:**

Computer with Python installed

Access to a Python IDE (like Jupyter Notebook or PyCharm)

Sample datasets for NP-complete problems (provided as CSV files)

#### **Activity Steps:**

##### **Problem Selection**

Divide students into small groups.

Each group will choose one NP-complete problem from the following:

1. 3-SAT
2. Vertex Cover
3. Subset Sum

## Understanding the Problem

Each group will research their chosen problem and answer the following questions:

What is the formal definition of the problem?

Why is it classified as NP-complete?

What are some real-world applications of this problem?

## Performing Reductions

Each group will demonstrate a polynomial-time reduction from one known NP-complete problem to their chosen problem. For example:

Reduce the Vertex Cover problem to the Clique problem.

Reduce the Subset Sum problem to the Knapsack problem.

Document the reduction process, including input transformations and explanations.

## Algorithm Implementation

Groups will implement a basic algorithm to solve their chosen NP-complete problem using Python. Suggested implementations:

For **3-SAT**: Write a simple backtracking algorithm to find a satisfying assignment.

For **Vertex Cover**: Implement a greedy algorithm to find a minimal vertex cover.

For **Subset Sum**: Write a recursive function to determine if there is a subset that sums to a target value.

Each group should provide comments in their code to explain their logic.

### **Submission Requirements:**

A Jupyter Notebook or Python script containing:

A detailed explanation of the chosen NP-complete problem.

Documentation of the reduction process.

The implemented algorithm code with comments.

Sample inputs and outputs demonstrating the functionality of the algorithm.

---

### FURTHER READING AND REFERENCES ON NP-COMPLETENESS

For further information and references on the topic refer to:

4. Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2022). *Introduction to algorithms* (4th ed.). The MIT Press.
  - NP-Completeness (p. 1042)
  - Approximation Algorithms (p. 1104)
5. Khan Academy. (n.d.). *Algorithms*.  
<https://www.khanacademy.org/computing/computer-science/algorithms>
6. Dimri, S. C., Malik, P., Ram, M. (2021). *Algorithms: Design and Analysis*. Germany: De Gruyter.

---

### ACTIVITY 3: END OF MODULE QUIZ

**1. Which of the following problems is classified as NP-complete?**

A) Sorting

**B) 3-SAT**

C) Searching

**2. What does NP stand for in NP-completeness?**

A) Non-Polynomial

B) Non-Procedural

**C) Nondeterministic Polynomial**

D) None of the above

**3. Which of the following statements is true about NP-complete problems?**

A) They can be solved in polynomial time.

**B) They can be verified in polynomial time.**

C) They are the easiest problems in NP.

D) They do not have real-world applications.

**4. The Cook-Levin theorem is significant because it:**

A) Proves that all NP problems are solvable in polynomial time.

**B) Establishes the first NP-complete problem, the Boolean satisfiability problem.**

- C) Introduces the concept of Turing machines.
- D) Provides a method for solving NP-complete problems.

**5. A polynomial-time reduction is used to:**

- A) Show that a problem is in P.
- B) Demonstrate that one problem is at least as hard as another.**
- C) Solve NP-complete problems in polynomial time.
- D) Verify the solutions to NP-complete problems.

**6. Which of the following is NOT an NP-complete problem?**

**A) Sorting**

- B) Hamiltonian Cycle
- C) Vertex Cover
- D) Subset Sum

**7. Which of the following algorithms is commonly used to find a satisfying assignment in the 3-SAT problem?**

- A) Greedy algorithm
- B) Backtracking algorithm**
- C) Dynamic programming
- D) Divide and conquer

**8. The Subset Sum problem can be described as:**

- A) Finding a subset of elements that equals a given target sum.

**B) Determining if any subset of integers sums to a specified target.**

C) Finding the largest subset of elements.

D) Identifying unique elements in a set.

**9. An approximation algorithm is used when:**

A) A polynomial-time solution is guaranteed.

**B) An exact solution is difficult or impossible to find efficiently.**

C) The problem is in P.

D) The solution must be optimal.

**10. What is the primary characteristic of NP-hard problems?**

A) They are always solvable in polynomial time.

**B) They are at least as hard as NP-complete problems.**

C) They can be solved using dynamic programming.

D) They can be verified in polynomial time.



