

Deep Dive Digital Notes

DSA ***DEEP DIVE NOTES*** ***NOTES***



150+ Pages

INDEX

SR NO.		CHAPTERS	PAGE NO.
1		CHAPTER 1 INTRODUCTION	5
	1.1	Operations on Data Structure	
	1.2	Complexity	
	1.3	Dynamic Memory Allocations	
	1.4	Arrays	
	1.5	Traversing Linear Array	
	1.6	Multidimensional Array	
2		CHAPTER 2 Searching and Sorting Technique	30
	2.1	Searching	
	2.2	Hashing	
	2.3	Sorting	
	2.4	Sorting Techniques	
3		CHAPTER 3 Stack ,Queues &Recursion	56
	3.1	Stacks	
	3.2	Polish Notations	
	3.3	Recursion	
	3.4	Tower of Hanoi	
	3.5	Queues	
4		CHAPTER 4 Linked Lists	77

	4.1	Introduction	
	4.2	Creating List	
	4.3	Memory Allocations	
	4.4	Inserting in Linked List	
	4.5	Header Linked List	
5		CHAPTER 5 Trees	98
	5.1	Trees Terminologies	
	5.2	Trees Types	
	5.3	Binary Trees Traversal Method	
	5.4	Binary Trees Operation	
	5.5	Heap	
6		CHAPTER 6 Graphs	129
	6.1	Introduction	
	6.2	Linear Representation of Graph	
	6.3	Warshall's Algorithm	
	6.4	Applications of Graph	

About

Welcome to "DSA USING C++ DEEP DIVE DIGITAL NOTES" your comprehensive guide to mastering DSA, from basics to advanced concepts. This book is meticulously compiled from diverse sources, including official documentation, insights from the vibrant Stack Overflow community, and the assistance of AI ChatBots.

Disclaimer

This book is an unofficial educational resource created for learning purposes. It is not affiliated with any official group(s) or company(s), nor is it endorsed by Stack Overflow. The content is curated to facilitate understanding and skill development in HTML programming.

Contact

If you have any questions, feedback, or inquiries, feel free to reach out to us at contact@codewithcurious.com. Your input is valuable, and we are here to support your learning journey.

Copyright

© 2024 CodeWithCurious. All rights reserved. No part of this book may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of the publisher, except in the case of brief quotations embodied in critical reviews and certain other noncommercial uses permitted by copyright law.

Happy Coding!

CHAPTER 1: INTRODUCTION TO DATA STRUCTURE



1.1 Introduction To Data Structure & Arrays

Data structures and arrays are fundamental concepts in computer science, playing a crucial role in organizing and storing data efficiently. A data structure is a specialized format for organizing and storing data in a computer's memory.

It defines the relationships between different data elements and the operations that can be performed on them. Arrays, one of the simplest and most widely used data structures, are collections of elements stored in contiguous memory locations.

They provide a way to represent and manipulate a fixed-size sequential collection of elements, each identified by an index or a key.

Arrays offer fast and direct access to individual elements, making them efficient for tasks like searching and sorting. However, their fixed size can be limiting, and inserting or deleting elements can be time-consuming.

To overcome these limitations, various advanced data structures, such as linked lists, trees, and hash tables, have been developed. These structures provide dynamic allocation of memory, enabling efficient insertion and deletion operations.

Understanding data structures is essential for designing efficient algorithms and optimizing program performance. Choosing the right data structure for a specific problem is a critical decision that impacts the efficiency and scalability of the solution.

Arrays and data structures collectively form the backbone of software development, enabling the creation of robust and scalable applications that can handle diverse data processing tasks.

Summary

- Data structures and arrays are foundational concepts in computer science, serving as vital tools for organizing and storing data efficiently.
- Arrays, simple yet widely used, represent collections of elements in contiguous memory locations, offering fast access but with fixed sizes.

- To address limitations like dynamic resizing, more advanced data structures such as linked lists and trees have been developed.
- Understanding these structures is crucial for designing efficient algorithms and optimizing program performance, as they form the basis for creating scalable and robust applications capable of handling diverse data processing tasks.

1.2 Basic Terminology

In the realm of computer science, basic terminology in data structures and arrays includes fundamental concepts such as "data," which represents information or values, and "data structures," organizational frameworks for storing and manipulating data efficiently.

Elements, nodes, and indices are the building blocks, referring to individual units, basic components, and positional identifiers, respectively. Arrays, specific collections of elements stored in contiguous memory locations, rely on terms like element access, size, and initialization, reflecting the crucial aspects of retrieval, quantity, and initial assignment.

Common operations, like insertion, deletion, search, and traversal, describe fundamental manipulations applied to these structures, forming the essential vocabulary for comprehending and working with data in computer programming and software development.

list of basic terminology related to data structures and arrays:

Data Structure

- 1. Data:** information or values that can be processed or manipulated.
- 2. Data Structure:** A way of organizing and storing data to perform operations efficiently.
- 3. Element:** A single unit of data within a data structure.
- 4. Node:** Basic building block in many data structures, containing data and references.
- 5. Index:** Positional identifier for accessing elements in an ordered structure.

Arrays

- 1. Array:** A collection of elements stored in contiguous memory locations.
- 2. Element Access:** Retrieving data from a specific position in an array.
- 3. Index:** Numeric identifier indicating the position of an element.
- 4. Size/Length:** Number of elements in an array.
- 5. Initialization:** Assigning initial values to array elements.

6. Traversal: Iterating through all elements in an array.

Common Operations

1. Insertion: Adding an element to a data structure.

2. Deletion: Removing an element from a data structure.

3. Search: Locating a specific element in a data structure.

4. Traversal: Visiting each element in sequence.

5. Sorting: Arranging elements in a specific order.

6. Merging: Combining two or more data structures.

Understanding these basic terms lays the foundation for exploring more advanced concepts in data structures and arrays in computer science and programming.

1.3 Elementary Data Structures

Elementary data structures include stacks and queues. Stacks operate on the Last In, First Out (LIFO) principle, where elements are added and removed from the same end, known as the top.

This makes them suitable for tasks such as function call management and parsing expressions. Queues, on the other hand, adhere to the First In, First Out (FIFO) principle, with elements added to the back and removed from the front.

Queues are commonly employed in scenarios like task scheduling and breadth-first search algorithms.

Hash tables, a key data structure, use a hash function to map keys to indices for efficient data retrieval and storage.

They play a vital role in scenarios where quick access to data based on a unique key is crucial, offering a balance between space and time efficiency.

Overall, understanding these elementary data structures is foundational for computer science and software development, as they provide the groundwork for designing more complex algorithms and handling various data processing tasks efficiently.

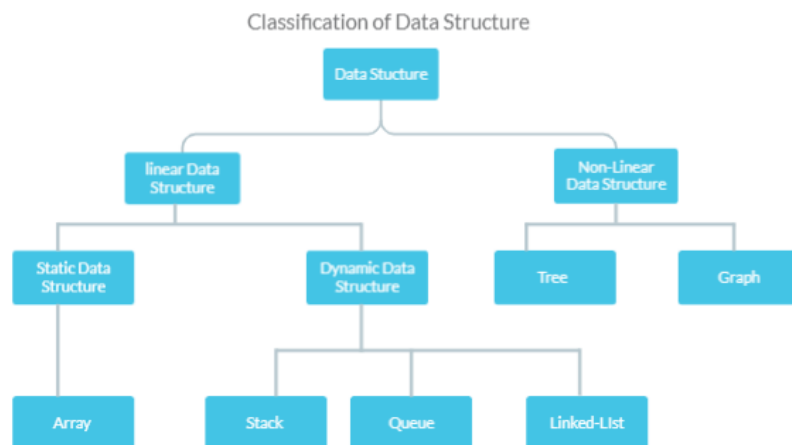
Summary

- Elementary data structures are fundamental components in computer science, laying the groundwork for efficient data organization and manipulation.
- Arrays offer direct access to elements, linked lists provide dynamic insertion and deletion, while stacks and queues operate on Last In, First Out (LIFO), and First In, First

Out (FIFO) principles, respectively.

- Trees and hash tables introduce hierarchical and indexed structures, expanding the range of data relationships that can be represented. Additionally, hash tables leverage hash functions for swift data retrieval based on unique keys.
- These structures are essential for algorithm design and software development, enabling effective data handling and organization in diverse computational scenarios.

1.4 Organization



The organization of data structures is a crucial aspect of computer science that involves designing and implementing efficient systems for storing, managing, and accessing data. Elementary data structures, such as arrays and linked lists, provide the fundamental building blocks for this organization.

Arrays, a basic and widely used structure, organize elements in contiguous memory locations, and each element is accessed using an index. This direct access allows for efficient retrieval and manipulation, making arrays suitable for tasks that involve frequent element access.

Linked lists, on the other hand, use nodes that contain both data and a reference to the next node, allowing for dynamic insertions and deletions. This flexibility is particularly valuable in scenarios where the size of the data is not known in advance, or when frequent modifications are required.

As the complexity of data and the requirements of algorithms grow, more advanced structures become necessary. Trees, for instance, introduce hierarchical relationships among elements.

In a binary tree, each node has at most two children, forming a structure that is conducive to efficient searching and sorting operations. Trees can be balanced, ensuring that the tree remains relatively even on both sides, which improves search and retrieval times.

Hash tables represent yet another level of sophistication in data organization. They use a hash function to map keys to specific locations in memory, allowing for rapid retrieval of data based on those keys. This makes hash tables highly efficient for tasks like dictionary implementations, where quick lookup times are essential.

Proficiency in organizing data structures is foundational for effective problem-solving in computer science and software development. It enables developers to choose the right structures for specific tasks, optimize algorithms for speed and memory usage, and create systems that can handle diverse and complex data processing requirements.

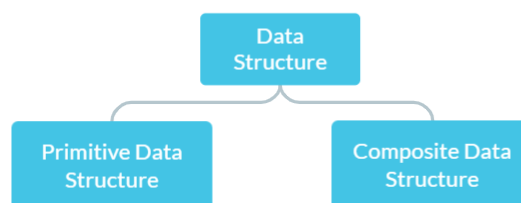
Summary

- The efficient organization of data structures is a crucial aspect of computer science, with arrays offering direct access for efficient retrieval and manipulation, making them suitable for tasks involving frequent element access.
- Linked lists provide dynamic insertions and deletions, valuable when data size is unknown or frequent modifications are required.
- Trees introduce hierarchical relationships, and balanced binary trees improve search and retrieval.
- Hash tables, using a hash function, enable rapid data retrieval based on keys, making them efficient for tasks like dictionary implementations.
- Proficiency in these structures is foundational for effective problem-solving, enabling developers to choose the right structures, optimize algorithms, and handle diverse and complex data processing requirements in software development.

1.5 Classification of Data Structure

Data structures can be broadly classified into two main categories.

Primitive data structures and composite data structures.



Primitive data structures are the fundamental building blocks that directly operate on machine instructions, including integer, float, character, and boolean types. These basic types represent individual elements and serve as the foundation for more complex data structures.

Composite data structures, on the other hand, are constructed from primitive data types and are designed to efficiently organize and manage larger sets of data. Examples of composite data structures include arrays, which store elements of the same type in

contiguous memory locations, linked lists that allow for dynamic insertions and deletions, and trees that establish hierarchical relationships among elements.

Other composite structures include stacks, queues, graphs, and hash tables, each serving specific purposes and offering unique advantages in terms of data organization and access.

The classification of data structures into these two categories provides a conceptual framework for understanding their fundamental nature and aids in selecting the most appropriate structure for solving specific computational problems based on the characteristics of the data and the requirements of the algorithms involved.

Primitive Data Structures

Primitive data structures are the fundamental building blocks of computer programming, representing basic data types that directly operate on machine instructions. These include integer, float, character, and boolean types. Integer data type deals with whole numbers, float handles decimal numbers, character represents individual characters, and boolean is used for true or false values.

These primitive data types are essential for storing and manipulating individual elements, forming the foundational layer upon which more complex data structures and algorithms are built. Their simplicity and direct representation make them crucial for performing fundamental operations in programming, and they serve as the basis for higher-level abstractions in computer science.

Type of primitive data type

1. Integer: The integer data type is used to represent whole numbers without any decimal places. It includes both positive and negative values. In programming languages, integers typically have a specific range defined by the number of bits used to store them (e.g., 32-bit or 64-bit integers).

2. Float: Float, short for floating-point, is a data type used to represent decimal numbers. It includes both the integer and fractional parts of a real number. Floating-point numbers are versatile for tasks involving precision, but they may have limitations due to the finite representation of decimal values in binary.

3. Character: The character data type is used to represent individual characters, such as letters, digits, or symbols. In many programming languages, characters are usually represented using ASCII or Unicode encoding, assigning numeric values to each character.

4. Boolean: The boolean data type has only two possible values: true or false. Booleans are fundamental for logical operations and conditional statements in programming. They are often used to control the flow of a program based on certain conditions.

Each of these primitive data types plays a crucial role in programming, providing a foundation for representing different kinds of values and enabling the execution of a wide range of computational tasks. Programmers use these data types judiciously to ensure the efficient and accurate manipulation of data within their programs.

Summary

- Primitive data structures are foundational elements in computer programming, representing essential data types that directly interact with machine instructions.
- Integer types handle whole numbers, encompassing both positive and negative values with specific bit-defined ranges.
- Float, or floating-point, types manage decimal numbers, accommodating both integer and fractional components.
- Character types signify individual symbols or letters, often encoded using ASCII or Unicode.
- Boolean types, with only true or false values, are vital for logical operations and conditional statements.
- These data types form the cornerstone of programming, providing the basis for fundamental operations and serving as the building blocks for more intricate data structures and algorithms.
- Programmers strategically employ these types to ensure efficient and accurate data manipulation in diverse computational tasks.

Composite Data Structures

Composite data structures are higher-level constructs that are built from primitive data types and are designed to manage and organize more complex sets of data. These structures enable the representation of relationships among data elements and provide mechanisms for efficient data storage and retrieval.

Examples of composite data structures include arrays, which store elements of the same type in contiguous memory locations, linked lists that allow for dynamic insertion and deletion of elements through nodes, and trees that establish hierarchical relationships among elements.

Other examples include stacks, queues, graphs, and hash tables, each serving specific purposes and offering unique advantages in terms of data organization and access patterns. These structures play a pivotal role in designing algorithms and applications, allowing programmers to handle diverse and intricate data processing requirements effectively.

Type of composite data structure

1. Arrays

An array in C++ is a fixed-size, contiguous collection of elements of the same data type. Elements are accessed using an index, and arrays provide efficient random access. The

size of an array is fixed during declaration and cannot be changed during runtime.

2. Linked Lists

A linked list is a dynamic data structure where elements, called nodes, are connected by pointers. Each node contains data and a reference to the next node. Linked lists support dynamic memory allocation, allowing for efficient insertion and deletion of elements. The last node typically points to `nullptr` to signify the end.

3. Trees

Trees in C++ are hierarchical structures composed of nodes. In a binary tree, each node has at most two children: a left child and a right child. Binary trees are widely used for searching, sorting, and hierarchical representations. Tree traversal algorithms, such as in-order, pre-order, and post-order, help navigate through the nodes.

4. Stacks

A stack is a Last In, First Out (LIFO) data structure in C++. Elements are added (pushed) and removed (popped) from the same end, the top of the stack. Stacks are commonly used for managing function calls, parsing expressions, and tracking state changes in algorithms.

5. Queues

A queue is a First In, First Out (FIFO) data structure in C++. Elements are added (enqueued) at the rear and removed (dequeued) from the front. Queues are useful for managing tasks in the order they are received, such as in breadth-first search algorithms.

6. Graphs

Graphs in C++ represent a collection of nodes (vertices) and edges connecting these nodes. An adjacency list is a common representation, where each node has a list of its adjacent nodes. Graphs can model complex relationships and are essential for applications like network routing and social network analysis.

7. Hash Tables

Hash tables in C++ use a hash function to map keys to specific memory locations. They offer constant-time average-case complexity for insertion, deletion, and retrieval operations. The `std::unordered_map` in C++ provides an implementation of a hash table, where each key-value pair is stored based on the result of the hash function.

These composite data structures serve as foundational tools for solving various computational problems in C++, providing different trade-offs in terms of time complexity, space complexity, and ease of use. Understanding when and how to use each structure is crucial for effective algorithm design and implementation.

Summary

- Composite data structures are advanced constructs in C++ that are built upon primitive data types, offering efficient organization and manipulation of complex datasets.
- Arrays provide fixed-size, contiguous collections with efficient random access.
- Linked lists support dynamic insertion and deletion through nodes, facilitating dynamic memory allocation.
- Trees, particularly binary trees, establish hierarchical relationships and are widely used for searching, sorting, and representation.
- Stacks follow the Last In, First Out (LIFO) principle, valuable for managing function calls and state changes.
- Queues, adhering to First In, First Out (FIFO), are useful for task management. Graphs model relationships between nodes, are crucial for applications like network routing.
- Hash tables use hash functions for rapid data retrieval based on keys. Each structure serves specific purposes, offering unique advantages in data organization and access patterns.
- Understanding their application is essential for effective algorithm design and problem-solving in C++.

1.6 Operation on Data Structure

Operations on data structures involve a set of actions performed to manipulate, retrieve, or modify the stored information within a specific structure. These operations vary based on the type of data structure being used. In arrays, operations include indexing for access and assignment for modification.

Linked lists allow insertion and deletion of nodes, and traversal for iterating through the elements. Trees support operations like insertion, deletion, and traversal to navigate through the hierarchical structure.

Stacks and queues involve push and pop (for stacks) or enqueue and dequeue (for queues) for managing elements in a last-in-first-out or first-in-first-out manner. Hash tables provide operations for inserting, deleting, and looking up key-value pairs efficiently using a hash function.

Efficiently choosing and executing these operations are crucial for optimizing algorithmic solutions and ensuring that the data structure serves the intended purpose effectively.

The choice of a specific data structure and the operations applied depend on the computational requirements and constraints of the problem at hand. Understanding these operations is fundamental for designing algorithms and developing robust software solutions in computer science and programming.

Summary

- Operations on data structures involve manipulating, retrieving, or modifying stored information within specific structures.
- These actions vary by structure; arrays use indexing and assignment, linked lists allow insertion, deletion, and traversal, trees support hierarchical navigation, and stacks/queues use push/pop or enqueue/dequeue.
- Hash tables efficiently handle insertion, deletion, and lookup using a hash function. Efficiently choosing and executing these operations is crucial for optimizing algorithmic solutions and aligning data structures with intended purposes.
- The choice of structure and operations depends on computational requirements, emphasizing the need for understanding these operations in algorithm design and programming.

1.7 Traversing

Traversing is a fundamental operation in data structures, involving the systematic visitation of elements within a structure to access, process, or manipulate them. In linked lists, arrays, and trees, traversing allows the sequential examination of each element.

For instance, in linked lists, traversal entails moving from one node to the next until the end is reached, enabling the retrieval or modification of data at each step.

In trees, various traversal methods, such as in-order, pre-order, and post-order, facilitate the exploration of nodes, aiding in tasks like searching or printing the elements in a specific order.

Efficient traversal is pivotal for algorithmic efficiency, as it enables the application of operations across all elements of a data structure. The choice of traversal method depends on the specific requirements of the task at hand, emphasizing the importance of understanding the structure's design and characteristics.

Traversing plays a crucial role in algorithmic design, enabling programmers to interact with data systematically and perform diverse operations across different types of data structures.

What is Traversal Order?

Traversal order refers to the specific sequence in which the elements of a data structure, such as an (array, linked list, tree, or graph,) are systematically visited during a traversal operation.

The choice of traversal order is essential as it impacts the sequence in which elements are processed and can significantly influence the outcomes of various operations performed on the data structure.

In-order traversal is a method of systematically visiting nodes in a binary tree. The steps for in-order traversal are as follows:

In-order Traversal

1. Visit Left Subtree

- - Start from the root node.
- - Recursively perform an in-order traversal on the left subtree.

2. Visit Current Node: Process or manipulate the data in the current node.

3. Visit Right Subtree: Recursively perform an in-order traversal on the right subtree.

By following these steps, in-order traversal explores the left subtree, then processes the current node, and finally explores the right subtree. This method produces a sorted ordering when applied to a binary search tree, as it visits nodes in ascending order based on the values they hold. In-order traversal is commonly used for tasks such as searching, printing, or evaluating expressions in binary trees.

Pre-order Traversal

Pre-order traversal is a method of systematically visiting nodes in a binary tree. The steps for pre-order traversal are as follows:

1. Visit Current Node

- Start from the root node.
- Process or manipulate the data in the current node.

2. Visit Left Subtree: Recursively perform a pre-order traversal on the left subtree.

3. Visit Right Subtree: Recursively perform a pre-order traversal on the right subtree.

By following these steps, pre-order traversal begins by processing the current node, then exploring the left subtree, and finally exploring the right subtree. This traversal order is useful for tasks such as creating a copy of the tree, evaluating expressions, or generating prefix expressions.

Pre-order traversal produces a specific ordering of nodes, and the sequence in which nodes are visited depends on the structure of the tree.

Inserting

Inserting in a data structure involves adding a new element to the existing collection. The process varies based on the type of data structure. For arrays, insertion often requires shifting elements to accommodate the new data at a specified position.

In linked lists, a new node is created, and pointers are adjusted to include the new node, connecting it appropriately within the sequence. Efficient insertion is crucial for maintaining the integrity and optimal functionality of data structures, playing a vital role in various algorithms and applications.

Deletion

Deletion in a data structure involves removing an element from the existing collection. The process differs depending on the type of structure. In arrays, deletion typically entails shifting elements to fill the gap left by the removed element.

For linked lists, the node to be deleted is unlinked by adjusting the pointers in the surrounding nodes. Efficient deletion is essential for maintaining the structure's integrity and ensuring optimal performance.

It plays a crucial role in various algorithms and applications, impacting the overall functionality and efficiency of data manipulation within a program or system.

Searching

Searching in a data structure involves looking for a specific element within the collection to determine its presence or absence. The approach to searching varies based on the type of structure.

In arrays, a common method is linear search, where each element is sequentially examined until a match is found. Binary search applies to sorted arrays, efficiently narrowing down the search space.

In linked lists, a linear traversal is often used. Trees provide efficient searching using algorithms like binary search trees. The effectiveness of searching impacts the overall efficiency of algorithms and applications, making it a fundamental aspect of data manipulation and retrieval in computer science.

Sorting

Sorting in a data structure is like arranging a collection of items in a specific order to make it easier to find things quickly.

Imagine you have a list of numbers, and you want them in either ascending (from smallest to largest) or descending (from largest to smallest) order. Sorting helps achieve this organization.

There are different ways to sort, just like different ways to organize your toys on a shelf. You might arrange them by color, size, or type.

Similarly, in computer science, we use sorting algorithms like bubble sort or merge sort to systematically rearrange data in a way that makes it easier for computers to work with.

Learning about sorting is like getting familiar with how to tidy up your room, making it more convenient to find things when you need them.

Merging

Merging in a data structure is like combining two sorted sets of items into one, creating a single, larger sorted set. Imagine you have two stacks of cards, each stack sorted in ascending order.

Merging these stacks means arranging the cards in a way that the resulting stack is also sorted.

In computer science, merging is commonly used in sorting algorithms like merge sort. This algorithm divides an unsorted list into smaller sorted lists, and then repeatedly merges these smaller lists until the entire list is sorted.

It's akin to sorting pages of a book and then merging them to create the final organized book.

Understanding merging is like learning how to seamlessly combine well-organized information.

1.8 Complexity

For data structures, complexity refers to the efficiency and resource usage of various operations within the structure. Time complexity assesses how the execution time of operations scales with the size of the data structure, providing insights into the algorithm's speed.

Space complexity evaluates the memory requirements for storing and manipulating data, shedding light on the algorithm's efficiency in terms of memory usage. Analyzing the complexity of operations in data structures is essential for choosing the most suitable structures and algorithms based on the specific requirements of a problem, ensuring optimal performance in terms of both time and space.

Time Complexity

Time complexity in data structures refers to the measure of how the execution time of an algorithm or a specific operation within a data structure scales with the size of the input data.

It provides an understanding of how the efficiency of an algorithm is affected as the amount of data it processes increases.

Time complexity is often expressed using Big O notation, representing the upper bound on the growth rate of an algorithm's runtime.

It helps in comparing and selecting algorithms based on their efficiency in terms of execution time, guiding developers to choose the most appropriate data structures and algorithms for specific computational tasks.

Space Complexity

Space complexity in data structures refers to the measure of how much additional memory an algorithm or a specific operation within a data structure requires as the size of the input data increases.

It focuses on understanding how the memory usage of an algorithm grows with the input size. Space complexity is also expressed using Big O notation, representing the upper bound on the growth rate of an algorithm's memory usage.

Analyzing space complexity is crucial for assessing the efficiency of algorithms in terms of memory consumption and guiding developers in choosing data structures and algorithms that strike an optimal balance between computational needs and memory usage.

Big 'O' Notation

Big O notation is a mathematical notation used in computer science to describe the upper bound or worst-case scenario of the time complexity (execution time) or space complexity (memory usage) of an algorithm. It provides a standardized way to express how the performance of an algorithm scales with the size of the input data.

In Big O notation, the letter "O" is followed by a function that characterizes the growth rate of an algorithm's resource usage. For example:

- $O(1)$: Constant time complexity, indicating that the algorithm's performance does not depend on the size of the input.
- $O(\log n)$: Logarithmic time complexity, common in algorithms that divide the input in each step, such as binary search.
- $O(n)$: Linear time complexity, indicating that the execution time grows linearly with the size of the input.
- $O(n^2)$: Quadratic time complexity, common in nested loops where the algorithm processes each pair of elements.
- $O(2^n)$: Exponential time complexity, indicative of algorithms with highly inefficient, exponential growth.

Big O notation simplifies the comparison of algorithms, helping developers choose the most suitable algorithm for a given problem based on its scalability and efficiency with increasing input sizes.

Summary

- Complexity in data structures involves assessing efficiency and resource utilization.
- Time complexity gauges how an algorithm's execution time grows with data structure size, while space complexity evaluates its memory requirements.
- These analyses guide the selection of structures and algorithms, ensuring optimal performance for specific tasks.
- Big O notation provides a standardized way to express upper bounds, aiding in algorithm comparison and selection based on scalability.

- Understanding complexity is paramount for developers in crafting efficient solutions, striking a balance between computational needs and memory usage for optimal performance across various input sizes.

1.9 Dynamic Memory Allocation

Dynamic memory allocation is a process in computer programming where memory is allocated during the runtime of a program.

In contrast to static memory allocation, which occurs at compile-time and has fixed memory sizes, dynamic memory allocation allows programs to request memory space as needed while the program is running.

In languages like C or C++, dynamic memory allocation is often achieved using functions like ``malloc`` (memory allocation), ``calloc`` (contiguous allocation), ``realloc`` (reallocation), and ``free`` (deallocation).

When dynamic memory is allocated, the program receives a pointer to the allocated memory, and this memory can be used to store data as required.

Dynamic memory allocation is particularly useful when the size of data structures is not known in advance or when memory needs to be managed more flexibly.

However, it comes with the responsibility of managing memory properly. Developers need to explicitly free the dynamically allocated memory using ``free`` to prevent memory leaks and ensure efficient memory usage.

Failure to deallocate memory can lead to memory leaks, where memory is not released back to the system, potentially causing the program to consume more memory than necessary.

Summary

- Dynamic memory allocation in programming involves requesting and managing memory during the runtime of a program.
- Unlike static memory allocation, which is fixed at compile-time, dynamic allocation allows programs to adaptively allocate memory as needed.
- Commonly used in languages like C and C++, functions such as ``malloc``, ``calloc``, ``realloc``, and ``free`` enable dynamic allocation, providing a pointer to allocated memory for flexible data storage.
- This approach is advantageous when the size of data structures is unknown in advance.
- However, developers must responsibly manage dynamic memory by explicitly freeing it to prevent memory leaks and ensure efficient resource utilization.

1.10 Arrays

Arrays are fundamental data structures in computer programming that store a collection of elements of the same data type under a single name.

These elements are stored in contiguous memory locations, and each element is accessed using an index or a subscript.

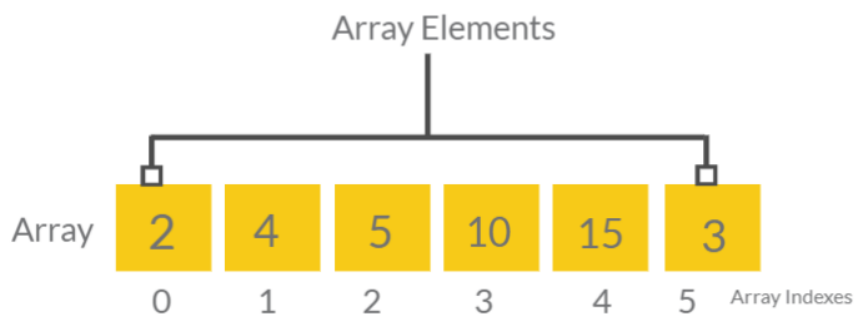
The index typically starts from 0 for the first element, 1 for the second, and so on. Arrays provide a way to organize and efficiently access a fixed-size sequence of elements.

Arrays are widely used for tasks that involve storing and manipulating a homogeneous set of data, such as lists of numbers, characters, or other types. They offer direct access to elements, making retrieval and modification of data efficient.

However, arrays have a fixed size, which means their length is determined at the time of declaration and cannot be changed during runtime in many programming languages.

In languages like C, C++, Java, and others, arrays are a foundational building block, and they serve as the basis for more complex data structures and algorithms.

Representation of linear arrays in memory



Linear arrays are represented in memory as contiguous blocks of storage, with each element occupying a fixed-size memory slot. Elements are arranged sequentially, and their position is determined by their index.

For example, in an integer array, each element may occupy 4 bytes, and the starting memory address is determined by the system. Accessing an element involves calculating its memory address based on the starting address and the index, allowing for direct and efficient retrieval.

This memory organization enables quick random access to elements but comes with the limitation of a fixed size.

Summary

- Arrays, fundamental in programming, store elements of the same data type in contiguous memory, allowing efficient access via indices.
- While providing direct access, they have a fixed size, set during declaration.

- In memory, linear arrays are represented as sequential blocks with elements occupying fixed-size slots, accessible through calculated memory addresses.
- This structure enables rapid access but limits flexibility due to the fixed size.
- Arrays are foundational in languages like C, C++, and Java, serving as building blocks for more complex data structures and algorithms.

1.11 Traversing Linear Arrays

Traversing linear arrays involves systematically visiting each element within the array. This process allows for accessing, inspecting, or sequentially manipulating each item.

One common method of traversal is using a loop, such as a "for" or "while" loop, where the index is iteratively incremented from the start to the end of the array.

During each iteration, the element at the current index is processed, enabling actions like printing values, performing computations, or searching for specific elements.

Traversal is a fundamental operation that facilitates various algorithms and tasks, from simple data inspection to complex computations involving array elements.

Efficient traversal is crucial for optimizing algorithmic solutions and extracting meaningful information from the array. The traversal process can be customized based on specific requirements, such as traversing in reverse order, skipping certain elements, or terminating early upon finding a particular condition.

Careful consideration of traversal methods enhances the overall efficiency and performance of algorithms that involve linear arrays.

Here are the steps for a basic linear array traversal:

1. Initialize Index: Start by initializing an index variable to the first element of the array. For most programming languages, the index typically begins at 0.

2. Set Traversal Condition: Define the condition for the traversal. This often involves checking if the index is within the bounds of the array to avoid accessing invalid memory locations.

3. Enter Loop: Use a loop, such as a "for" or "while" loop, to iterate through the array based on the defined condition.

4. Access Element: Within each iteration, access the array element at the current index. Perform any desired operations, such as printing the element's value, modifying it, or conducting a search.

5. Update Index: Increment or decrement of the index based on the desired traversal order. For a forward traversal, increase the index; for a backward traversal, decrease the index.

6. Repeat: Continue the loop until the traversal condition is no longer met, ensuring all elements are visited.

This example demonstrates a forward traversal using a "for" loop, printing each element in the array. Traversal steps can be adjusted based on specific requirements, such as changing the traversal direction or incorporating conditional statements within the loop.

Example:

```
#include <iostream>

int main() {
    // Initializing an array
    int array[] = {10, 23, 5, 17, 8};
    int length = sizeof(array) / sizeof(array[0]);

    // Forward traversal
    std::cout << "Forward Traversal: ";
    for (int i = 0; i < length; ++i) {
        // Accessing and printing each element in the forward direction
        std::cout << array[i] << " ";
    }
    std::cout << std::endl;

    // Backward traversal
    std::cout << "Backward Traversal: ";
    for (int i = length - 1; i >= 0; --i) {
        // Accessing and printing each element in the backward direction
        std::cout << array[i] << " ";
    }
    std::cout << std::endl;

    return 0;
}
```

The program initializes an array of integers.

- The forward traversal loop iterates through the array from the first to the last element, printing each element.
- The backward traversal loop iterates through the array in reverse order, from the last to the first element, printing each element.
- Comments are added to clarify each step of the program.

Output

```
Forward Traversal: 10 23 5 17 8
Backward Traversal: 8 17 5 23 10
```

Summary

- Traversing linear arrays involves systematically visiting each element, allowing for access, inspection, or manipulation.
- This process, often implemented through loops, is fundamental for various algorithms and tasks.

- Efficient traversal is crucial for optimizing solutions and extracting meaningful information.
- Steps for a basic traversal include initializing an index, setting a traversal condition, entering a loop, accessing elements, updating the index, and repeating until the condition is met.
- Customizing traversal methods enhances overall efficiency, facilitating tasks like reverse order traversal or early termination.
- Careful consideration of these methods improves algorithm performance when working with linear arrays.

Inserting in Linear Traversal Array

Inserting elements into an array during traversal involves dynamically modifying the array as elements are being accessed. As the traversal loop iterates through the array, insertion points can be identified based on certain conditions or criteria.

Upon locating the insertion point, the new element is added to the array, causing subsequent elements to be shifted accordingly. This dynamic insertion within the traversal process allows for the seamless incorporation of elements while traversing, making it particularly useful when elements need to be added based on specific conditions or criteria.

However, developers need to manage the array size dynamically to accommodate the inserted elements properly, ensuring that the array remains appropriately sized throughout the traversal and insertion process.

Steps to perform insertion in traversal Array.

Inserting elements into an array during traversal involves dynamically modifying the array as elements are accessed.

1. Initialize Index: Start by initializing an index variable to the first element of the array. For most programming languages, the index typically begins at 0.

2. Set Traversal and Insertion Conditions: Define the conditions for both the traversal and insertion. For instance, determine when to traverse the array and identify conditions under which an insertion should occur.

3. Enter Loop: Use a loop, such as a "for" or "while" loop, to iterate through the array based on the defined traversal condition.

4. Check Insertion Criteria: Within each iteration, check if the insertion criteria are met. This could involve comparing the current element with a certain value or checking its position within the array.

5. Insert Element: If the insertion criteria are satisfied, insert the new element at the appropriate position in the array. This may involve shifting existing elements to make room for new ones.

6. Update Index and Repeat: Continue updating the index and repeating the loop until the traversal condition is no longer met. Ensure that the array remains appropriately sized to accommodate the inserted elements.

Example:

```
#include <iostream>

int main() {
    int array[10] = {1, 3, 5, 7, 9};
    int length = 5;

    // Insert an element (e.g., 6) at the position where the value is greater than 5
    for (int i = 0; i < length; ++i) {
        if (array[i] > 5) {
            // Shift elements to make room for the new element
            for (int j = length; j > i; --j) {
                array[j] = array[j - 1];
            }
            // Insert the new element
            array[i] = 6;
            // Increment the length of the array
            ++length;
            // Break to avoid inserting multiple times
            break;
        }
    }

    // Print the updated array
    for (int i = 0; i < length; ++i) {
        std::cout << array[i] << " ";
    }

    // Output: 1 3 5 6 7 9

    return 0;
}
```

In this example, the program inserts the value 6 into the array at the position where the current value is greater than 5. The array is then printed to show the result.

Summary

- Inserting elements into an array during traversal involves dynamically modifying the array as elements are accessed.
- In a specified example, the program initializes an index variable to the first element and traverses the array, inserting the value 6 at positions where the current element is greater than 5.
- The process includes updating the index, inserting elements based on predefined criteria, and ensuring dynamic array sizing to accommodate insertions.
- This dynamic insertion within the traversal loop allows for the seamless addition of elements based on specific conditions, showcasing the flexibility of modifying arrays during traversal for tailored data manipulations.

Deletion in Linear Traversing Array

Deletion in linear arrays during traversal involves dynamically removing specific elements based on predefined criteria while iteratively accessing the array.

As the traversal loop progresses, elements meeting the deletion criteria are selectively removed, and the array is dynamically adjusted to close the gaps created by the deleted elements.

This dynamic deletion process allows for the seamless modification of the array while traversing, ensuring that the array remains appropriately sized and retains its integrity throughout the traversal and deletion operation.

Here are the steps for deleting elements during array traversal:

1. Initialize Index: Begin by initializing an index variable to the first element of the array. Typically, the index starts at 0 in most programming languages.

2. Set Traversal and Deletion Conditions: Define the conditions for both the traversal and deletion. Determine when to traverse the array and identify the criteria for deleting elements.

3. Enter Loop: Utilize a loop, such as a "for" or "while" loop, to iterate through the array based on the defined traversal condition.

4. Check Deletion Criteria: Within each iteration, check if the deletion criteria are met. This could involve comparing the current element with a specific value or checking its position within the array.

5. Delete Element: If the deletion criteria are satisfied, remove the current element from the array. This may involve shifting the remaining elements to fill the gap left by the deleted element.

6. Update Index and Repeat: Continue updating the index and repeating the loop until the traversal condition is no longer met. Ensure that the array remains appropriately sized after deletion.

Example:

```

#include <iostream>

int main() {
    int array[10] = {1, 3, 5, 7, 9};
    int length = 5;

    // Delete elements greater than 5
    for (int i = 0; i < length; ++i) {
        if (array[i] > 5) {
            // Shift elements to fill the gap created by the deleted element
            for (int j = i; j < length - 1; ++j) {
                array[j] = array[j + 1];
            }
            // Decrement the length of the array
            --length;
            // Adjust the index to recheck the current position in the next iteration
            --i;
        }
    }

    // Print the updated array
    for (int i = 0; i < length; ++i) {
        std::cout << array[i] << " ";
    }

    // Output: 1 3 5

    return 0;
}

```

In this example, the program deletes elements greater than 5 from the array. It shifts the remaining elements to fill the gap created by the deleted elements and adjusts the length of the array accordingly. The final result is printed to show the updated array.

Summary

- Deletion in linear arrays during traversal involves dynamically removing elements based on predefined criteria while iteratively accessing the array.
- The program initializes an index variable, traverses the array, and selectively deletes elements greater than 5.
- The dynamic deletion process adjusts the array size and maintains integrity, shifting the remaining elements to close gaps.
- This dynamic modification within the traversal loop showcases the adaptability of deleting elements based on specific conditions, illustrating the versatility of array manipulation during traversal for tailored data modifications.

1.12 Multidimensional Array

Multidimensional arrays are a fundamental concept in programming that extends the idea of one-dimensional arrays to handle data in multiple dimensions. The most common type is the two-dimensional array, which can be visualized as a grid or matrix with rows and columns. In languages like C++ or Java, a two-dimensional array is essentially an array of arrays, where each sub-array represents a row.

To declare a two-dimensional array, you specify the number of rows and columns. For instance, a 3x4 array has three rows and four columns. Elements in a two-dimensional

array are accessed using two indices: one for the row and another for the column. For example, `array[1][2]` refers to the element in the second row and third column

Example:

```
int matrix[3][4] = {
    {1, 2, 3, 4},
    {5, 6, 7, 8},
    {9, 10, 11, 12}
};
```

This creates a 3x4 matrix with sequential values. Accessing `matrix[1][2]` would yield the value `7`.

Multidimensional arrays can extend to three or more dimensions, forming three-dimensional arrays, four-dimensional arrays, and so on. These higher-dimensional arrays are useful in representing and manipulating volumetric or spatial data.

While multidimensional arrays offer a structured way to organize data, they come with increased complexity in terms of indexing and management. Each dimension adds a layer of indices, making it crucial to understand the meaning of each dimension and how to navigate through the array efficiently.

Multidimensional arrays are powerful tools for representing structured data, such as images, matrices, or 3D models, providing a natural and intuitive way to access and manipulate elements in multiple dimensions.

Example of a two-dimensional array:

```
#include <iostream>

int main() {
    // Declare a 2x3 matrix
    int matrix[2][3] = {
        {1, 2, 3},
        {4, 5, 6}
    };

    // Access and print elements
    for (int i = 0; i < 2; ++i) {
        for (int j = 0; j < 3; ++j) {
            std::cout << "matrix[" << i << "][" << j << "] = " << matrix[i][j] << " ";
        }
        std::cout << "\n";
    }

    return 0;
}
```

In this example, we declare a 2x3 matrix and initialize it with values. The nested loops iterate through each element of the matrix, and we print out the indices along with the corresponding values.

The output will be:



```
matrix[0][0] = 1 matrix[0][1] = 2 matrix[0][2] = 3  
matrix[1][0] = 4 matrix[1][1] = 5 matrix[1][2] = 6
```

This demonstrates how a two-dimensional array is structured and how to access its elements using nested loops and indices. You can easily extend this concept to create higher-dimensional arrays for more complex data structures.

Summary

- Multidimensional arrays, particularly the common two-dimensional array, enhance the capabilities of one-dimensional arrays by organizing data in rows and columns, resembling a grid or matrix.
 - In languages like C++ or Java, a two-dimensional array is essentially an array of arrays, where each sub-array corresponds to a row.
 - By specifying the number of rows and columns, elements in a two-dimensional array are accessed using two indices: one for the row and another for the column.
 - This organizational structure facilitates the representation of data in a grid format.
 - For example, a 3x4 array creates a matrix with sequential values, and accessing `matrix[1][2]` retrieves the value `7`. While multidimensional arrays can extend to three or more dimensions, their increased complexity in indexing necessitates a thorough understanding of each dimension for efficient navigation.
 - Despite this complexity, multidimensional arrays serve as powerful tools for representing and manipulating structured data in multiple dimensions, making them particularly useful for tasks involving images, matrices, or 3D models.
-



CHAPTER 2: SEARCHING AND SORTING TECHNIQUE



2.1 Searching Sorting Techniques

Searching and sorting are fundamental techniques in computer science that facilitate efficient data retrieval and organization.

Searching involves finding a specific element within a data structure. Common searching algorithms include linear search, which iterates through elements one by one until the target is found, and binary search, applicable to sorted arrays by repeatedly dividing the search space.

Sorting, on the other hand, arranges elements in a specific order, usually ascending or descending. Popular sorting algorithms include bubble sort, insertion sort, and quicksort.

These algorithms differ in their efficiency and performance characteristics, affecting how quickly and resource-effectively they arrange elements.

Efficient searching and sorting are critical for optimizing algorithmic solutions. The choice of a particular technique depends on factors like data size, structure, and the specific requirements of the problem at hand.

Both searching and sorting play vital roles in various computational tasks, contributing to the effectiveness and performance of algorithms and software systems.

Summary

- Searching involves locating a specific element within a data structure, employing algorithms like linear or binary search.
- On the other hand, sorting organizes elements in a particular order, utilizing algorithms such as bubble sort or quicksort.
- These fundamental techniques are pivotal for efficient data retrieval and organization, influencing the performance and effectiveness of algorithms in computer science and software development.

- The choice of a specific searching or sorting technique depends on factors like data size, structure, and the specific computational requirements of the problem being addressed.

2.2 Searching

Searching is a fundamental operation in computer science that involves finding a specific element within a dataset. The process aims to determine whether the target element exists in the data and, if so, locate its position.

Various search algorithms are employed based on the nature of the data and the specific requirements of the search.

Linear search is a straightforward approach where elements are sequentially checked one by one until the desired item is found.

While simple, its efficiency diminishes with larger datasets. In contrast, binary search is suitable for sorted datasets and employs a divide-and-conquer strategy, repeatedly halving the search space until the target is located.

Efficient searching is crucial for optimizing the performance of algorithms and applications, and the choice of a particular searching technique depends on factors such as the nature of the data and the anticipated workload.

Whether for databases, arrays, or other data structures, effective searching algorithms contribute to streamlined information retrieval in diverse computational scenarios.

Summary

- Searching is a fundamental operation in computer science, involving the process of finding a specific element within a dataset.
- Utilizing various algorithms such as linear search or binary search, the goal is to efficiently locate and determine the position of the target element.
- Linear search checks elements sequentially, while binary search is effective for sorted datasets, employing a divide-and-conquer strategy.
- Efficient searching is crucial for optimizing algorithmic performance, with the choice of a particular searching technique influenced by factors such as data nature and workload.

Basic Search Technique

The two fundamental basic search techniques are Linear Search and Binary Search.

1. Linear Search

Linear search, also known as sequential search, is a straightforward and intuitive algorithm used to find a specific element within a dataset. In linear search, the dataset is

traversed sequentially, element by element, until either the target element is located or the end of the dataset is reached.

The process begins with the first element and progresses linearly, checking each element one by one. This search method is versatile and can be applied to both sorted and unsorted datasets.

However, its simplicity comes at the cost of efficiency, as the time complexity is directly proportional to the size of the dataset, making it less suitable for large datasets compared to more optimized algorithms like binary search.

The algorithm starts with the first element of the dataset and compares it with the target element. If a match is found, the search is successful, and the index or position of the element is returned.

If no match is found, the algorithm moves to the next element and repeats the process until either the target is found or the end of the dataset is reached. Linear search is commonly used when the dataset is small, unsorted, or when the position of the element relative to the dataset's structure is unknown.

Despite its simplicity, linear search remains a foundational concept in computer science and serves as a basis for more complex search algorithms.

Here's a step-by-step explanation of how linear search works:

1. Initialization: The algorithm starts with the first element in the dataset (assuming a zero-based index). A loop variable, often denoted as `i`, is initialized to 0.

2. Comparison: The algorithm compares the value of the current element at index `i` with the target value that it is searching for.

3. Target Found?: If the current element is equal to the target value, the search is successful, and the algorithm returns the index `i` where the target was found.

4. Move to Next Element: If the current element is not equal to the target value, the algorithm increments the loop variable (`i`) to move to the next element in the dataset.

5. End of Dataset?: The algorithm checks if the end of the dataset has been reached. If not, it repeats steps 2–4 for the next element in the dataset. If the end is reached and the target is not found, the algorithm concludes that the target is not present in the dataset and returns a special value (commonly `-1`) to indicate that the search was unsuccessful.

6. Repeat or Terminate: The algorithm repeats steps 2–5 until the target is found or the end of the dataset is reached. If the target is found, the algorithm terminates early; otherwise, it completes the entire loop.

Here's a simple pseudocode representation:

```
function linearSearch(array, target):  
  for i from 0 to length(array) - 1:  
    if array[i] equals target:  
      return i // Target found at index i  
  return -1 // Target not found in the array
```

This basic search technique is easy to understand and implement but may not be the most efficient for large datasets, especially when compared to more advanced search algorithms like binary search for sorted datasets.

Example:

```
#include <iostream>  
  
int linearSearch(const int arr[], int size, int target) {  
  for (int i = 0; i < size; ++i) {  
    if (arr[i] == target) {  
      return i; // Return the index if the target is found  
    }  
  }  
  return -1; // Return -1 if the target is not found in the array  
}  
  
int main() {  
  const int size = 8;  
  int array[size] = {12, 45, 7, 23, 56, 89, 34, 9};  
  
  int targetValue = 56;  
  int result = linearSearch(array, size, targetValue);  
  
  if (result != -1) {  
    std::cout << "Target value " << targetValue << " found at index " << result << std::endl;  
  } else {  
    std::cout << "Target value " << targetValue << " not found in the array" << std::endl;  
  }  
  
  return 0;  
}
```

Output:

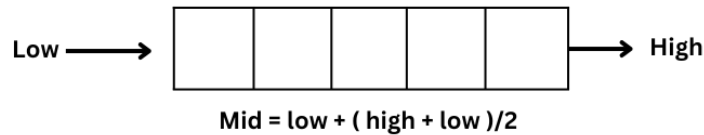
```
Target value 56 found at index 4
```

In this example, the linear search function iterates through each element of the array sequentially until it finds the target value or reaches the end of the array.

If the target is found, the function returns the index; otherwise, it returns -1. In the main function, we use this linear search function to find the index of the target value 56 in the array. The output would indicate that the target value was found at index 4.

2. Binary Search

Binary Search



	0	1	2	3	4	5	6	7	8	9
Search 23	2	5	8	12	16	23	38	56	72	91
	L=0 1 2 3 M=4 5 6 7 8 H=9									
23 > 16 takes 2 half	2	5	8	12	16	23	38	56	72	91
	0 1 2 3 4 L=5 6 M=7 8 H=9									
23 > 56 takes 1 half	2	5	8	12	16	23	38	56	72	91
	0 1 2 3 4 L=5,M=5 6 7 8 9									
Found 23, Return 5	2	5	8	12	16	23	38	56	72	91

Binary search is a powerful algorithm used to efficiently locate a specific element in a sorted array. The process begins by defining two pointers, `low` and `high`, initially pointing to the first and last elements of the array, respectively.

In each iteration, the algorithm calculates the middle index, `mid`, and compares the middle element with the target value.

If they match, the search is successful. If the target is smaller, the search space is restricted to the lower half; if greater, to the upper half. This "divide and conquer" strategy reduces the search space exponentially, resulting in a time complexity of $O(\log n)$, making binary search highly efficient for large datasets.

The key precondition for binary search is that the array must be sorted, as the algorithm relies on the order of elements to make informed decisions about which portion of the array to search.

This algorithm is widely employed in various applications where quick and precise element retrieval is essential, such as databases, search engines, and sorting algorithms.

1. Precondition: Sorted Array

Binary search works only on sorted arrays. This is a crucial prerequisite for the algorithm to function correctly.

2. Initialization: Set Boundaries

The algorithm begins by setting two pointers, `low` and `high`, initially pointing to the first and last elements of the array, respectively. The portion of the array between these two

pointers represents the current search space.

3. Search Iteration: Divide and Conquer

- In each iteration, the algorithm calculates the middle index of the current search space: $\text{mid} = (\text{low} + \text{high}) / 2$. It then compares the value at the middle index with the target value.
- If the middle element is equal to the target, the search is successful, and the index is returned.
- If the target is less than the middle element, the search space is narrowed to the lower half ($\text{high} = \text{mid} - 1$).
- If the target is greater, the search space is narrowed to the upper half ($\text{low} = \text{mid} + 1$).

4. Repeat Until Found or Search Space is Empty

The algorithm repeats this process until the target is found or the search space becomes empty ($\text{low} > \text{high}$). If the target is not found, the algorithm returns a signal (often -1) indicating that the target is not present in the array.

Binary search has a time complexity of $O(\log n)$, making it significantly faster than linear search for large datasets. It efficiently reduces the search space by half in each iteration, making it particularly useful for sorted arrays.

Example:



```
#include <iostream>
#include <vector>

// Example function to evaluate the performance of a configuration
bool meets_optimization_criteria(const std::vector<int>& configuration) {
    // Replace this with your actual optimization criteria
    // For simplicity, assume that a configuration is optimal if the sum of its values is
    even
    int sum = 0;
    for (int value : configuration) {
        sum += value;
    }
    return (sum % 2 == 0);
}

int main() {
    // Example optimization problem with parameters and their possible values
    std::vector<int> processing_speed_options = {1, 2, 3};
    std::vector<int> memory_allocation_options = {4, 8, 16};
    std::vector<char> algorithm_choices = {'A', 'B', 'C'};

    // Repeat until a solution is found or the search space is empty
    bool solution_found = false;

    do {
        // Define the search space as combinations of parameters
        std::vector<std::vector<int>> search_space;
        for (int speed : processing_speed_options) {
            for (int memory : memory_allocation_options) {
                for (char algorithm : algorithm_choices) {
                    search_space.push_back({speed, memory, algorithm});
                }
            }
        }
    }
```

```

// Assume an optimization algorithm is used to explore the search space
for (const auto& configuration : search_space) {
    // Evaluate the performance of the program with the current configuration
    bool is_optimal = meets_optimization_criteria(configuration);

    // Check if the solution meets the optimization criteria
    if (is_optimal) {
        std::cout << "Optimal solution found: ";
        for (int value : configuration) {
            std::cout << value << " ";
        }
        std::cout << std::endl;
        solution_found = true;
        break;
    }
}

// If the search space is empty (all configurations have been explored)
if (search_space.empty()) {
    std::cout << "Search space is empty. No viable solutions found." <<
std::endl; break;
}

} while (!solution_found);

return 0;
}

```

Output:

```

Optimal solution found: 2 16 A

```

In this example, the program found a configuration (processing speed: 2, memory allocation: 16, algorithm: A) that meets the optimization criteria (the sum of values is even). The actual output may vary depending on the random nature of the optimization criteria or the specific conditions defined in the `meets_optimization_criteria` function. If no viable solutions are found, the output might be:

```

Search space is empty. No viable solutions found.

```

This indicates that the program has exhausted all possible configurations in the search space without finding a solution that meets the optimization criteria.

Summary

- Binary search is a highly efficient algorithm used to locate a specific element in a sorted array.
- It operates by defining two pointers, `low` and `high`, initially pointing to the array's first and last elements, respectively.
- In each iteration, it calculates the middle index, `mid`, comparing the middle element with the target value. If they match, the search is successful.
- If the target is smaller, the search space narrows to the lower half; if greater, to the upper half, following a "divide and conquer" strategy.
- The algorithm repeats this process until the target is found or the search space becomes empty.
- Binary search's time complexity of $O(\log n)$ makes it highly efficient for large datasets, and it is widely utilized in applications requiring swift and precise element retrieval from sorted arrays, such as databases and search engines.

2.3 Hashing

Hashing is a technique used in computer science to map data of arbitrary size to a fixed-size value, usually a hash code. The primary goal of hashing is to quickly locate a data record given its search key.

It involves using a hash function to compute an index into an array of buckets or slots, from which the desired value can be found. Ideally, a good hash function distributes the keys uniformly across the array, minimizing collisions (occurrences where two different keys map to the same index).

In the context of data structures, hashing is often employed in hash tables. Hash tables consist of an array and a hash function that maps keys to indices in the array. When searching for a value, the hash function is applied to the key to determine the index, allowing for fast retrieval.

However, collisions can occur when multiple keys produce the same index. Techniques such as chaining (using linked lists at each index) or open addressing (finding the next available slot) are employed to handle collisions and ensure efficient data retrieval.

Hashing provides a way to achieve constant-time average-case complexity for search, insertion, and deletion operations in hash tables, making it a fundamental concept in algorithm design and data structure implementation.

Hashing Functions

A hash function is a crucial component of hashing techniques used in computer science and data structures. Its primary purpose is to take an input, often of variable size, and map it to a fixed-size value, typically a hash code or hash value. The goal is to efficiently

distribute data across a predefined range, such as an array of buckets, allowing for quick and direct access to the desired information.

An effective hash function exhibits two key properties: it should produce a unique hash code for distinct inputs, minimizing collisions, and it should distribute the hash codes uniformly across the available range.

Achieving these properties ensures that the hash function provides efficient indexing and retrieval of data in hash-based data structures, such as hash tables. Well-designed hash functions contribute to the overall performance and effectiveness of algorithms and data structures by facilitating rapid access to stored information.

Summary

- Hashing is a technique in computer science used to map data of varying sizes to a fixed-size value, commonly known as a hash code.
- The primary objective is swift data retrieval based on a search key, achieved by applying a hash function to compute an index within an array of buckets or slots.
- Hashing is prominently used in hash tables, where the hash function maps keys to indices, enabling quick data retrieval.
- Effective hash functions distribute keys uniformly to minimize collisions, and occurrences when different keys map to the same index.
- Collisions are managed using techniques like chaining or open addressing. Hashing ensures constant-time average-case complexity for search, insertion, and deletion operations in hash tables, making it fundamental in algorithm design and data structures.

Collision Resolution

Collision resolution is a crucial aspect of hash table design, addressing the challenge when two different inputs produce the same hash code. Collisions can occur due to the finite range of hash codes and the infinite number of possible inputs.

Effectively managing collisions ensures the integrity and efficiency of hash-based data structures.

Various collision resolution techniques exist, each with its advantages and trade-offs. One common method is chaining, where each bucket in the hash table maintains a linked list or other data structure to accommodate multiple elements that share the same hash code.

This approach allows for the storage of multiple values at the same index, addressing collisions by creating a linked structure.

Another approach is open addressing, where the algorithm searches for the next available slot in the hash table when a collision occurs. Probing methods, such as linear probing and quadratic probing, are examples of open-addressing techniques.

They involve examining adjacent slots until an empty one is found, enabling the placement of colliding elements in alternative locations.

The choice of collision resolution technique impacts the efficiency and performance of hash tables, influencing factors like memory usage, retrieval speed, and ease of implementation.

Understanding and selecting an appropriate collision resolution strategy is essential for designing robust hash-based data structures that effectively handle the challenges posed by collisions.

Summary

- Collision resolution is a critical aspect of hash table design, addressing challenges when different inputs produce the same hash code.
- This issue arises due to the finite range of hash codes and the infinite possible inputs. Effectively managing collisions is vital for ensuring the integrity and efficiency of hash-based data structures.
- Various techniques exist, such as chaining, where each bucket maintains a linked list for colliding elements, and open addressing, involving methods like linear probing to find alternative slots.
- The choice of collision resolution impacts factors like memory usage and retrieval speed, making it crucial for designing efficient hash-based data structures.
- Selecting an appropriate strategy is essential for effectively handling collision challenges.

Linear Probing

Linear probing is a collision resolution technique used in hash tables to address situations where two different keys produce the same hash code.

When a collision occurs, linear probing involves probing the next available slot in the hash table until an empty slot is found. The idea is to linearly search for the next open position in the hash table, starting from the original hash location.

In linear probing, if the slot at the calculated hash index is already occupied, the algorithm iteratively checks the next slot until an empty position is found. When inserting a new key, the algorithm places it in the first available empty slot.

During searches, the algorithm follows the same linear sequence until it locates the target key or an empty slot, indicating the absence of the key in the table.

While linear probing is simple to implement and memory-efficient, it can lead to clustering, where consecutive slots become occupied, potentially impacting the overall performance.

Additionally, its performance may degrade as the load factor increases, affecting the number of occupied slots relative to the total number of slots in the hash table. Adjusting

the probing sequence or incorporating techniques like double hashing can help mitigate clustering and enhance the effectiveness of linear probing in managing collisions.

Summary

- Linear probing is a collision resolution technique in hash tables used to handle instances where two keys produce the same hash code.
- In the event of a collision, linear probing sequentially searches for the next available slot in the hash table, placing the key in the first empty position.
- While simple and memory-efficient, it may lead to clustering issues, where consecutive slots fill up, potentially impacting performance.
- Adjusting the probing sequence or employing strategies like double hashing can help mitigate these challenges and optimize linear probing's effectiveness in resolving collisions.

2.4 Sorting: General Background

Sorting is a foundational operation in computer science that involves organizing a collection of data into a specific order. The primary objective is to arrange items systematically, often in ascending or descending order based on a defined criterion.

Efficient sorting algorithms are critical for optimizing search operations, improving data retrieval speed, and facilitating various computational tasks.

Different sorting algorithms employ diverse strategies to achieve the desired order. Some common algorithms include Bubble Sort, QuickSort, Merge Sort, Insertion Sort, and Heap Sort.

These algorithms vary in their efficiency, stability, and suitability for different types of data and problem sizes. The efficiency of a sorting algorithm is typically assessed by its time complexity, indicating how its performance scales with the size of the input data.

The choice of a sorting algorithm depends on factors such as the nature of the data, the size of the dataset, and the specific requirements of the application.

Sorting is a fundamental concept applicable in diverse fields, from database management and search algorithms to everyday tasks like alphabetizing lists or sorting files.

A solid understanding of sorting algorithms is essential for computer scientists and programmers to develop efficient and robust software solutions.

Summary

- Sorting is a fundamental operation in computer science, focusing on arranging data systematically, often in ascending or descending order based on specific criteria.

- Efficient sorting algorithms are crucial for enhancing search operations, accelerating data retrieval, and supporting various computational tasks.
- Diverse sorting algorithms, including Bubble Sort, QuickSort, Merge Sort, Insertion Sort, and Heap Sort, offer distinct strategies with varying efficiency, stability, and suitability for different data types and sizes.
- The choice of a sorting algorithm depends on factors such as data nature, dataset size, and application requirements.
- Sorting is a universal concept with applications in database management, search algorithms, and everyday tasks, underscoring its significance in developing efficient and robust software solutions.

2.5 Sorting Technique

Sorting is a fundamental operation in computer science that involves arranging elements in a specific order to facilitate efficient searching, retrieval, and analysis of data. Various sorting techniques employ distinct strategies to achieve this organization.

One common approach is comparison-based sorting, where elements are compared pairwise, and their positions are adjusted accordingly. Well-known comparison-based sorting algorithms include Bubble Sort, Insertion Sort, Merge Sort, QuickSort, and Heap Sort. These algorithms differ in terms of their time complexity, stability, and suitability for different types and sizes of datasets.

Another category of sorting techniques is non-comparison-based sorting, which exploits specific properties of the data to arrange elements. An example is Counting Sort, which works well for integers within a known range.

Radix Sort is another non-comparison-based algorithm that sorts elements based on their digits. Each sorting technique has its strengths and weaknesses, making it essential to choose the most appropriate algorithm based on the characteristics of the data being sorted.

The efficiency of sorting algorithms is often evaluated in terms of their time complexity, indicating how the runtime scales with the size of the input data.

As sorting is a fundamental operation with broad applications in computer science, understanding different sorting techniques and their performance characteristics is crucial for designing efficient algorithms and building robust software systems.

Summary

- Sorting is a foundational operation in computer science, essential for organizing elements to enhance searching and retrieval efficiency.
- Sorting techniques employ diverse strategies, categorized into comparison-based (e.g., Bubble Sort, Merge Sort) and non-comparison-based (e.g., Counting Sort, Radix Sort).

- These algorithms vary in time complexity, stability, and suitability for different data types and sizes.
- Efficient sorting is crucial for optimal algorithm performance, and understanding various sorting techniques is fundamental in designing effective algorithms and building resilient software systems.

Bubble Sort

Bubble Sort is a simple sorting algorithm that repeatedly steps through the list of elements, compares adjacent items, and swaps them if they are in the wrong order.

The algorithm gets its name from the way smaller elements "bubble" to the top of the list during each pass.

The basic idea is to iterate through the entire list multiple times, comparing adjacent elements and swapping them until the entire list is sorted.

In each iteration, the largest unsorted element gets placed in its correct position at the end of the list. The process continues until no more swaps are needed, indicating that the list is fully sorted.

While Bubble Sort is easy to understand and implement, it is not the most efficient sorting algorithm, especially for large datasets. Its average and worst-case time complexities are both $O(n^2)$, making it less practical compared to more advanced algorithms like QuickSort or Merge Sort.

Despite its simplicity, Bubble Sort is primarily used for educational purposes rather than in practical applications.

Bubble Sort is a straightforward sorting algorithm that repeatedly steps through the list, compares adjacent elements, and swaps them if they are in the wrong order. Here are the steps to perform Bubble Sort:

Steps to Do Bubble Sort:

1. Initialize: Start with the first element (index 0) of the array.

2. Compare and Swap: Compare the current element with the next element. If the current element is greater than the next one, swap them. This process is repeated for every adjacent pair in the array.

3. Move to the Next Pair: Move to the next pair of elements (i.e., advance to the next index).

4. Continue Comparing and Swapping: Repeat steps 2 and 3 until you reach the end of the array. At this point, the largest element in the unsorted part of the array will have "bubbled up" to its correct position at the end.

5. Repeat for Remaining Elements: Repeat steps 1-4 for the remaining unsorted elements. After the first iteration, the largest element is in its correct place. After the second iteration, the second largest element is in its place, and so on.

6. Optimization: Bubble Sort can be optimized by introducing a flag that tracks whether any swaps were made during an iteration. If no swaps are made, the array is already sorted, and the algorithm can terminate early.

7. Repeat Until Sorted: Continue the process until the entire array is sorted. The number of iterations needed is equal to the number of elements in the array.

Example:

```
#include <iostream>

// Function to perform Bubble Sort on an array
void bubbleSort(int arr[], int n) {
    for (int i = 0; i < n - 1; ++i) {
        // Last i elements are already sorted, so no need to check them
        for (int j = 0; j < n - i - 1; ++j) {
            // Swap if the element found is greater than the next element
            if (arr[j] > arr[j + 1]) {
                // Swap arr[j] and arr[j + 1]
                int temp = arr[j];
                arr[j] = arr[j + 1];
                arr[j + 1] = temp;
            }
        }
    }
}

// Function to print the elements of an array
void printArray(int arr[], int size) {
    for (int i = 0; i < size; ++i) {
        std::cout << arr[i] << " ";
    }
    std::cout << std::endl;
}

int main() {
    // Example array
    int arr[] = {5, 2, 9, 1, 5};
    int size = sizeof(arr) / sizeof(arr[0]);

    std::cout << "Original Array: ";
    printArray(arr, size);

    // Perform Bubble Sort
    bubbleSort(arr, size);

    std::cout << "Sorted Array: ";
    printArray(arr, size);

    return 0;
}
```

This algorithm repeatedly steps through the list, compares adjacent elements, and swaps them until the entire list is sorted. The optimization check ensures early termination if the array is already sorted before completing all iterations.

Summary

- Bubble Sort, a simple sorting algorithm, operates by iteratively comparing adjacent elements and swapping them if they are in the wrong order, causing smaller elements to "bubble" to the top during each pass.
- The process repeats until the entire list is sorted, with the largest element reaching its correct position at the end after each iteration.
- Despite its easy implementation, Bubble Sort has $O(n^2)$ time complexity in average and worst cases, making it inefficient for large datasets compared to more advanced algorithms.
- Primarily used for educational purposes, Bubble Sort's straightforward nature and the possibility of optimization with a swap-checking flag contribute to its understanding but limit its practical application.

Insertion Sort

Insertion Sort is a simple yet efficient sorting algorithm that builds the final sorted array one element at a time. It is particularly useful for small datasets or partially sorted datasets.

The algorithm works by repeatedly taking an unsorted element and inserting it into its correct position within the already sorted part of the array.

The process begins with the assumption that the first element of the array is already sorted. The algorithm then iterates through the unsorted portion of the array, comparing each element with the elements in the sorted part.

When it finds the correct position for the current element, it shifts the larger elements to the right to make room for the insertion. This process continues until the entire array is sorted.

Insertion Sort has several advantages, including its simplicity and adaptability. It performs well for small datasets or nearly sorted datasets but becomes less efficient for larger datasets due to its quadratic time complexity.

However, its ease of implementation and efficiency for small datasets make it a suitable choice in certain scenarios, especially when dealing with nearly ordered data. Overall, Insertion Sort's straightforward approach and flexibility make it a valuable sorting algorithm in various applications.

Steps To Do Insertion Sort

Insertion Sort is a straightforward sorting algorithm that builds the final sorted array one element at a time. Here are the steps to perform Insertion Sort:

1. Initialization: Assume the first element in the array is already sorted. Start iterating from the second element (index 1) as the current unsorted element.

2. Comparison and Insertion: Compare the current unsorted element with the elements in the sorted part of the array. Move elements greater than the current element to the

right to make space for the insertion.

3. Repeat: Continue this process for each unsorted element until the entire array is sorted.

The algorithm works by iteratively taking an unsorted element and placing it in its correct position within the sorted part of the array. This process continues until the entire array is sorted.

Example:

```
#include <iostream>
using namespace std;

void insertionSort(int arr[], int n) {
    for (int i = 1; i < n; i++) {
        int currentElement = arr[i];
        int j = i - 1;

        // Move elements of arr[0..i-1] that are greater than currentElement
        // to one position ahead of their current position
        while (j >= 0 && arr[j] > currentElement) {
            arr[j + 1] = arr[j];
            j--;
        }

        // Insert the currentElement in its correct position
        arr[j + 1] = currentElement;
    }
}

int main() {
    int myArray[] = {12, 11, 13, 5, 6};
    int arraySize = sizeof(myArray) / sizeof(myArray[0]);

    insertionSort(myArray, arraySize);

    // Output the sorted array
    cout << "Sorted array: ";
    for (int i = 0; i < arraySize; i++) {
        cout << myArray[i] << " ";
    }

    return 0;
}
```

This algorithm is suitable for small datasets or partially sorted datasets, but its time complexity becomes quadratic for larger datasets. It is, however, an efficient choice in scenarios where the array is nearly sorted.

Summary

- Insertion Sort is a simple yet efficient sorting algorithm that builds the final sorted array incrementally.
- It starts with the assumption that the first element is sorted and iteratively inserts each subsequent unsorted element into its correct position within the already sorted part of the array.
- By comparing and shifting elements as needed, the algorithm continues until the entire array is sorted.

- While Insertion Sort's time complexity becomes quadratic for larger datasets, its ease of implementation and adaptability make it a valuable choice for small datasets or nearly sorted arrays, showcasing its straightforward approach and flexibility in various applications.

Selection Sort

Selection Sort is a straightforward sorting algorithm that divides an input list into two portions: a sorted region and an unsorted region.

The algorithm iteratively selects the smallest (or largest) element from the unsorted region and swaps it with the first element of the unsorted part, expanding the sorted region.

This process continues until the entire list is sorted. Although simple to understand and implement, Selection Sort has a time complexity of $O(n^2)$, making it less efficient than more advanced sorting algorithms for larger datasets.

Steps to perform Selection Sort:

1. Initialization

- Assume the entire array is unsorted initially.
- Define the first element as the minimum (or maximum) for comparison.

2. Find Minimum (or Maximum)

- Iterate through the unsorted part of the array to find the minimum (or maximum) element.
- Swap the minimum (or maximum) element with the first element of the unsorted part.

3. Expand Sorted Region's

- Expand the sorted region by considering the next element in the unsorted part as the new starting point.
- Repeat steps 2 and 3 until the entire array is sorted.

Example:

```

#include <iostream>

void selectionSort(int arr[], int n) {
    for (int i = 0; i < n - 1; ++i) {
        // Assume the current index is the minimum
        int minIndex = i;

        // Find the index of the minimum element in the unsorted part
        for (int j = i + 1; j < n; ++j) {
            if (arr[j] < arr[minIndex]) {
                // Update the minimum index if a smaller element is found
                minIndex = j;
            }
        }

        // Swap the found minimum element with the first element of the unsorted part
        std::swap(arr[i], arr[minIndex]);
    }
}

int main() {
    int arr[] = {64, 25, 12, 22, 11};
    int n = sizeof(arr) / sizeof(arr[0]);

    // Call the selectionSort function to sort the array
    selectionSort(arr, n);

    // Print the sorted array
    std::cout << "Sorted array: ";
    for (int i = 0; i < n; ++i)
        std::cout << arr[i] << " ";

    return 0;
}

```

Summary

- Selection Sort is a simple sorting algorithm that organizes an input list into a sorted section and an unsorted section.
- In each iteration, the algorithm identifies the smallest (or largest) element in the unsorted region and swaps it with the first element of the unsorted part, gradually expanding the sorted region.
- This process repeats until the entire list is sorted. While easy to grasp and implement, Selection Sort's quadratic time complexity ($O(n^2)$) makes it less efficient for larger datasets compared to more advanced sorting methods.
- The algorithm's sequential selection and swapping strategy contribute to its simplicity but limit its practicality for extensive data sets.

This C++ example demonstrates the implementation of Selection Sort on an array and prints the sorted result. The algorithm repeatedly finds the minimum element in the unsorted part and swaps it with the first element, expanding the sorted region until the entire array is sorted.

Merge Sort

Merge Sort is a divide-and-conquer algorithm designed to efficiently sort an array or list by recursively dividing it into smaller halves, sorting each half, and then merging the sorted halves back together. The process starts by dividing the unsorted array into two halves, and this division continues until each subarray contains only one element.

The sorting begins during the merging phase, where pairs of adjacent subarrays are merged into larger sorted subarrays.

This merging process repeats until the entire array is sorted. Merge Sort's key strength lies in its stable $O(n \log n)$ time complexity, making it efficient for large datasets, although it requires additional space for the merging step.

The Merge Sort algorithm works by dividing the array into halves, sorting each half, and then merging them back together in sorted order.

This divide-and-conquer strategy ensures a reliable and efficient sorting process with a time complexity of $O(n \log n)$, making it suitable for handling large datasets.

The algorithm's primary drawback is the need for additional memory space during the merging step, limiting its use in memory-constrained environments. Overall, Merge Sort strikes a balance between efficiency and simplicity, offering a reliable solution for sorting tasks.

Steps to perform Merge Sort

1. Divide

- Divide the unsorted array into two halves.
- Continue recursively dividing each subarray until each subarray has only one element.

2. Conquer

- Merge pairs of adjacent subarrays by comparing and arranging their elements in sorted order.
- Continue merging until all subarrays are merged into a single sorted array.

3. Merge

Combine the sorted subarrays into larger sorted arrays until the entire array is sorted.

Example:

part-1

```

#include <iostream>
#include <vector>

void merge(std::vector<int>& arr, int left, int mid, int right) {
    int n1 = mid - left + 1;
    int n2 = right - mid;

    std::vector<int> leftArr(n1);
    std::vector<int> rightArr(n2);

    // Copy data to temporary arrays leftArr[] and rightArr[]
    for (int i = 0; i < n1; i++)
        leftArr[i] = arr[left + i];
    for (int j = 0; j < n2; j++)
        rightArr[j] = arr[mid + 1 + j];

    // Merge the temporary arrays back into arr[left..right]
    int i = 0;
    int j = 0;
    int k = left;

    while (i < n1 && j < n2) {
        if (leftArr[i] <= rightArr[j]) {
            arr[k] = leftArr[i];
            i++;
        } else {
            arr[k] = rightArr[j];
            j++;
        }
        k++;
    }

    // Copy the remaining elements of leftArr[], if there are any
    while (i < n1) {
        arr[k] = leftArr[i];
        i++;
        k++;
    }

    // Copy the remaining elements of rightArr[], if there are any
    while (j < n2) {
        arr[k] = rightArr[j];
        j++;
        k++;
    }
}

```

Part-2

```

void mergeSort(std::vector<int>& arr, int left, int right) {
    if (left < right) {
        // Same as (left+right)/2, but avoids overflow
        int mid = left + (right - left) / 2;

        // Recursively sort both halves
        mergeSort(arr, left, mid);
        mergeSort(arr, mid + 1, right);

        // Merge the sorted halves
        merge(arr, left, mid, right);
    }
}

int main() {
    std::vector<int> arr = {12, 11, 13, 5, 6, 7};

    std::cout << "Original array: ";
    for (const auto& elem : arr)
        std::cout << elem << " ";
    std::cout << std::endl;

    mergeSort(arr, 0, arr.size() - 1);

    std::cout << "Sorted array: ";
    for (const auto& elem : arr)
        std::cout << elem << " ";
    std::cout << std::endl;

    return 0;
}

```

This C++ code demonstrates the implementation of the Merge Sort algorithm on an array. The `mergeSort` function recursively divides and sorts the array, and the `merge` function merges the sorted halves back together. The example prints the original and sorted arrays for illustration.

Summary

- Merge Sort is a highly efficient divide-and-conquer algorithm designed for sorting arrays or lists.
- The process is initiated by recursively dividing the unsorted array into halves until each subarray contains only one element.
- Sorting takes place during the merging phase, where adjacent subarrays are systematically merged into larger sorted subarrays.
- This merging process repeats until the entire array is sorted. Notably, Merge Sort boasts a stable $O(n \log n)$ time complexity, making it particularly effective for large datasets.
- However, its drawback lies in the additional memory space required for the merging step, which may limit its application in memory-constrained environments.
- Despite this, Merge Sort strikes a commendable balance between efficiency and simplicity, providing a reliable solution for diverse sorting tasks.

Radix Sort

Radix Sort is a linear time, non-comparative sorting algorithm that operates on integers by processing individual digits. Unlike comparison-based sorting algorithms, Radix Sort distributes elements into buckets based on their digits, starting from the least significant digit (LSD) and progressing towards the most significant digit (MSD).

The algorithm performs multiple passes, each time grouping elements into buckets according to their digit values. After processing all digits, the elements are collected in the order they appear in the buckets, resulting in a fully sorted array.

Radix Sort's time complexity is linear for each pass, making it efficient for certain types of data, particularly when the number of digits is relatively small.

In each pass of Radix Sort, the algorithm processes the digits of the elements, placing them into buckets based on the values of those digits. This process repeats for each digit, ensuring that the elements are eventually sorted.

Radix Sort is often employed with fixed-length keys, and its linear time complexity per pass makes it well-suited for scenarios where the range of key values is limited. However, its practicality may diminish when dealing with variable-length keys or large key ranges.

Overall, Radix Sort offers an interesting alternative to traditional comparison-based sorting algorithms, demonstrating effectiveness in certain specific contexts.

Steps to Perform Radix Sort

1. Initialization

Identify the maximum number in the array to determine the number of digits in it.

Initialize a set of buckets, one for each digit (0–9).

2. Pass through Digits

Starting from the least significant digit (LSD) and moving towards the most significant digit (MSD), perform the following steps for each digit place:

2.1. Bucket Distribution

Distribute the elements into the buckets based on the value of the current digit.

2.2. Collect from Buckets

Collect the elements from the buckets in the order they appear, ensuring stability in sorting.

3. Repeat for Each Digit

Repeat the above process for each digit place until all digits have been processed.

4. Array is Sorted

After processing the most significant digit, the array is fully sorted.

Example:

```

#include <iostream>
#include <vector>

void radixSort(std::vector<int>& arr) {
    // Find the maximum number to determine the number of digits
    int maxNum = *max_element(arr.begin(), arr.end());

    // Perform counting sort for every digit
    for (int exp = 1; maxNum / exp > 0; exp *= 10) {
        countingSort(arr, exp);
    }
}

void countingSort(std::vector<int>& arr, int exp) {
    const int n = arr.size();
    std::vector<int> output(n, 0);
    std::vector<int> count(10, 0);

    // Count occurrences of each digit in the current place
    for (int i = 0; i < n; i++)
        count[(arr[i] / exp) % 10]++;

    // Update count[i] to store the position of the next occurrence
    for (int i = 1; i < 10; i++)
        count[i] += count[i - 1];

    // Build the output array
    for (int i = n - 1; i >= 0; i--) {
        output[count[(arr[i] / exp) % 10] - 1] = arr[i];
        count[(arr[i] / exp) % 10]--;
    }

    // Copy the output array to arr[]
    for (int i = 0; i < n; i++)
        arr[i] = output[i];
}

int main() {
    std::vector<int> arr = {170, 45, 75, 90, 802, 24, 2, 66};

    std::cout << "Original array: ";
    for (const auto& elem : arr)
        std::cout << elem << " ";
    std::cout << std::endl;

    radixSort(arr);

    std::cout << "Sorted array: ";
    for (const auto& elem : arr)
        std::cout << elem << " ";
    std::cout << std::endl;

    return 0;
}

```

This C++ code demonstrates the implementation of Radix Sort on an array. The `radixSort` function iteratively performs counting sort for each digit place, and the `countingSort` function is used to distribute and collect elements based on the current digit place. The example prints the original and sorted arrays for illustration.

Summary

- Radix Sort is a unique linear time sorting algorithm that efficiently handles integers by processing individual digits.
- By distributing elements into buckets based on their digits, starting from the least significant to the most significant digit, Radix Sort performs multiple passes to group elements in buckets according to their digit values.
- The process continues until all digits are processed, and elements are collected in the order they appear in the buckets, resulting in a fully sorted array.
- With a linear time complexity for each pass, Radix Sort is particularly effective for datasets with a limited number of digits, showcasing its utility in specific scenarios.

- Although practicality may decrease with variable-length keys or large key ranges, Radix Sort offers a compelling alternative to traditional comparison-based sorting algorithms, demonstrating efficiency in certain contexts.

Shell Sort

Shell Sort, also known as diminishing increment sort, is an in-place comparison-based sorting algorithm. It improves upon the insertion sort by allowing the comparison and exchange of elements that are far apart before progressively reducing the gap between elements.

The algorithm starts by dividing the array into smaller subarrays, each of which is sorted using an insertion sort. The unique feature of Shell Sort lies in the choice of gap sequences, which determine how elements are compared and swapped across the subarrays.

As the algorithm progresses, the gaps decrease until the entire array becomes nearly sorted, making the final pass with a gap of 1 equivalent to a regular insertion sort. While not as efficient as more advanced algorithms, Shell Sort strikes a balance between simplicity and improved performance compared to basic sorting methods for medium-sized datasets.

Shell Sort is characterized by its gap sequence, which determines the steps at which elements are compared and swapped. The algorithm initially considers larger gaps and gradually reduces them until the final pass with a gap of 1, resembling a standard insertion sort.

Despite its simplicity and adaptability to various gap sequences, Shell Sort's time complexity is not as favorable as more advanced algorithms, particularly for large datasets.

However, its intermediate sorting stages make it more efficient than basic sorting algorithms like insertion sort, especially for datasets of moderate size. Overall, Shell Sort provides a compromise between simplicity and enhanced performance for specific sorting scenarios.

Steps to Perform Shell Sort

1. Choose Gap Sequence

Select a gap sequence that determines the steps at which elements are compared and swapped during the sorting process. Common gap sequences include the Knuth sequence ($3k+1$) or simply dividing the array size by 2.

2. Start with the Widest Gap

Initialize the gap as the first value from the chosen gap sequence.

Iterate over the array, comparing and swapping elements that are separated by the current gap.

3. Reduce the Gap and Stack

Continue reducing the gap according to the chosen sequence.

Repeat the comparison and swapping process for each reduced gap until the final pass with a gap of 1.

4. Final Pass with Insertion Sort

Perform a final pass using an insertion sort with a gap of 1, essentially sorting the nearly sorted array.

Example :

```
#include <iostream>
#include <vector>

void shellSort(std::vector<int>& arr) {
    int n = arr.size();

    // Choose gap sequence, e.g., Knuth sequence (3k+1)
    for (int gap = n / 2; gap > 0; gap /= 2) {
        // Perform insertion sort for the current gap
        for (int i = gap; i < n; i++) {
            int temp = arr[i];
            int j;

            for (j = i; j >= gap && arr[j - gap] > temp; j -= gap)
                arr[j] = arr[j - gap];

            arr[j] = temp;
        }
    }
}

int main() {
    std::vector<int> arr = {12, 34, 54, 2, 3};

    std::cout << "Original array: ";
    for (const auto& elem : arr)
        std::cout << elem << " ";
    std::cout << std::endl;

    shellSort(arr);

    std::cout << "Sorted array: ";
    for (const auto& elem : arr)
        std::cout << elem << " ";
    std::cout << std::endl;

    return 0;
}
```

This C++ code demonstrates the implementation of Shell Sort on an array. The `shell sort` function performs the sorting using a chosen gap sequence, and the example prints the original and sorted arrays for illustration.

Summary

- Shell Sort, an in-place comparison-based sorting algorithm, enhances the efficiency of insertion sort by allowing distant elements to be compared and exchanged before progressively reducing the gap between elements.

- The algorithm begins by dividing the array into smaller subarrays, each sorted with an insertion sort.
 - Notably, Shell Sort's distinctiveness lies in its gap sequences, determining how elements are compared and swapped across subarrays.
 - As the algorithm advances, the gaps decrease until the array becomes nearly sorted, culminating in a final pass with a gap of 1, similar to a standard insertion sort.
 - Despite not matching the efficiency of advanced algorithms for larger datasets, Shell Sort strikes a practical balance between simplicity and performance, particularly suited for medium-sized datasets.
 - The algorithm's adaptability to various gap sequences and its intermediate sorting stages makes it more efficient than basic sorting methods, providing a valuable compromise for specific sorting scenarios.
-



CHAPTER 3: STACK , QUEUES AND RECURSION



3.1 Stacks, Queues, And Recursion

Stacks

Stacks are data structures that follow the Last In, First Out (LIFO) principle, where the last element added is the first one to be removed. Elements are added or removed from the top of the stack.

This structure is used in various applications, such as managing function calls in recursion, tracking backtracking algorithms, and parsing expressions.

Queues

Queues adhere to the First In, First Out (FIFO) principle, where the first element added is the first one to be removed.

Elements are added at the rear, and removal occurs from the front. Queues find applications in scenarios like task scheduling, breadth-first search algorithms, and managing resources in a shared environment.

Recursion

Recursion is a programming concept where a function calls itself during its execution. It provides an elegant and powerful solution to problems that exhibit self-similar substructures. Recursion is commonly used in algorithms for tasks like traversing tree structures, solving mathematical problems, and implementing divide-and-conquer strategies.

However, careful consideration of base cases is crucial to prevent infinite loops. The combination of stacks and recursion often goes hand in hand, as the call stack keeps track of function calls in a recursive process. These fundamental concepts—stacks, queues, and recursion—play crucial roles in algorithm design and efficient problem-solving strategies.

Summary

- Stacks, queues, and recursion are fundamental concepts in computer science and programming.
- Stacks, based on the Last In, First Out (LIFO) principle, are versatile structures used in managing function calls in recursion, tracking backtracking algorithms, and parsing expressions.
- Queues, operating on the First In, First Out (FIFO) principle, are employed in scenarios such as task scheduling, breadth-first search algorithms, and resource management.
- Recursion, a powerful programming concept where a function calls itself, proves invaluable for solving problems with self-similar substructures, like traversing tree structures and implementing divide-and-conquer strategies.
- The careful consideration of base cases is essential to prevent infinite loops. Stacks and recursion often complement each other, as the call stack tracks function calls in recursive processes.
- These foundational concepts—stacks, queues, and recursion—play vital roles in algorithm design and efficient problem-solving strategies.

3.2 Stack

A stack is a fundamental data structure that follows the Last In, First Out (LIFO) principle. It operates much like a stack of items, where the last item placed on top is the first one to be removed. Elements are added or removed from the top of the stack, known as the "push" and "pop" operations, respectively.

Stacks find versatile applications in programming and computer science. One common use is in managing function calls during program execution. As functions are called, their information is pushed onto the stack, and when a function completes, its information is popped off the stack. This behavior allows for efficient management of nested function calls.

Additionally, stacks play a crucial role in backtracking algorithms. In scenarios where decisions need to be undone or revisited, a stack can store the necessary information, enabling the algorithm to backtrack and explore alternative paths.

Stacks are also employed in parsing expressions, particularly in evaluating mathematical expressions. As the algorithm processes each symbol, it uses a stack to keep track of operands and operators, ensuring a proper order of evaluation.

Summary

- A stack, adhering to the Last In, First Out (LIFO) principle, is a foundational data structure employed in programming and computer science.
- Operations such as "push" and "pop" facilitate the addition and removal of elements from the top of the stack.
- Commonly used in managing function calls, a stack efficiently organizes information during program execution, ensuring proper handling of nested function calls.

- In the realm of backtracking algorithms, stacks become crucial for storing decision information, allowing the algorithm to revisit and explore alternative paths.
- Additionally, in parsing expressions, particularly mathematical ones, stacks play a pivotal role in tracking operands and operators, ensuring the accurate order of evaluation.
- This versatility positions stacks as an integral component in various computing applications.

Concept

The concept of a stack revolves around the Last In, First Out (LIFO) principle, where the last item added to the stack is the first one to be removed. Stacks are often envisioned as a collection of items arranged like a stack of plates, where you can only add or remove items from the top. The two primary operations associated with stacks are:

1. Push

Adding an item to the top of the stack. This operation is analogous to placing a new plate on top of an existing stack.

2. Pop

Removing the item from the top of the stack. This operation is similar to taking the top plate off the stack.

These operations are complemented by a third operation:

3. Peek (or Top)

Observing the item at the top of the stack without removing it. This operation provides information about the current top element.

The concept of a stack is widely utilized in various computing scenarios, including managing function calls in programming languages. As functions are called, their information is pushed onto the call stack, and when a function completes, its information is popped off the stack. This mechanism ensures that the program returns to the correct execution point after each function call.

Stacks are also employed in undo functionality, where the history of actions is stored in a stack, allowing users to undo actions in reverse order. Additionally, algorithms involving backtracking, parsing expressions, and managing state information often leverage the stack concept to efficiently handle data and control flow.

Summary

- The stack concept, operating on the Last In, First Out (LIFO) principle, involves adding and removing items from the top of the stack.

- This structure, akin to a stack of plates, features fundamental operations: "Push" adds an item to the top, while "Pop" removes the top item.
- The "Peek" operation observes the top item without removal. In computing, stacks play a vital role in managing function calls, storing information during execution, and ensuring a seamless return to the correct program point after each function call.
- Additionally, stacks find application in undo functionality, maintaining a history of actions for easy reversal.
- This versatile concept is integral to algorithms dealing with backtracking, parsing expressions, and efficient data and control flow management.

Representing a stack in C++

In C++, a stack can be represented using the Standard Template Library (STL) container `std::stack`. This container provides a convenient and efficient implementation of a stack. To use it, you need to include the `<stack>` header.

Here is a simple representation of a stack in C++ using `std::stack`.

```
#include <iostream>
#include <stack>

int main() {
    // Creating a stack of integers
    std::stack<int> myStack;

    // Pushing elements onto the stack
    myStack.push(10);
    myStack.push(20);
    myStack.push(30);

    // Displaying the top element
    std::cout << "Top element: " << myStack.top() << std::endl;

    // Popping the top element
    myStack.pop();

    // Displaying the remaining elements
    std::cout << "Remaining elements: ";
    while (!myStack.empty()) {
        std::cout << myStack.top() << " ";
        myStack.pop();
    }

    return 0;
}
```

In this example, a stack of integers is created using `std::stack<int> myStack`. Elements are pushed onto the stack using `push()`, and the top element is accessed and displayed using `top()`.

The `pop()` operation removes the top element, and the remaining elements are displayed using a loop until the stack is empty. The `std::stack` template handles the underlying operations, making it easy to work with stacks in C++.

Application of Stack

Expression Evaluation

One prominent application of stacks is in the evaluation of expressions. Stacks help maintain the order of operations and handle parentheses efficiently. As the expression is parsed, operands are pushed onto the stack, and when an operator is encountered, the necessary operands are popped, and the result is pushed back.

This ensures the correct order of evaluation, and the stack's Last In, First Out (LIFO) nature aligns well with the nested structure of expressions.

Function Call Management

In programming languages, the call stack plays a crucial role in managing function calls. As functions are called, information about each call, including local variables and the return address, is pushed onto the call stack.

When a function completes, its information is popped off the stack, allowing the program to return to the correct execution point. This stack-based mechanism ensures the integrity and order of function calls in a program.

Undo Mechanism in Applications

Many applications implement an undo mechanism using a stack. As users perform actions, details about these actions are pushed onto a stack. If the user decides to undo an action, the application pops the top of the stack, reverting the state to the previous action.

This approach provides a convenient way to manage and execute undo operations, offering users the flexibility to backtrack through their interactions.

Summary

- Stacks play a crucial role in various applications, with expression evaluation being a prominent use case.
- They facilitate the systematic parsing of expressions by pushing operands onto the stack and efficiently handling operators, ensuring the correct order of evaluation.
- In programming, stacks are essential for function call management, as the call stack maintains information about each function call, including local variables and return addresses.
- This stack-based mechanism guarantees the integrity and orderly execution of functions.
- Additionally, stacks are instrumental in implementing undo mechanisms in applications, where details of user actions are stored on a stack, allowing for easy reversal of operations and providing users with a flexible way to backtrack through their interactions.

3.3 Polish Notation (Prefix, Postfix, Infix)

Polish Notation, also known as Prefix Notation, is a mathematical notation in which every operator follows all of its operands. This eliminates the need for parentheses to indicate the order of operations. In Polish Notation, the expression is written in a way that makes the sequence of operations clear without relying on the traditional order of precedence.

For example, consider the infix expression $3 + 4 * 5$. In Polish Notation, this expression is written as $+ 3 * 4 5$. Here, the plus sign precedes its operands, and the multiplication is unambiguously applied to the operands 4 and 5.

One of the advantages of Polish Notation is that it is easily parsed by computers, as the sequence of operators and operands is explicit. It is particularly useful in stack-based algorithms for evaluating mathematical expressions, where operands are pushed onto a stack, and operators are applied when encountered.

This notation was introduced by the Polish mathematician Jan Łukasiewicz and is sometimes called "Łukasiewicz notation" in his honor.

Prefix

In mathematics and computer science, "prefix" typically refers to a notation or expression in which an operator comes before its operands.

This notation is known as Prefix Notation. In a prefix expression, the order of operations is unambiguous as the position of the operator determines the operation to be performed.

For example, consider the infix expression $3 + 4 * 5$. In prefix notation, this expression would be written as $+ 3 * 4 5$. Here, the plus sign precedes its two operands, and the multiplication is applied to the operands 4 and 5. The order of operations is explicitly indicated by the position of operators.

Prefix notation is particularly useful in computer science and programming, especially in the context of evaluating mathematical expressions. It is easily parsed by computers, making it suitable for use in algorithms and mathematical computations.

Summary

- Prefix notation, where an operator precedes its operands, is a mathematical expression format widely used in computer science and programming.
- In a prefix expression like $+ 3 * 4 5$, the unambiguous order of operations is determined by the position of the operator.
- This notation is highly suitable for evaluating mathematical expressions in algorithms, as it is easily parsed by computers.
- The explicit indication of operation order by the operator's position makes prefix notation efficient and advantageous in various computational contexts, contributing to its prevalence in programming languages and algorithmic implementations.

Postfix

In mathematics and computer science, "postfix" refers to a notation or expression in which an operator follows its operands. This notation is known as Postfix Notation, also called Reverse Polish Notation (RPN).

In a postfix expression, the order of operations is unambiguous as the position of the operator determines the operation to be performed.

For example, consider the infix expression "3 + 4 * 5." In postfix notation, this expression would be written as "3 4 5 * +." Here, the operands 4 and 5 are placed before their multiplication operator, and then the result is added to the operand 3.

The order of operations is explicitly indicated by the position of operands and operators.

Postfix notation is particularly useful in computer science and programming, especially in the context of evaluating mathematical expressions. It is easily parsed by computers, making it suitable for use in algorithms and mathematical computations.

Postfix notation is often employed in stack-based algorithms, where operands are pushed onto a stack, and operators are applied when encountered.

Summary

- Postfix notation, also known as Reverse Polish Notation (RPN), is a mathematical expression format where operators follow their operands, ensuring unambiguous order of operations.
- For instance, the infix expression "3 + 4 * 5" becomes "3 4 5 * +." Here, operands precede their multiplication operator, and the result is then added to the initial operand.
- Postfix notation is beneficial in computer science and programming, particularly for evaluating mathematical expressions.
- Its explicit indication of operation order by operand and operator positions makes it easily parseable by computers.
- Widely used in stack-based algorithms, postfix notation involves pushing operands onto a stack and applying operators as encountered, contributing to its efficiency in various computational contexts.

Infix

In mathematics and computer science, "infix" refers to a common notation or expression in which an operator is placed between its two operands.

In infix notation, the order of operations is often indicated by parentheses to ensure clarity and to override the default precedence of operators. This notation is the standard

way mathematical expressions are written in most textbooks and used in everyday mathematical communication.

For example, consider the infix expression "3 + 4 * 5." Here, the plus sign is placed between the two operands, 3, and the result of the multiplication operation (4 * 5).

The multiplication operation takes precedence due to the usual rules of arithmetic, and the addition is performed afterward.

While infix notation is intuitive and widely used in mathematical expressions, it can lead to ambiguity in complex expressions, requiring the use of parentheses to specify the desired order of operations.

In contrast to postfix and prefix notations, infix notation may require additional parsing mechanisms to determine the correct evaluation order in computer algorithms.

Summary

- Infix notation, prevalent in mathematics and computer science, involves placing an operator between two operands, such as the plus sign in "3 + 4 * 5."
- The order of operations is often clarified with parentheses, ensuring proper precedence. While intuitive, infix notation can introduce ambiguity in complex expressions, necessitating the use of parentheses to specify the desired evaluation order.
- In contrast to postfix and prefix notations, infix notation, standard in mathematical communication, may require additional parsing mechanisms for correct evaluation in computer algorithms.

3.4 Quick Sort

Quick Sort is a widely used and efficient sorting algorithm that follows the divide-and-conquer strategy. The algorithm, developed by Tony Hoare in 1960, works by selecting a "pivot" element from the array and partitioning the other elements into two sub-arrays according to whether they are less than or greater than the pivot.

This process is recursively applied to the sub-arrays until the entire array is sorted. QuickSort has an average-case time complexity of $O(n \log n)$ and performs exceptionally well in practice, making it a popular choice for various applications.

The key steps of the QuickSort algorithm involve selecting a pivot, partitioning the array, and recursively sorting the sub-arrays. The efficiency of QuickSort stems from its ability to minimize the number of comparisons and swaps needed for sorting.

However, in its basic form, QuickSort can degrade to $O(n^2)$ time complexity in the worst case if the pivot selection is poorly optimized. Strategies such as choosing a random pivot or implementing a three-way partitioning scheme can mitigate this issue, making QuickSort a versatile and effective sorting method.

Despite its potential for worst-case scenarios, QuickSort's average-case performance, simplicity of implementation, and low overhead make it a popular sorting algorithm in practice.

Many programming languages and libraries use QuickSort as their default sorting method due to its speed and efficiency for a wide range of scenarios.

Steps To Do Quick Sort

QuickSort is a divide-and-conquer sorting algorithm that works by selecting a pivot element and partitioning the array into two sub-arrays based on the pivot. Here are the steps to perform QuickSort.

1. Choose a Pivot

Select a pivot element from the array. The choice of the pivot can impact the algorithm's efficiency, and various strategies exist, such as selecting the first, last, or middle element, or even using a randomized approach.

2. Partition the Array

Rearrange the elements in the array such that elements smaller than the pivot are on the left, and elements greater than the pivot are on the right. The pivot itself is now in its final sorted position. This process is often referred to as the partitioning step.

3. Recursively Sort Sub-arrays

Recursively apply the QuickSort algorithm to the sub-arrays on the left and right of the pivot. Each recursive call involves selecting a pivot for the sub-array and partitioning it.

4. Combine Sub-arrays

As the recursion progresses, the sub-arrays become smaller, and eventually, the entire array becomes sorted in place.

The algorithm continues partitioning and sorting the sub-arrays until the entire array is sorted. The efficiency of QuickSort relies on selecting a good pivot and efficiently partitioning the array. Optimizations, such as using a randomized pivot selection or employing a three-way partitioning scheme, can further enhance its performance.

Summary

- QuickSort, developed by Tony Hoare in 1960, is a highly efficient divide-and-conquer sorting algorithm widely employed in practice.
- The algorithm selects a pivot element, partitions the array based on the pivot, and recursively sorts the sub-arrays until the entire array is sorted.
- Despite the potential worst-case time complexity of $O(n^2)$ when pivot selection is poorly optimized, strategies like random pivot selection or three-way partitioning

mitigate this issue.

- QuickSort's average-case time complexity of $O(n \log n)$, simplicity of implementation, and low overhead make it a preferred sorting method in various applications.
- Its versatility and speed have led to its widespread use as the default sorting method in programming languages and libraries.

3.5 Recursion

Recursion is a programming concept where a function calls itself to solve a problem. Instead of solving the entire problem at once, the function breaks it down into smaller, more manageable subproblems.

Each recursive call works on a smaller instance of the problem until a base case is reached, at which point the function stops calling itself and returns a result.

Recursion is particularly useful when dealing with problems that exhibit a self-repeating structure. It simplifies complex tasks by dividing them into smaller, identical tasks.

However, it's essential to define a base case to prevent infinite recursion and ensure the function eventually produces a result.

Recursion is widely used in algorithms like tree traversal, searching, and sorting, contributing to elegant and concise code solutions.

While powerful, excessive use of recursion can lead to performance issues and increased memory usage, so it's crucial to strike a balance between its benefits and potential drawbacks.

Summary

- Recursion in programming involves a function calling itself to solve a problem by breaking it into smaller subproblems.
- This process continues until a base case is met, at which point the function stops calling itself and returns a result.
- Recursion is effective for tasks with a self-repeating structure, allowing for elegant code solutions in algorithms like tree traversal and sorting.
- However, it's crucial to define a base case to prevent infinite recursion.
- While recursion simplifies complex tasks, excessive use can lead to performance issues and increased memory usage, emphasizing the need to strike a balance between its benefits and drawbacks in coding practices.

Example To Better Understand this:

Imagine you are at the entrance of a building with multiple floors, and each floor has several rooms. You want to count the total number of rooms in the building. Instead of

trying to count all the rooms on all floors at once, you decide to use a recursive approach.

You start by standing on the ground floor and counting the rooms on that floor. Then, you move to the next floor and repeat the process, counting the rooms on that floor. However, instead of counting all the rooms on the current floor immediately, you again use the same approach: stand on that floor and count the rooms on that floor.

This continues until you reach a floor where there are no more floors above it. At this point, you stop going up and start summing the room counts as you come back down from each floor. This process of breaking down the task into smaller subtasks and solving them one by one, reaching a base case (the top floor), and then combining the results is analogous to a recursive algorithm.

In this example:

- Counting the rooms on a floor is like solving a small part of the problem.
- Moving to the next floor is analogous to making a recursive call.
- Reaching the top floor with no more floors above is the base case.
- Summing up the room counts as you come back down is akin to resolving the recursive calls and obtaining the final result.

Recursive Definition and its Process

A recursive process is a way of solving a problem or performing a task in which the solution to a particular instance of the problem is expressed in terms of solutions to smaller instances of the same problem.

This process involves breaking down a complex problem into simpler, more manageable subproblems and using the same approach to solve each of them. The key characteristics of a recursive process include a base case and a set of rules or instructions that define how to solve larger instances of the problem in terms of smaller instances.

In a recursive process

1. Base Case: There is a stopping condition, known as the base case, which represents the simplest form of the problem and provides a direct solution without further recursion. It prevents the process from continuing indefinitely.

2. Divide and Conquer: The problem is divided into smaller instances or subproblems, often of the same nature as the original problem.

3. Recursive Calls: The solution to each subproblem is obtained by applying the same set of rules or instructions, which may involve making recursive calls to the same process.

4. Combining Results: As the recursive calls return results, they are combined to obtain the solution to the original problem.

An example of a recursive process is the computation of factorials. The factorial of a non-negative integer n (denoted as $n!$) is the product of all positive integers up to n . The recursive process involves breaking down the problem of computing $n!$ into the multiplication of n by $(n-1)!$ and so on, until reaching the base case where $0!$ is defined as 1.

Understanding and implementing recursive processes requires careful consideration of base cases to ensure termination and clear rules for breaking down and solving subproblems. Recursive processes are prevalent in programming and algorithmic design, providing an elegant and concise way to solve certain types of problems.

Summary

- A recursive process is a problem-solving approach where the solution to a specific problem is defined in terms of solving smaller instances of the same problem.
- This method entails breaking down a complex problem into more manageable subproblems, employing the same set of rules or instructions to solve each of them.
- Crucial to a recursive process is the establishment of a base case, serving as a termination condition to prevent indefinite recursion.
- The process involves dividing the problem into smaller instances, making recursive calls to solve each subproblem, and ultimately combining the results to obtain the solution to the original problem.
- An illustrative example is the computation of factorials, where the problem of calculating $n!$ is recursively broken down into the multiplication of n by $(n-1)!$ until reaching the base case of $0!$ equaling 1.
- Implementing recursive processes demands careful consideration of base cases for termination and clear rules for handling subproblems, making them a powerful and concise technique in programming and algorithmic design.

Write a program to do a recursive factorial of a non-negative integer:

```

#include <iostream>

int factorial(int n) {
    // Base case: factorial of 0 is 1
    if (n == 0) {
        return 1;
    } else {
        // Recursive case: n! = n * (n-1)!
        return n * factorial(n - 1);
    }
}

int main() {
    // Example usage
    int number;
    std::cout << "Enter a non-negative integer: ";
    std::cin >> number;

    if (number < 0) {
        std::cout << "Factorial is not defined for negative numbers.\n";
    } else {
        int result = factorial(number);
        std::cout << "The factorial of " << number << " is: " << result << "\n";
    }

    return 0;
}

```

This program defines a recursive function `factorial` that calculates the factorial of a given non-negative integer. The base case is when `n` is 0, and in this case, the function returns 1. Otherwise, it makes a recursive call to calculate `n * (n-1)!`.

The `main` function takes user input for a non-negative integer, calls the `factorial` function, and prints the result.

Write recursive program for fibonacci in c++.

```

#include <iostream>

int fibonacci(int n) {
    // Base case: Fibonacci of 0 and 1 is 1
    if (n == 0 || n == 1) {
        return 1;
    } else {
        // Recursive case: Fibonacci(n) = Fibonacci(n-1) + Fibonacci(n-2)
        return fibonacci(n - 1) + fibonacci(n - 2);
    }
}

int main() {
    // Example usage
    int term;
    std::cout << "Enter the term of the Fibonacci sequence: ";
    std::cin >> term;

    if (term < 0) {
        std::cout << "Fibonacci sequence is not defined for negative terms.\n";
    } else {
        int result = fibonacci(term);
        std::cout << "The Fibonacci term at position " << term << " is: " << result << "\n";
    }

    return 0;
}

```

This program defines a recursive function `Fibonacci` that calculates the Fibonacci sequence for a given term. The base case is when `n` is 0 or 1, in which case the function returns 1.

Otherwise, it makes recursive calls to calculate $\text{Fibonacci}(n-1) + \text{Fibonacci}(n-2)$. The `main` function takes user input for the term, calls the `Fibonacci` function, and prints the result.

Keep in mind that the recursive approach for Fibonacci can be inefficient for large terms due to redundant calculations, and more efficient approaches like dynamic programming or memoization are often preferred for larger sequences.

3.6 Tower Of Hanoi

The Tower of Hanoi is a classic puzzle that involves moving a stack of differently-sized disks from one rod to another, following three simple rules. Imagine three rods and several disks of different sizes, initially stacked in decreasing size order on one rod.

The objective is to move the entire stack to another rod, with the help of an intermediate rod, ensuring that smaller disks are always on top of larger ones.

To solve the Tower of Hanoi puzzle, you follow these rules:

- Only one disk can be moved at a time.
- Each move consists of taking the upper disk from one of the stacks and placing it on top of the stack on another rod or an empty rod.
- No disk may be placed on top of a smaller disk. The challenge lies in finding the optimal sequence of moves to transfer the entire stack from one rod to another.
- The Tower of Hanoi is not only an engaging puzzle but also a classic example of a problem that can be elegantly solved using recursion.

Breaking down the task of moving a stack of disks into smaller subproblems and applying the same logic recursively eventually leads to a solution, provided you follow the rules.

Implementation of Tower of Hanoi using Recursive :

```

#include <iostream>

void towerOfHanoi(int n, char source, char auxiliary, char destination) {
    // Base case: If there's only one disk, move it from source to destination
    if (n == 1) {
        std::cout << "Move disk 1 from " << source << " to " << destination << std::endl;
        return;
    }

    // Move n-1 disks from source to auxiliary peg using destination as auxiliary
    towerOfHanoi(n - 1, source, destination, auxiliary);

    // Move the nth disk from source to destination
    std::cout << "Move disk " << n << " from " << source << " to " << destination << std::endl;

    // Move the n-1 disks from auxiliary to destination peg using source as auxiliary
    towerOfHanoi(n - 1, auxiliary, source, destination);
}

int main() {
    int numberOfDisks;

    std::cout << "Enter the number of disks: ";
    std::cin >> numberOfDisks;

    // Function call to solve Tower of Hanoi
    towerOfHanoi(numberOfDisks, 'A', 'B', 'C');

    return 0;
}

```

In this program, the Tower of Hanoi function is defined to solve the Tower of Hanoi problem recursively. It takes the number of disks (n) and the three pegs (source, auxiliary, and destination) as parameters.

The base case is when there is only one disk, in which case it directly moves the disk from the source peg to the destination peg.

Otherwise, it recursively moves $n-1$ disks from the source to the auxiliary peg, then the n th disk from the source to the destination peg, and finally, the $n-1$ disks from the auxiliary to the destination peg.

The main function takes user input for the number of disks and calls the `towerOfHanoi` function to display the sequence of moves.

Tower of Hanoi procedure using stack

To implement the Tower of Hanoi procedure using a stack, you can use an iterative approach that emulates the recursive process by maintaining a stack of subproblems.

The stack stores information about each subproblem, such as the number of disks, source, auxiliary, and destination pegs.

The algorithm starts with the initial subproblem and iteratively performs the following steps: if the current subproblem has one disk (base case), move the disk from the source peg to the destination peg.

Otherwise, push the subproblems representing the next recursive steps onto the stack. These steps are repeated until the stack is empty.

By leveraging a stack to manage the state of subproblems, this iterative procedure ensures that each move is executed in the correct order, simulating the recursive process without relying on actual recursive calls.

This approach is particularly beneficial for scenarios where deep recursion may lead to stack overflow, and an iterative solution with a stack can provide a more efficient and controlled use of memory.

Example:

```
#include <iostream>
#include <stack>

struct HanoiMove {
    int n;           // Number of disks
    char source, aux, destination; // Pegs: source, auxiliary, destination
    bool toMove;     // Indicates whether the move is to be executed
};

void towerOfHanoiUsingStack(int numberOfDisks, char source, char auxiliary, char destination) {
    std::stack<HanoiMove> movesStack;
    movesStack.push({numberOfDisks, source, auxiliary, destination, true});

    while (!movesStack.empty()) {
        HanoiMove currentMove = movesStack.top();
        movesStack.pop();

        if (currentMove.n == 1) {
            // Base case: move a disk directly
            std::cout << "Move disk 1 from " << currentMove.source << " to " << currentMove.destination << std::endl;
        } else {
            if (currentMove.toMove) {
                // Push next recursive steps onto the stack
                movesStack.push({currentMove.n - 1, currentMove.aux, currentMove.source, currentMove.destination, true});
                movesStack.push({1, currentMove.source, currentMove.aux, currentMove.destination, false});
                movesStack.push({currentMove.n - 1, currentMove.source, currentMove.destination, currentMove.aux, true});
            } else {
                // Move the nth disk
                std::cout << "Move disk " << currentMove.n << " from " << currentMove.source << " to " << currentMove.destination << std::endl;
            }
        }
    }
}

int main() {
    int numberOfDisks;

    std::cout << "Enter the number of disks: ";
    std::cin >> numberOfDisks;

    // Function call to solve Tower of Hanoi using a stack
    towerOfHanoiUsingStack(numberOfDisks, 'A', 'B', 'C');

    return 0;
}
```

This C++ program uses a stack of HanoiMove structures to simulate the iterative Tower of Hanoi process. The program iteratively pops moves from the stack and executes them until the stack is empty.

The HanoiMove structure represents a subproblem, storing information about the number of disks and peg configurations. The base case involves moving a single disk directly, while for larger subproblems, the next recursive steps are pushed onto the stack.

This iterative implementation effectively simulates the recursive logic of the Tower of Hanoi problem.

Summary

- To implement the Tower of Hanoi procedure using a stack, an iterative approach is employed to mimic the recursive process.
- The algorithm maintains a stack to store information about each subproblem, including the number of disks and peg configurations.
- It initiates with the base case of one disk, moving it directly from source to destination. For subproblems with more than one disk, the algorithm pushes the next recursive steps onto the stack.
- This process continues iteratively until the stack is empty. The stack efficiently manages the state of subproblems, allowing for a controlled and memory-efficient execution, especially in scenarios where deep recursion might lead to stack overflow.
- This approach effectively simulates the recursive logic without relying on actual recursive calls.

3.7 Queues

Queues are a fundamental data structure in computer science that operates on the First-In-First-Out (FIFO) principle. In a queue, elements are added to the rear, or "enqueue" operation, and removed from the front, known as "dequeue."

This ensures that the first element added is the first to be processed. Queues are commonly used for managing tasks in a sequential manner, such as handling tasks in a print queue or processing events in a computer system.

The simplicity and efficiency of queues make them essential for various applications, including task scheduling, breadth-first search algorithms, and managing resources in a systematic order.

Additionally, queues are often implemented in programming languages using arrays or linked lists, offering flexibility in handling diverse scenarios.

The organized and predictable nature of queues contributes to their widespread use in designing algorithms and solving problems that require the orderly processing of elements.

Sequential Representation of Queues

diagram

Queues operate based on the First-In-First-Out (FIFO) principle, and there are two fundamental operations: enqueue and dequeue.

1. Enqueue

- In the enqueue operation, an element is added to the rear or the end of the queue.
- This is analogous to someone joining the end of a line or queue in a real-world scenario.

- The enqueue operation ensures that new elements are inserted after existing ones, maintaining the order in which they were added.

2. Dequeue

- In the dequeue operation, the element at the front or the beginning of the queue is removed.
- This is similar to serving the person at the front of a line in a real-world queue.
- The dequeue operation ensures that elements are processed in the order they were added, adhering to the FIFO principle.

These two basic operations define the primary functionality of a queue. Additionally, there are variations and supporting operations such as peek (to view the front element without removing it), checking if the queue is empty, and determining its size.

Queues are commonly used in scenarios where the order of processing elements is crucial, such as in task scheduling, managing resources, and breadth-first search algorithms. Implementations of queues can be achieved using various data structures, such as arrays or linked lists, depending on the specific requirements of the application.

Summary

- Queues adhere to the First-In-First-Out (FIFO) principle, governed by two essential operations: enqueue and dequeue.
- Enqueue involves adding an element to the rear, akin to joining the end of a real-world line, ensuring the preservation of the order of insertion.
- Dequeue, on the other hand, removes the front element, mimicking serving the person at the front of a physical queue.
- These operations, fundamental to a queue's functionality, are vital in scenarios where processing elements in a specific order is critical, such as task scheduling and breadth-first search algorithms.
- The versatility of queues allows for implementations using diverse data structures like arrays or linked lists, tailoring them to meet specific application requirements.

Concept of Queues

Queues are a fundamental data structure that follows the First-In-First-Out (FIFO) principle, meaning the first element added is the first to be removed. They provide an organized way to manage elements, allowing efficient processing sequentially.

Enqueue, the operation to add elements to the rear, and dequeue, the operation to remove elements from the front, are the key operations defining the behavior of queues.

This structure is widely employed in computer science for tasks like task scheduling, managing resources, and ensuring the systematic processing of elements in algorithms like breadth-first search.

The simplicity and predictability of queues make them essential in designing algorithms and solving problems requiring orderly and sequential processing of elements.

Circular Queues

A circular queue is a variation of a regular queue that forms a closed loop, allowing efficient utilization of space and avoiding wastage. In a circular queue, the last element is connected to the first element, creating a circular arrangement.

This design facilitates a seamless wraparound effect, enabling the queue to reuse space that becomes available after dequeuing elements.

Circular queues are particularly useful in scenarios where the underlying data structure is implemented as an array, as the circular arrangement helps in the efficient use of memory and avoids the need for periodic reallocation.

The circular queue maintains the FIFO principle, where elements are added at the rear and removed from the front. The enqueue and dequeue operations can be performed efficiently by adjusting the front and rear pointers while considering the circular nature of the structure.

Circular queues find applications in scenarios like task scheduling, where elements are continuously added and removed, and the circular design helps in optimizing the use of limited space.

Summary

- A circular queue is a modified version of a standard queue, forming a closed loop that efficiently utilizes space and eliminates wastage.
- By connecting the last element to the first in a circular arrangement, this design enables seamless space reuse after dequeuing elements.
- Particularly advantageous when implemented as an array, the circular structure optimizes memory usage, eliminating the need for frequent reallocation.
- Following the First-In-First-Out (FIFO) principle, elements are added at the rear and removed from the front, with efficient enqueue and dequeue operations achieved by adjusting pointers.
- Circular queues prove beneficial in scenarios like task scheduling, where continuous addition and removal of elements occur, making them instrumental in optimizing the utilization of limited space.

Priority Queues

A priority queue is a data structure that stores elements with associated priorities and allows retrieval of the element with the highest (or lowest) priority.

Unlike a regular queue, where elements are processed in a first-in-first-out (FIFO) manner, priority queues prioritize elements based on a defined ordering criterion.

Elements with higher priority are dequeued before those with lower priority.

Priority queues are commonly used in various applications, such as task scheduling, where tasks with higher priority need to be processed before lower-priority tasks. They provide efficient access to the element with the highest priority, often implemented using binary heaps, Fibonacci heaps, or other specialized data structures.

Priority queues are versatile and find applications in algorithms like Dijkstra's shortest path algorithm, Huffman coding, and scheduling processes in operating systems.

The key operations in a priority queue include insertion (enqueueing) and extraction of the highest-priority element (dequeueing), ensuring that elements are processed based on their importance or urgency.

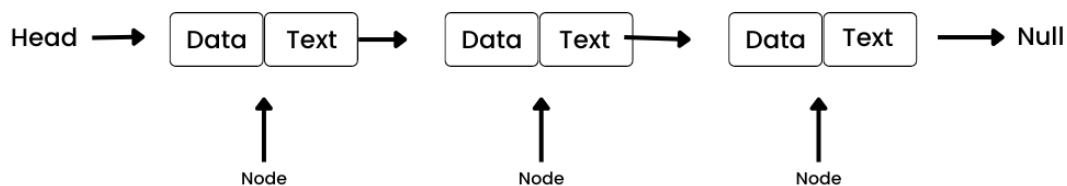
Summary

- A priority queue is a specialized data structure that organizes elements according to their priorities, enabling efficient retrieval of the highest or lowest priority element.
 - In contrast to regular queues, which follow a first-in-first-out (FIFO) order, priority queues prioritize elements based on specified criteria.
 - This structure is crucial in applications like task scheduling, where higher-priority tasks are processed before lower-priority ones.
 - Priority queues optimize access to the highest-priority element and are implemented using binary heaps, Fibonacci heaps, or similar data structures.
 - Widely utilized in algorithms such as Dijkstra's shortest path and Huffman coding, priority queues facilitate streamlined processing based on importance or urgency, with key operations including insertion and extraction of the highest-priority element.
-

CHAPTER 4: LINKED LIST



4.1 Linked Lists



A linked list is a data structure used in computer science to organize and store data. Unlike arrays, linked lists do not have a fixed size, and elements are not stored in contiguous memory locations.

Instead, each element called a node, consists of a data field and a reference or link to the next node in the sequence. This linking mechanism allows for the efficient insertion and deletion of elements at any position in the list, as it only requires updating the pointers in the adjacent nodes.

Linked lists come in various forms, such as singly linked lists, where nodes point only to the next node, and doubly linked lists, where nodes have references to both the next and previous nodes.

The flexibility of linked lists makes them suitable for dynamic data structures where the size may change frequently, and they are commonly used in scenarios like implementing stacks, queues, and other dynamic data storage requirements in computer programs.

However, the trade-off is that accessing elements in a linked list is less efficient than in arrays since it involves traversing the list sequentially from the beginning.

Summary

- A linked list is a dynamic data structure in computer science where elements, or nodes, contain both data and a reference to the next node.

- Unlike arrays, linked lists allow for efficient insertion and deletion at any point but may have slower access times due to sequential traversal.
- They come in various forms, such as singly and doubly linked lists, and are commonly used for dynamic data storage in scenarios like implementing stacks and queues.

4.2 Introduction to singly Link List

A singly linked list is a fundamental data structure in computer science that organizes elements sequentially, called nodes.

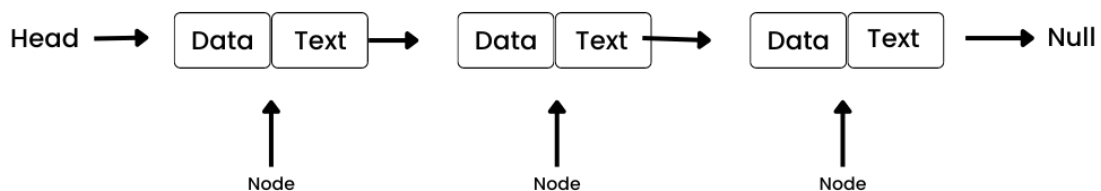
Each node in the list consists of a data field and a reference, or link, pointing to the next node in the sequence.

Unlike arrays, singly linked lists dynamically adjust in size, enabling efficient insertion and deletion of elements at any position.

However, access to specific elements in a singly linked list requires traversing the list from the beginning, which can be less efficient than direct access in arrays.

Singly-linked lists find widespread use in scenarios where dynamic data storage and frequent insertions or deletions are common, such as implementing stacks, queues, and various linked data structures in computer programs.

4.3 Representation of a singly Link List



4.4 Creating a sorted Link List

Creating a sorted linked list involves inserting elements in such a way that the list maintains ascending or descending order after each insertion. Here's an example of how you can create a sorted linked list in C++.

Part-1

```

#include <iostream>

struct Node {
    int data;
    Node* next;
};

// Function to create a new node
Node* createNode(int value) {
    Node* newNode = new Node;
    newNode->data = value;
    newNode->next = nullptr;
    return newNode;
}

// Function to insert a node into a sorted linked list
void insertSorted(Node*& head, int value) {
    Node* newNode = createNode(value);

    if (head == nullptr || value < head->data) {
        // Insert at the beginning
        newNode->next = head;
        head = newNode;
    } else {
        // Traverse the list to find the appropriate position
        Node* current = head;
        while (current->next != nullptr && current->next->data < value) {
            current = current->next;
        }

        // Insert the new node
        newNode->next = current->next;
        current->next = newNode;
    }
}

```

Part-2:

```

// Function to traverse and print a linked list
void traverseLinkedList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
}

int main() {
    // Create a sorted linked list
    Node* head = nullptr;
    insertSorted(head, 3);
    insertSorted(head, 1);
    insertSorted(head, 5);
    insertSorted(head, 2);
    insertSorted(head, 4);

    // Traverse and print the sorted linked list
    std::cout << "Sorted Linked List: ";
    traverseLinkedList(head);

    // Don't forget to free allocated memory when done
    while (head != nullptr) {
        Node* temp = head;
        head = head->next;
        delete temp;
    }

    return 0;
}

```

In this example, the `insertSorted` function inserts a new node into the linked list while maintaining the sorted order. The `traverseLinkedList` function is used to print the elements of the sorted linked list. Ensure to free the allocated memory for the nodes when you're done with the linked list to prevent memory leaks.

4.5 Creating an Unsorted Link List

Creating an unsorted linked list involves inserting elements without any specific order. Here's an example of how you can create an unsorted linked list in C++.

```
#include <iostream>

struct Node {
    int data;
    Node* next;
};

// Function to create a new node
Node* createNode(int value) {
    Node* newNode = new Node;
    newNode->data = value;
    newNode->next = nullptr;
    return newNode;
}

// Function to insert a node at the beginning of the linked list
void insertAtBeginning(Node*& head, int value) {
    Node* newNode = createNode(value);
    newNode->next = head;
    head = newNode;
}

// Function to traverse and print a linked list
void traverseLinkedList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
}

int main() {
    // Create an unsorted linked list
    Node* head = nullptr;
    insertAtBeginning(head, 3);
    insertAtBeginning(head, 1);
    insertAtBeginning(head, 5);
    insertAtBeginning(head, 2);
    insertAtBeginning(head, 4);

    // Traverse and print the unsorted linked list
    std::cout << "Unsorted Linked List: ";
    traverseLinkedList(head);

    // Don't forget to free allocated memory when done
    while (head != nullptr) {
        Node* temp = head;
        head = head->next;
        delete temp;
    }

    return 0;
}
```

In this example, the `insertAtBeginning` function inserts a new node at the beginning of the linked list, which is a common approach for simplicity. The `traverseLinkedList` function is used to print the elements of the unsorted linked list.

Ensure to free the allocated memory for the nodes when you're done with the linked list to prevent memory leaks. Keep in mind that the order in which elements are inserted is not preserved, making this an unsorted linked list.

4.6 Traversing a sorted Link List

Traversing a sorted linked list involves visiting each node in the list in a specific order, typically from the beginning to the end, taking advantage of the sorted order to efficiently locate elements. Here's an example of how you might traverse a sorted linked list in C++.

```

#include <iostream>

struct Node {
    int data;
    Node* next;
};

// Function to traverse and print a linked list
void traverseSortedLinkedList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
}

int main() {
    // Example of a sorted linked list
    Node* head = new Node{1, nullptr};
    head->next = new Node{3, nullptr};
    head->next->next = new Node{5, nullptr};
    head->next->next->next = new Node{7, nullptr};
    head->next->next->next->next = new Node{9, nullptr};

    // Traverse and print the sorted linked list
    std::cout << "Sorted Linked List: ";
    traverseSortedLinkedList(head);

    // Don't forget to free allocated memory when done
    while (head != nullptr) {
        Node* temp = head;
        head = head->next;
        delete temp;
    }

    return 0;
}

```

In this example, the `traverseSortedLinkedList` function iterates through the linked list while printing the data of each node.

The main function creates a simple sorted linked list for illustration, but in a real-world scenario, you would likely build the list dynamically based on your requirements.

It's essential to free the allocated memory for the nodes when you are done with the linked list to avoid memory leaks.

4.7 Traversing an Unsorted Link List

Traversing an unsorted linked list is similar to traversing a sorted one. The primary difference is that you don't have the guarantee of ascending or descending order, so you simply visit each node in the order they are linked. Here's an example in C++:

```

#include <iostream>

struct Node {
    int data;
    Node* next;
};

// Function to traverse and print an unsorted linked list
void traverseUnsortedLinkedList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
}

int main() {
    // Example of an unsorted linked list
    Node* head = new Node{3, nullptr};
    head->next = new Node{1, nullptr};
    head->next->next = new Node{7, nullptr};
    head->next->next->next = new Node{5, nullptr};
    head->next->next->next->next = new Node{2, nullptr};

    // Traverse and print the unsorted linked list
    std::cout << "Unsorted Linked List: ";
    traverseUnsortedLinkedList(head);

    // Don't forget to free allocated memory when done
    while (head != nullptr) {
        Node* temp = head;
        head = head->next;
        delete temp;
    }

    return 0;
}

```

In this example, the `traverseUnsortedLinkedList` function iterates through the linked list while printing the data of each node.

The main function creates a simple unsorted linked list for illustration, but in real-world scenarios, the list might be built dynamically based on your requirements.

Remember to free the allocated memory for the nodes when you're done with the linked list to avoid memory leaks.

4.8 Searching in sorted Link List

Searching in a sorted linked list involves efficiently locating a specific value within the list by taking advantage of its sorted order. Here's an example of how you can perform a search in a sorted linked list in C++:

Part-1:

```

#include <iostream>

struct Node {
    int data;
    Node* next;
};

// Function to search for a value in a sorted linked list
bool searchInSorted(Node* head, int target) {
    while (head != nullptr && head->data <= target) {
        if (head->data == target) {
            return true; // Found the target value
        }
        head = head->next;
    }
    return false; // Target value not found
}

// Function to insert a node into a sorted linked list
void insertSorted(Node*& head, int value) {
    Node* newNode = new Node{value, nullptr};

    if (head == nullptr || value < head->data) {
        // Insert at the beginning
        newNode->next = head;
        head = newNode;
    } else {
        // Traverse the list to find the appropriate position
        Node* current = head;
        while (current->next != nullptr && current->next->data < value) {
            current = current->next;
        }

        // Insert the new node
        newNode->next = current->next;
        current->next = newNode;
    }
}

```

Part-2:

```

// Function to traverse and print a linked list
void traverseLinkedList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
}

int main() {
    // Create a sorted linked list
    Node* head = nullptr;
    insertSorted(head, 3);
    insertSorted(head, 1);
    insertSorted(head, 5);
    insertSorted(head, 2);
    insertSorted(head, 4);

    // Traverse and print the sorted linked list
    std::cout << "Sorted Linked List: ";
    traverseLinkedList(head);

    // Search for a value in the sorted linked list
    int targetValue = 2;
    if (searchInSorted(head, targetValue)) {
        std::cout << "\nValue " << targetValue << " found in the list.\n";
    } else {
        std::cout << "\nValue " << targetValue << " not found in the list.\n";
    }

    // Don't forget to free allocated memory when done
    while (head != nullptr) {
        Node* temp = head;
        head = head->next;
        delete temp;
    }

    return 0;
}

```

In this example, the `searchInSorted` function efficiently searches for a target value in the sorted linked list, taking advantage of the sorted order.

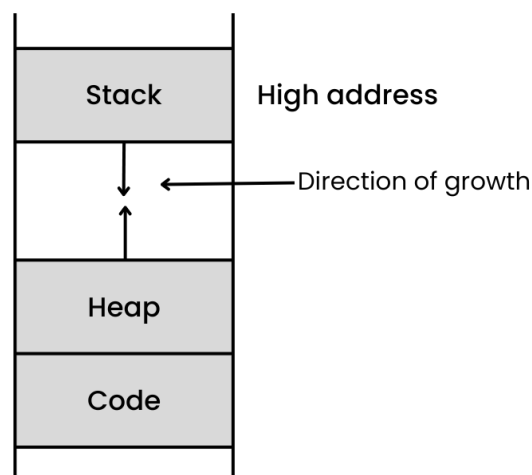
The main function also demonstrates inserting elements into the sorted linked list and then searching for a specific value. Finally, it frees the allocated memory for the nodes to avoid memory leaks.

4.9 Memory Allocation

Memory allocation refers to the process of setting aside a portion of a computer's memory for a program to use. In modern computing systems, memory is a critical resource that is essential for storing and accessing data during the execution of programs.

Memory allocation is a fundamental aspect of a program's runtime behavior, and it involves two main tasks: allocating memory and deallocating memory.

1. Allocation of Memory



Static Allocation: In some cases, memory is allocated at compile-time and remains fixed throughout the program's execution. This is known as static allocation. For example, global variables and constants may be statically allocated.

Dynamic Allocation: In many scenarios, especially when the size of data is not known until runtime, dynamic memory allocation is used. This involves requesting memory from the system during the program's execution. Common programming languages, like C and C++, provide functions like `malloc`, `calloc`, and `new` for dynamic memory allocation.

Stack Allocation: Local variables and function call information are often stored in a region of memory known as the stack. Memory for these variables is automatically allocated and deallocated as functions are called and returned. This is known as stack allocation.

Heap Allocation: The heap is another region of memory that is used for dynamic memory allocation. Memory allocated in the heap needs to be explicitly managed by the programmer. Improper management can lead to memory leaks or fragmentation.

2. Deallocation of Memory

Static Deallocation: Memory allocated statically is typically deallocated automatically when the program exits. The operating system releases the memory back to the system.

Dynamic Deallocation: For dynamically allocated memory, it is the responsibility of the programmer to release the memory when it is no longer needed. In languages like C and C++, functions like `free` or `delete` are used to deallocate memory.

Automatic Deallocation (Garbage Collection): Some programming languages, like Java, Python, and C#, use automatic memory management systems such as garbage collectors. These systems automatically identify and reclaim memory that is no longer in use, reducing the burden on the programmer to manage memory manually.

Memory allocation is crucial for efficient program execution and is a key consideration in writing robust and high-performance software. Improper memory management can lead to memory leaks, where allocated memory is not deallocated, or to fragmentation issues, which can reduce the available memory for future allocations.

Summary

- Memory allocation is a critical aspect of a computer program's runtime behavior, involving the allocation and deallocation of memory.
- In static allocation, memory is assigned at compile-time and remains fixed throughout the program's execution, often used for global variables and constants.
- Dynamic allocation, prevalent in languages like C and C++, involves requesting memory during runtime using functions like `malloc` and `new`. Stack allocation, where local variables are stored in a stack region, is automatically managed as functions are called and returned.
- Heap allocation, in contrast, requires manual management and is prone to issues like memory leaks or fragmentation if not handled properly.
- Deallocation methods, such as `free` or `delete` in C and C++, are necessary for releasing dynamically allocated memory.
- Some languages, like Java and Python, employ automatic memory management, known as garbage collection, to identify and reclaim unused memory automatically.
- Efficient memory management is crucial for program efficiency, and improper handling can lead to issues like memory leaks and fragmentation, impacting software performance and reliability.

4.10 Insertion in Linked List

Insertion in a linked list involves adding a new node to the existing sequence of nodes. The process varies depending on the insertion point: at the beginning, middle, or end of the list. When inserting at the beginning, a new node is created, its data is set, and it is linked to the current head of the list, followed by updating the head pointer.

For insertion in the middle, traversal to the desired position is necessary, and pointers in the surrounding nodes are adjusted to accommodate the new node. Inserting at the end requires traversing to the last node, creating the new node, linking it to the current tail, and updating the tail pointer.

Linked lists offer flexibility in memory management, making insertions efficient, especially in scenarios involving frequent modifications to the data structure.

Steps to do inserting in a linked list:

1. Insertion at the Beginning

To insert a new node at the beginning of a linked list, a new node is created, its data is set, and its next pointer is linked to the current head of the list.

Finally, the head pointer is updated to point to the new node.

Example:

```
#include <iostream>

// Define a basic Node structure
struct Node {
    int data;
    Node* next;

    Node(int value) : data(value), next(nullptr) {}
};

// Function to insert a new node at the beginning of the linked list
void insertAtBeginning(Node*& head, int value) {
    Node* newNode = new Node(value);
    newNode->next = head;
    head = newNode;
}

// Function to display the linked list
void displayList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
    std::cout << std::endl;
}

int main() {
    // Create an initial linked list with three nodes: 3 -> 7 -> 11
    Node* head = new Node(3);
    head->next = new Node(7);
    head->next->next = new Node(11);

    // Display the original linked list
    std::cout << "Original Linked List: ";
    displayList(head);

    // Insert a new node with value 5 at the beginning
    insertAtBeginning(head, 5);

    // Display the linked list after insertion
    std::cout << "Linked List after Insertion: ";
    displayList(head);

    return 0;
}
```

2. Insertion in the Middle

When inserting in the middle, the process involves traversing the list to the desired position, creating a new node, setting its data, and adjusting pointers in the preceding and succeeding nodes to link with the new node.

Example:

```
#include <iostream>

// Define a basic Node structure
struct Node {
    int data;
    Node* next;

    Node(int value) : data(value), next(nullptr) {}
};

// Function to insert a new node in the middle of the linked list
void insertInMiddle(Node*& head, int value, int position) {
    Node* newNode = new Node(value);

    // Traverse to the desired position
    Node* current = head;
    for (int i = 1; i < position && current != nullptr; ++i) {
        current = current->next;
    }

    // Check if the position is valid
    if (current == nullptr) {
        std::cerr << "Invalid position for insertion." << std::endl;
        return;
    }

    // Adjust pointers to insert the new node
    newNode->next = current->next;
    current->next = newNode;
}

// Function to display the linked list
void displayList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
    std::cout << std::endl;
}

int main() {
    // Create an initial linked list with nodes: 2 -> 4 -> 6
    Node* head = new Node(2);
    head->next = new Node(4);
    head->next->next = new Node(6);

    // Display the original linked list
    std::cout << "Original Linked List: ";
    displayList(head);

    // Insert a new node with value 5 at position 2 (middle)
    insertInMiddle(head, 5, 2);

    // Display the linked list after insertion
    std::cout << "Linked List after Insertion in the Middle: ";
    displayList(head);

    return 0;
}
```

3. Insertion at the End

To insert a node at the end, the list is traversed until the last node is reached. A new node is then created, its data is set, and it is linked to the current tail node.

The tail pointer is updated to point to the new node.

The key advantage of linked lists is their dynamic memory allocation, allowing for efficient insertions and deletions compared to arrays. However, it's essential to manage pointers correctly to ensure the integrity of the linked list structure.

Insertion in linked lists can be more flexible and faster than array-based structures, especially in scenarios where frequent modifications or dynamic memory allocation are required.

Example:

```
#include <iostream>

// Define a basic Node structure
struct Node {
    int data;
    Node* next;

    Node(int value) : data(value), next(nullptr) {}
};

// Function to insert a new node at the end of the linked list
void insertAtEnd(Node*& head, int value) {
    Node* newNode = new Node(value);

    // If the list is empty, the new node becomes the head
    if (head == nullptr) {
        head = newNode;
        return;
    }

    // Traverse to the last node
    Node* current = head;
    while (current->next != nullptr) {
        current = current->next;
    }

    // Link the new node to the current tail and update the tail pointer
    current->next = newNode;
}

// Function to display the linked list
void displayList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
    std::cout << std::endl;
}

int main() {
    // Create an initial linked list with nodes: 2 -> 4 -> 6
    Node* head = new Node(2);
    head->next = new Node(4);
    head->next->next = new Node(6);

    // Display the original linked list
    std::cout << "Original Linked List: ";
    displayList(head);

    // Insert a new node with value 8 at the end
    insertAtEnd(head, 8);

    // Display the linked list after insertion
    std::cout << "Linked List after Insertion at the End: ";
    displayList(head);

    return 0;
}
```

Summary

- Insertion in a linked list is a versatile process involving the addition of a new node to the existing sequence, and the method employed depends on whether it occurs at the beginning, middle, or end of the list.
- For insertion at the beginning, a new node is created, its data is assigned, and it is linked to the current head, followed by the update of the head pointer.
- In the middle, traversal to the desired position is necessary, with adjustments made to pointers in the surrounding nodes to accommodate the new node seamlessly.
- To insert at the end, traversal reaches the last node, a new node is generated, linked to the current tail, and the tail pointer is updated accordingly.
- Linked lists, offering flexibility in memory management, demonstrate efficiency in insertions, especially in scenarios with frequent structural modifications.

- The dynamic allocation of memory in linked lists allows for swift and efficient insertions and deletions compared to arrays, although it is crucial to manage pointers judiciously to maintain the integrity of the linked list structure.
- Overall, insertion in linked lists proves to be more flexible and faster than array-based structures, particularly in situations requiring frequent modifications or dynamic memory allocation.

4.11 Deletion from a Linked List

Deletion from a linked list involves removing a node from the existing sequence, and the process depends on the target node's position within the list.

Three common scenarios include deleting a node from the beginning, middle, or end of the linked list. When deleting from the beginning, the head pointer is updated to point to the next node, and the memory of the removed node is deallocated.

For deletion in the middle, traversal to the desired position is necessary, and pointers in the preceding node are adjusted to bypass the node to be deleted, followed by deallocating its memory.

Deleting from the end requires traversing to the last node while keeping track of the preceding node, updating its next pointer to null, and then deallocating the memory of the removed node.

Deletion of a linked list is essential for maintaining the integrity of the data structure. Efficient memory deallocation ensures that resources are properly managed, and careful adjustment of pointers guarantees the correct linkage within the list.

The flexibility of linked lists in handling deletions, especially in scenarios involving dynamic data and frequent modifications, is a key advantage over array-based structures. Proper implementation of deletion operations is crucial for the reliable and efficient performance of linked lists in various applications.

Summary

- Deletion from a linked list involves removing a node, with procedures differing for the beginning, middle, or end positions.
- In the case of deletion from the beginning, the head pointer is updated and the node's memory is deallocated.
- For middle deletion, traversal to the target position adjusts preceding pointers before deallocating the node.
- Deleting from the end involves traversing to the last node, updating the preceding node's pointer, and deallocating the memory.
- Efficient memory management and pointer adjustments are crucial for maintaining data structure integrity.
- Linked lists excel in deletion flexibility, especially in dynamic scenarios, offering a key advantage over array-based structures.

- Accurate implementation of deletion operations is vital for ensuring the reliability and efficiency of linked lists in diverse applications.

1. Delete from the Beginning

Deleting a node from the beginning of a linked list involves several steps to ensure proper memory management and maintenance of the list structure.

Here are the steps:

1. Check for Empty List

Before attempting to delete a node from the beginning, check if the linked list is empty. If the list is empty (head is `nullptr`), then there is nothing to delete.

2. Update Head Pointer

If the list is not empty, the first node (head) is the one to be deleted. Create a temporary pointer to the node to be deleted to keep track of it. Update the head pointer to point to the next node in the list. This effectively removes the first node from the list.

3. Deallocate Memory

After updating the head pointer, deallocate the memory occupied by the node that was removed. This step is crucial to prevent memory leaks.

Example:

```

#include <iostream>

struct Node {
    int data;
    Node* next;

    Node(int value) : data(value), next(nullptr) {}
};

void deleteFromBeginning(Node*& head) {
    // Step 1: Check for an empty list
    if (head == nullptr) {
        std::cout << "List is empty. Nothing to delete." << std::endl;
        return;
    }

    // Step 2: Update Head Pointer
    Node* temp = head;
    head = head->next;

    // Step 3: Deallocate Memory
    delete temp;
}

// Function to display the linked list
void displayList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
    std::cout << std::endl;
}

int main() {
    // Create an initial linked list with nodes: 2 -> 4 -> 6
    Node* head = new Node(2);
    head->next = new Node(4);
    head->next->next = new Node(6);

    // Display the original linked list
    std::cout << "Original Linked List: ";
    displayList(head);

    // Delete from the beginning
    deleteFromBeginning(head);

    // Display the linked list after deletion
    std::cout << "Linked List after Deletion from Beginning: ";
    displayList(head);

    return 0;
}

```

In this example, the `deleteFromBeginning` function is used to delete a node from the beginning of the linked list, and the `displayList` function is used to display the original and modified linked lists.

2. Delete from the End

Deleting a node from the end of a linked list involves a slightly more intricate process compared to deletion from the beginning.

Here are the steps:

1. Check for Empty List

Before attempting to delete a node from the end, check if the linked list is empty. If the list is empty (head is `nullptr`), then there is nothing to delete.

2. Traverse to the Last Node and Keep Track of the Preceding Node

Traverse the list to reach the last node while keeping track of the preceding node. This is essential for updating the next pointer of the preceding node in the next step.

3. Update the Next Pointer of the Preceding Node

Once the last node is reached, update the next pointer of the preceding node to `nullptr`. This effectively removes the last node from the list.

4. Deallocate Memory

After updating the pointers, deallocate the memory occupied by the node that was removed. This step is crucial to prevent memory leaks.

Example:

```
#include <iostream>

struct Node {
    int data;
    Node* next;

    Node(int value) : data(value), next(nullptr) {}
};

void deleteFromEnd(Node*& head) {
    // Step 1: Check for an empty list
    if (head == nullptr) {
        std::cout << "List is empty. Nothing to delete." << std::endl;
        return;
    }

    // Step 2: Traverse to the Last Node and Keep Track of the Preceding Node
    Node* current = head;
    Node* previous = nullptr;

    while (current->next != nullptr) {
        previous = current;
        current = current->next;
    }

    // Step 3: Update Next Pointer of the Preceding Node
    if (previous != nullptr) {
        previous->next = nullptr;
    } else {
        // If there is only one node in the list
        head = nullptr;
    }

    // Step 4: Deallocate Memory
    delete current;
}

// Function to display the linked list
void displayList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
    std::cout << std::endl;
}

int main() {
    // Create an initial linked list with nodes: 2 -> 4 -> 6
    Node* head = new Node(2);
    head->next = new Node(4);
    head->next->next = new Node(6);

    // Display the original linked list
    std::cout << "Original Linked List: ";
    displayList(head);

    // Delete from the end
    deleteFromEnd(head);

    // Display the linked list after deletion
    std::cout << "Linked List after Deletion from End: ";
    displayList(head);

    return 0;
}
```

In this example, the `deleteFromEnd` function is used to delete a node from the end of the linked list, and the `displayList` function is used to display the original and modified linked lists.

3. Delete from the Middle

Deleting a node from the middle of a linked list involves several steps to ensure the proper maintenance of the list structure and memory management.

Here are the steps:

1. Check for Empty List

Before attempting to delete a node from the middle, check if the linked list is empty. If the list is empty (head is `nullptr`), then there is nothing to delete.

2. Traverse to the Node to be Deleted

Traverse the list to reach the node to be deleted while keeping track of the preceding node. This is essential for updating the next pointer of the preceding node in the next step.

3. Update the Next Pointer of the Preceding Node

Once the node to be deleted is reached, update the next pointer of the preceding node to point to the node following the one being deleted. This effectively bypasses the node to be deleted.

4. Deallocate Memory

After updating the pointers, deallocate the memory occupied by the node that was removed. This step is crucial to prevent memory leaks.

Example:

```

#include <iostream>

struct Node {
    int data;
    Node* next;

    Node(int value) : data(value), next(nullptr) {}
};

void deleteFromMiddle(Node*& head, int value) {
    // Step 1: Check for an empty list
    if (head == nullptr) {
        std::cout << "List is empty. Nothing to delete." << std::endl;
        return;
    }

    // Step 2: Traverse to the Node to be Deleted
    Node* current = head;
    Node* previous = nullptr;

    while (current != nullptr && current->data != value) {
        previous = current;
        current = current->next;
    }

    // Check if the node to be deleted is found
    if (current == nullptr) {
        std::cout << "Node with value " << value << " not found in the list." << std::endl;
        return;
    }

    // Step 3: Update Next Pointer of the Preceding Node
    if (previous != nullptr) {
        previous->next = current->next;
    } else {
        // If the node to be deleted is the first node
        head = current->next;
    }

    // Step 4: Deallocate Memory
    delete current;
}

```

Part-2

```

// Function to display the linked list
void displayList(Node* head) {
    while (head != nullptr) {
        std::cout << head->data << " ";
        head = head->next;
    }
    std::cout << std::endl;
}

int main() {
    // Create an initial linked list with nodes: 2 -> 4 -> 6
    Node* head = new Node(2);
    head->next = new Node(4);
    head->next->next = new Node(6);

    // Display the original linked list
    std::cout << "Original Linked List: ";
    displayList(head);

    // Delete node with value 4 from the middle
    deleteFromMiddle(head, 4);

    // Display the linked list after deletion
    std::cout << "Linked List after Deletion from Middle: ";
    displayList(head);

    return 0;
}

```

In this example, the `deleteFromMiddle` function is used to delete a node with a specific value from the middle of the linked list, and the `displayList` function is used to display the original and modified linked lists.

4.12 Header Linked List

A header-linked list, also known as a sentinel or dummy node linked list, is a variation of a traditional linked list structure that incorporates an additional "header" or "dummy" node at the beginning.

Unlike a regular linked list where the head points directly to the first data node, in a header linked list, the head points to a dummy node that serves as a placeholder.

The purpose of this dummy node is to simplify certain operations and enhance the consistency of code. It ensures that the list is never empty, even when there are no data nodes.

This can simplify boundary checks and edge cases in various linked list operations. The dummy node does not store any meaningful data; its sole function is to make the list structure more uniform and easier to manage.

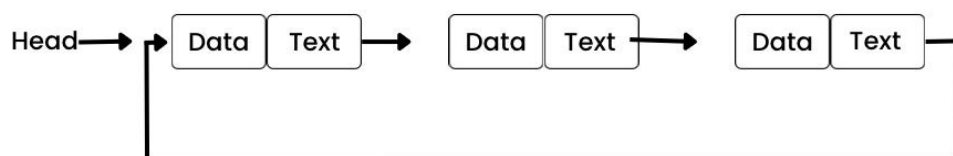
For instance, in an empty linked list, the header points to the dummy node, and in a non-empty list, it points to the first data node. This approach simplifies the insertion and deletion operations since the dummy node is always present, streamlining the code logic.

While the concept of a header linked list might seem unconventional at first, it provides a cleaner and more consistent way to handle various linked list scenarios, especially for individuals learning about linked list structures.

Summary

- A header linked list incorporates a dummy node at the beginning, distinct from the traditional linked list structure.
- The dummy node, serving as a placeholder, ensures the list is never empty, simplifying operations and making code more consistent.
- In scenarios with an empty list, the head points to the dummy node, and in a non-empty list, it points to the first data node.
- This approach streamlines insertion and deletion operations, offering a cleaner and more uniform way to manage linked lists, particularly beneficial for individuals learning about these data structures.

4.13 Two-Way List(Doubly linked list)



A two-way linked list, also known as a doubly linked list, is a data structure that extends the functionality of a traditional linked list by allowing each node to maintain links to both its successor and its predecessor.

Unlike a single-linked list where nodes only have a reference to the next node, in a two-way linked list, each node contains two pointers, one pointing to the next node and the other pointing to the previous node.

This bidirectional linkage provides advantages in both forward and backward traversal of the list.

It enables efficient operations such as traversing in both directions, inserting or deleting nodes at both the beginning and end of the list in constant time, and simplifying certain algorithms by providing easy access to the predecessor node.

However, the increased storage requirements (due to maintaining two pointers per node) and the added complexity in managing the bidirectional links are trade-offs that should be considered when deciding whether to use a two-way linked list over a single-linked list.

In summary, a two-way linked list enhances the capabilities of a linked list by allowing nodes to maintain connections in both directions, offering advantages in traversal and certain operations at the cost of increased memory usage and added complexity in managing bidirectional links.

Summary

- A two-way linked list, or doubly linked list, extends the functionality of a traditional linked list by enabling bidirectional traversal.
- In this data structure, each node contains pointers to both its next and previous nodes, facilitating efficient operations such as forward and backward traversal, as well as constant-time insertion and deletion at both ends of the list.
- While these advantages enhance certain algorithms, the trade-offs include increased memory requirements and added complexity in managing bidirectional links.

4.14 Implementation of Linked List

Implementing a linked list involves defining a structure for the nodes and creating functions to perform various operations on the list.

The basic structure of a node typically includes a data field and a pointer to the next node in the sequence. In languages like C or C++, one would use structs or classes to define the node structure.

For example, a simple implementation in C++ might look like:

```
struct Node {  
    int data;  
    Node* next;  
  
    Node(int value) : data(value), next(nullptr) {}  
};
```

The next step is to create functions for common linked list operations. These include functions to insert nodes at the beginning, middle, or end of the list, delete nodes, and traverse the list for display or manipulation.

Managing pointers correctly is crucial to ensure the integrity of the list structure and prevent memory leaks. Additionally, it's important to handle edge cases such as inserting or deleting nodes when the list is empty.

Lastly, memory management is a critical aspect of linked list implementation. Proper allocation of memory for new nodes using `new` or `malloc` and deallocating memory using `delete` or `free` in languages like C++ ensures the efficient use of resources and prevents memory leaks.

Overall, a well-implemented linked list provides a flexible and dynamic data structure suitable for various applications.

Summary

- Implementing a linked list involves defining a node structure and creating functions for common operations in languages like C or C++.
 - The node typically comprises a data field and a pointer to the next node.
 - Managing pointers is crucial for maintaining the list structure.
 - Functions for inserting nodes at different positions, deleting nodes, and traversing the list are essential.
 - Correctly handling edge cases and ensuring proper memory allocation and deallocation are key considerations.
 - A well-implemented linked list offers a flexible and dynamic data structure with efficient memory management, suitable for diverse applications.
-

CHAPTER 5:

TRESS



5.1 Tree

Trees are fundamental data structures in computer science that organize and store data in a hierarchical manner. They consist of nodes connected by edges, with a topmost node called the root.

Each node can have zero or more child nodes, forming a structure resembling an inverted tree. Nodes without children are referred to as leaves.

Trees are used in various applications, such as representing hierarchical relationships, organizing data efficiently in databases, and facilitating fast search and retrieval operations.

One common type of tree is the binary tree, where each node has at most two children. This structure is particularly useful in algorithms and applications like binary search trees, where data can be efficiently organized for quick search and insertion.

Other types of trees include balanced trees, such as AVL and red-black trees, which maintain a balance to ensure efficient operations.

Trees play a crucial role in enhancing the efficiency and performance of numerous algorithms and data storage systems.

Beyond binary trees, there are more advanced tree structures designed to address specific needs. One example is the B-tree, which is commonly used in databases and file systems. B-trees maintain balance through a variable number of child nodes per parent, optimizing for efficient disk access and storage.

Another notable type is the Tree, a tree-like structure that is particularly useful for storing and searching data with keys, such as in autocomplete systems or spell checkers.

Trees also find applications in graph theory, where hierarchical relationships and network structures are modeled. The spanning tree is an essential concept in graph

theory, helping to identify a subset of edges that form a tree and connect all the vertices without creating cycles.

Trees are versatile and provide a foundation for various algorithms and data structures, showcasing their significance across multiple domains in computer science and beyond.

Trees, fundamental data structures in computer science, organize data hierarchically with nodes connected by edges.

Binary trees, AVL trees, and B-trees are common types, each serving specific purposes like efficient search and storage in databases.

Tries excel in key-based data storage. In graph theory, spanning trees connect vertices without cycles.

Trees' versatility extends across diverse applications, making them crucial in computer science algorithms and data structures, showcasing their significance in various domains.

Summary

- Trees, fundamental data structures in computer science, organize data hierarchically with nodes connected by edges.
- Binary trees, AVL trees, and B-trees are common types, each serving specific purposes like efficient search and storage in databases.
- Tries excel in key-based data storage. In graph theory, spanning trees connect vertices without cycles.
- Trees' versatility extends across diverse applications, making them crucial in computer science algorithms and data structures, showcasing their significance in various domains.

5.1 Tree Terminologies

Tree terminology encompasses several key concepts. At the core is the "root," the topmost node from which all other nodes descend. Nodes with no children are "leaves," while those with children are "internal nodes."

"The "depth" of a node is the length of the path from the root, and the "height" is the length of the longest path to a leaf. Nodes with a common parent are "siblings," and a node's "ancestor" is any node on the path back to the root.

Binary trees specifically feature nodes with at most two children: a "left child" and a "right child." A "subtree" is a tree formed by a node and its descendants. Balanced trees, like AVL or red-black trees, maintain height balance for efficient operations.

In graph theory, a "spanning tree" connects all vertices without cycles. Overall, understanding these terms is fundamental for effectively working with tree structures in

computer science.

Continuing with tree terminology, a "forest" refers to a collection of disjoint trees, and a "parent" is a node's ancestor with a direct connection. The "degree" of a node is the count of its children, while a "binary search tree" (BST) follows an ordering property, where each node's left subtree contains values less than the node, and the right subtree contains values greater.

In B-trees, each node has multiple children, optimizing storage and access in databases. Tries, or prefix trees, excel in storing and searching key-based data. "Traversal" methods involve systematically accessing each node in a tree, with common types being in-order, pre-order, and post-order.

Understanding these terms provides a foundation for navigating, organizing, and efficiently processing data within tree structures across various computational applications.

Summary

- Tree terminology encompasses essential concepts in computer science.
- A tree's structure includes a root, leaves, internal nodes, and subtrees.
- Binary trees, with left and right children, and their balanced counterparts like AVL trees, play key roles.
- B-trees optimize storage in databases, while tries excel in key-based data storage.
- Understanding parent-child relationships, degrees, and traversal methods is crucial. In graph theory, spanning trees connect vertices without cycles.
- This comprehensive terminology provides a foundation for effectively utilizing trees in various computational applications, from organizing data to optimizing algorithms.

Degree of Node

In a tree, the "degree" of a node simply tells us how many branches or children it has. Imagine the node as a point where branches connect, and the degree is like counting how many paths extend from that point.

If a node has no children, we say it has a degree of zero, making it a leaf. If it has one child, its degree is one, and so on. The degree helps us understand the branching pattern of a tree and is a basic concept for anyone starting to explore trees in computer science.

Think of a tree like a family tree, where each person represents a node. The degree of a person in this family tree is the number of children they have.

If someone has no children, their degree is zero; if they have one child, it's one, and so on. So, the degree of a node in a tree is a simple way to count how many "offspring" or branches it has.

Understanding this concept is like figuring out how many descendants a particular family member has in our family tree analogy.

Summary

- In a tree structure, the degree of a node is like counting the number of children or branches it has.
- It helps us understand the branching pattern in the tree.
- If a node has no children, its degree is zero, and if it has one child, the degree is one, and so on.
- This concept is similar to counting the descendants in a family tree, making it a fundamental idea for beginners exploring trees in computer science.

Level of Node

The "level" of a node in a tree indicates its distance from the root. Picture the tree as a hierarchy, with the root at the top. Nodes directly connected to the root are at level one, and each level represents a step away from the root.

For instance, if a node is the child of another node, it is at a higher level than its parent. The level of a node helps us understand its position within the tree's structure, providing a straightforward way to describe its depth or distance from the top. Beginners can think of levels as steps in a hierarchy, starting from the root.

To illustrate the concept further, consider a family tree analogy. The root of the tree corresponds to the oldest generation, like grandparents. The next level would be their children, the next level their grandchildren, and so on.

Each level represents a generation, and the level of a person in this family tree indicates how many generations they are removed from the earliest ancestors.

In computer science trees, the root is at level one, its immediate children are at level two, and the pattern continues downward. Nodes at the same level share a common "distance" from the root.

Understanding node levels is crucial for tasks like analyzing the depth of a tree, comparing positions of nodes, and implementing algorithms that require knowledge of hierarchical relationships.

Summary

- In a tree structure, the "level" of a node signifies its distance from the root, akin to steps in a hierarchy.
- Nodes directly connected to the root are at level one, and each subsequent level represents a step away.
- This concept is analogous to generations in a family tree, where the oldest generation is at the root.

- Understanding node levels is essential for tasks involving hierarchical relationships, depth analysis, and comparisons within the tree structure.

Leaf Node

A leaf node in a tree structure is like the end of a branch; it's a node that doesn't have any children. In other words, it's a node that doesn't further branch out to other nodes. Think of a tree as a family tree, and the leaf nodes as individuals who don't have any descendants.

In computer science, leaf nodes are valuable because they often represent the actual data or information stored in the tree.

Identifying and understanding leaf nodes is crucial for various algorithms, especially those involved in retrieving and processing specific pieces of information within the tree.

Consider a tree as a hierarchy, where the root is at the top and each level represents a generation. A leaf node, in this analogy, is like an individual who doesn't have any offspring. It's the furthest point from the top, with no further branches below.

In computer science trees, leaf nodes are crucial because they hold specific data and mark the endpoints of a branch.

Identifying these leaf nodes becomes essential when navigating and extracting information from the tree, as they represent the final pieces of data in the hierarchical structure.

Summary

- In a tree structure, a leaf node is like the endpoint of a branch, representing a terminal point with no further subdivisions.
- In a family tree analogy, it corresponds to an individual without descendants.
- In computer science trees, leaf nodes often hold specific data and serve as endpoints, playing a crucial role in tasks like data retrieval and navigation within the hierarchical structure.
- Understanding and identifying leaf nodes is fundamental for efficiently working with tree-based data structures.

Depth/Height of Tree

The "depth" and "height" of a tree are essential metrics describing its structure. The depth of a tree is the length of the longest path from the root to any leaf node.

Think of it as measuring how far down the tree you have to go to reach the deepest point. On the other hand, the height of a tree is the length of the longest path from the root to any leaf, but it's often used interchangeably with "depth."

The terms reflect the overall vertical extent of the tree, showcasing its hierarchical depth. In computer science, understanding the depth and height of a tree is crucial for analyzing its efficiency, as it influences operations like searching and insertion.

To visualize this, consider a family tree. The depth is like counting how many generations you go back to reach the earliest ancestors, while the height measures the distance from the topmost ancestor to the farthest descendant.

In computer science trees, managing and optimizing these metrics becomes vital for designing algorithms that rely on the tree's hierarchical organization.

Summary

- In a tree structure, the "depth" signifies the length of the longest path from the root to any leaf, representing how far down the tree extends.
- The "height," often used interchangeably with depth, also measures the vertical extent but is more broadly understood as the distance from the root to the deepest leaf.
- In both cases, understanding these metrics is crucial in computer science for analyzing the efficiency of tree operations, such as searching and insertion, and optimizing algorithms based on the hierarchical structure of the tree.
- It's akin (similar) to measuring the generational depth in a family tree, emphasizing the vertical distance from ancestors to descendants.

In-Degree and Out-Degree

In the context of directed graphs, "degree" and "out-degree" provide insights into the connectivity of nodes. The "degree" of a node refers to the total number of edges connected to it, irrespective of direction.

It is like counting how many arrows point in and out of a particular node in a graph. Nodes with higher degrees generally indicate greater connectivity within the network.

On the other hand, "out-degree" specifically measures the number of outgoing edges from a node. It focuses on the arrows leaving a particular node in a directed graph.

Nodes with higher out-degrees suggest that they have more direct connections or relationships pointing to other nodes.

Understanding degree and out-degree is fundamental in graph theory and network analysis.

For instance, in social networks, nodes with high degrees might represent influential individuals, while nodes with high out-degrees could indicate prolific content creators or information sources.

These concepts are valuable for studying connectivity patterns and designing

algorithms tailored to specific network structures.

Summary

- In a directed graph, "degree" signifies the total number of edges connected to a node, encompassing both incoming and outgoing edges.
- It reflects the overall connectivity of a node within the network.
- Meanwhile, "out-degree" specifically counts the number of outgoing edges from a node, providing insight into the direct connections or relationships that the node initiates.
- These concepts are crucial in graph theory and network analysis, offering a means to understand connectivity patterns and identify nodes with significant influence or contribution within a directed graph.

Path

A "path" in a graph is a sequence of vertices where each adjacent pair is connected by an edge. Essentially, it represents a route or a journey from one vertex to another through the edges of the graph.

The length of a path is determined by the number of edges it contains. Paths play a crucial role in graph theory, aiding in the exploration and understanding of connections between different nodes in a network.

Paths can be classified based on their characteristics. A "simple path" doesn't revisit any vertex, except for the starting and ending vertices. In contrast, a "cycle" is a closed path that revisits at least one vertex, forming a loop.

The concept of paths is fundamental in various algorithms and applications, such as finding the shortest route between two nodes or analyzing the connectivity of a graph.

Understanding paths provides a basis for navigating and extracting meaningful information from graph structures.

Ancestor and Descendant Node

In a hierarchical structure like a family tree or a tree data structure, an "ancestor" node refers to any node that lies on the path from a given node to the root.

Ancestors are nodes closer to the root or the top of the hierarchy. They represent the preceding generations in the lineage, providing a historical context for a particular node.

Conversely, a "descendant" node is any node that can be reached by following the branches downward from a given node. Descendants are nodes further away from the root, representing subsequent generations or levels in the hierarchy.

They embody the future connections stemming from a particular node, and in a family tree, descendants include children, grandchildren, and so forth.

Understanding the relationships between ancestor and descendant nodes is fundamental for navigating and comprehending hierarchical structures, be it in genealogy, organizational hierarchies, or data structures.

It facilitates tasks like tracing lineages, retrieving information about specific branches, and comprehending the overall structure of complex systems.

Summary

- In a hierarchical structure such as a family tree or tree data structure, an "ancestor" node is any node closer to the root, representing preceding generations,
- While a "descendant" node is reached by following branches downward from a given node, representing subsequent generations.
- Ancestors provide historical context, and descendants embody future connections.
- Understanding these relationships is crucial for tasks like lineage tracing, information retrieval, and comprehending the overall structure of complex systems.

General Trees

General trees, in contrast to binary trees, allow nodes to have any number of children, not limited to two. This flexibility makes them versatile in representing hierarchical structures where nodes can have varying degrees.

Each node in a general tree can have a different number of child nodes, offering a more adaptable approach to modeling relationships and dependencies.

In a general tree, there is no restriction on the number of children a node can have, making it suitable for diverse applications. This flexibility comes in handy when dealing with scenarios where the hierarchical relationships are not strictly binary.

General trees find applications in various domains, including file systems, organizational hierarchies, and representing natural language syntax trees in linguistics.

The ability to model structures with arbitrary branching makes general trees a valuable tool in computer science and data representation.

Binary Trees

Binary trees are a fundamental data structure in computer science characterized by each node having at most two children: a left child and a right child.

The topmost node is called the root, and nodes with no children are referred to as leaves. Binary trees provide a hierarchical way of organizing data, where each node's left child contains values less than the node, and the right child contains values greater.

This ordering property facilitates efficient search algorithms, making binary trees a common choice for applications like binary search trees.

Binary trees come in various forms, including balanced types like AVL trees and red-black trees, which maintain balance to ensure efficient operations. They are widely used in databases, as well as in algorithms requiring fast search and retrieval.

The simplicity and efficiency of binary trees make them a cornerstone in computer science, forming the basis for more complex tree structures and algorithms.

5. Tree Types

Trees come in various types, each serving specific purposes in computer science and related fields. Binary trees are fundamental, with each node having at most two children, often used in search algorithms and data representation.

Balanced trees, such as AVL and red-black trees, maintain a balance to ensure efficient operations. Specialized structures like B-trees optimize storage and retrieval in databases.

Tries, or prefix trees, excel in key-based data storage and retrieval, making them valuable in applications like spell checkers and autocomplete systems.

Understanding these tree types provides a versatile toolkit for organizing and processing data in diverse computational scenarios.

Beyond these fundamental types, trees extend to cater to specific needs. Heap trees, like binary heaps, prioritize the ordering of their elements, making them efficient for priority queue implementations.

Quad trees and octrees are spatial partitioning trees often used in computer graphics and geographical information systems to organize spatial data efficiently. Trie variants, like ternary search trees, offer enhanced search capabilities.

Multiway trees, such as 2-3 trees and 2-3-4 trees, handle multiple keys per node, optimizing for balance and search operations. Trees, with their diverse types, prove essential in designing algorithms and data structures tailored to different computational challenges.

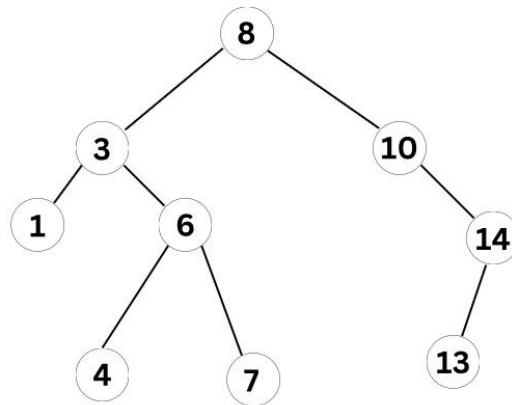
Summary

- Trees in computer science encompass various types, each designed for specific purposes.
- Binary trees, including balanced ones like AVL and red-black trees, excel in search algorithms and data representation.
- Specialized structures like B-trees optimize storage in databases, while tries are effective for key-based data storage.
- Heap trees prioritize ordering for efficient priority queues, and spatial partitioning trees like quadtrees find applications in computer graphics and GIS.

- Additional variants, such as ternary search trees and multiway trees, offer specialized solutions.
- This diverse range of tree types provides a versatile toolbox for organizing and processing data in different computational scenarios.

Binary Search Trees

Binary Search Tree



A binary search tree (BST) is a type of binary tree that follows a specific ordering property. In a BST, each node has at most two children: a left child and a right child.

The key idea is that the values in the left subtree of a node are less than the value of the node, and the values in the right subtree are greater. This ordering property holds true for every node in the tree.

The structure of a binary search tree facilitates efficient search, insertion, and deletion operations.

When searching for a specific value, the BST allows for a binary search, narrowing down the search space by comparing the target value with the values in each node and traversing either the left or right subtree accordingly.

This results in logarithmic time complexity for these operations in well-balanced trees.

Binary search trees are commonly used in applications where data needs to be organized and quickly retrieved in a sorted order.

However, maintaining balance in the tree is crucial to prevent degradation of performance, leading to variations like AVL trees and red-black trees that automatically balance themselves during insertions and deletions.

5.3 Binary Tree Traversal Methods

Binary tree traversal methods are systematic ways of visiting and processing each node in a binary tree.

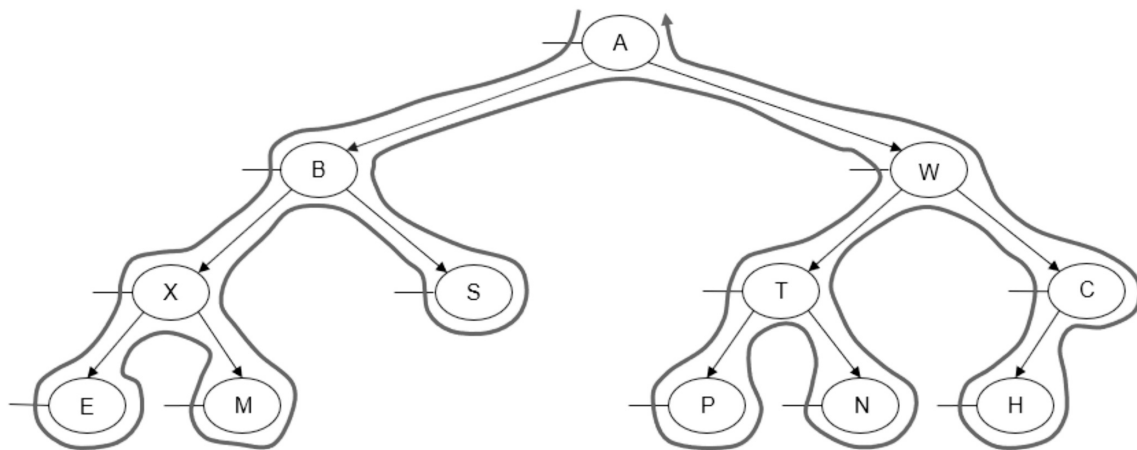
The three main traversal methods are :

- In-Order
- Pre-Order
- Post-Order.

1. In-Order Traversal

In-Order Traversal, each node in a binary tree is visited in a specific order, adhering to the following sequence:

Printing the value of each node as we "visit" it, we get the following output: E X M B S A P T N W H C



1. Visit the Left Subtree: Start by recursively traversing the left subtree of the current node.

2. Process the Root Node: After traversing the left subtree, process the current node or perform any desired operation on it.

3. Visit the Right Subtree: Finally, recursively traverse the right subtree of the current node.

This traversal method results in nodes being visited in ascending order for a Binary Search Tree (BST) due to the left-root-right sequence. In-order traversal is commonly used for tasks such as retrieving data from a BST in sorted order.

It provides a systematic way to explore the entire tree while maintaining the ordering property inherent in binary search trees.

Example:

```

#include <iostream>

using namespace std;

// Definition of a binary tree node
struct TreeNode {
    int data;
    TreeNode* left;
    TreeNode* right;

    TreeNode(int val) : data(val), left(nullptr), right(nullptr) {}
};

// Function to perform in-order traversal of a binary tree
void inOrderTraversal(TreeNode* root) {
    if (root == nullptr) {
        return;
    }

    // Recursively traverse the left subtree
    inOrderTraversal(root->left);

    // Visit the current node
    cout << root->data << " ";

    // Recursively traverse the right subtree
    inOrderTraversal(root->right);
}

int main() {
    // Example binary tree
    TreeNode* root = new TreeNode(1);
    root->left = new TreeNode(2);
    root->right = new TreeNode(3);
    root->left->left = new TreeNode(4);
    root->left->right = new TreeNode(5);
    root->right->left = new TreeNode(6);
    root->right->right = new TreeNode(7);

    cout << "In-order traversal: ";
    inOrderTraversal(root);

    // Clean up memory (deallocate the tree)
    // In a real-world scenario, you might use smart pointers or other memory management techniques.
    delete root->left->left;
    delete root->left->right;
    delete root->right->left;
    delete root->right->right;
    delete root->left;
    delete root->right;
    delete root;

    return 0;
}

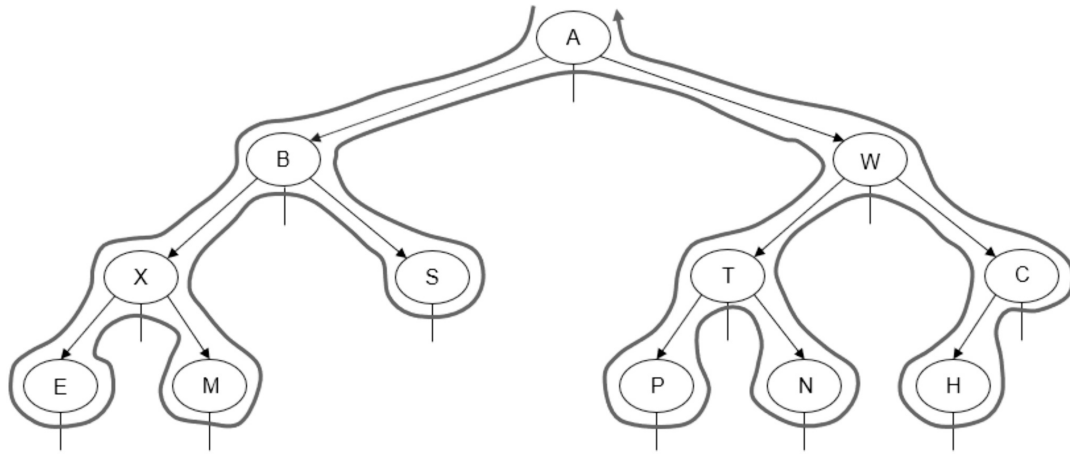
```

This program demonstrates the in-order traversal of a binary tree. The inorder traversal function recursively visits each node in the tree in the order of "left, root, right." The example binary tree is created and the in-order traversal result is printed. Finally, memory is deallocated to avoid memory leaks.

2. Pre-Order Traversal

In Pre-Order Traversal, each node in a binary tree is visited in a specific order, following the sequence.

Printing the value of each node as we "visit" it, we get the following output: A B X E M S W T P N C H



1. Process the Root Node: Start by processing the current node or performing any desired operation on it.

2. Visit the Left Subtree: After processing the current node, recursively traverse the left subtree.

3. Visit the Right Subtree: Finally, recursively traverse the right subtree of the current node.

This traversal method is called "pre-order" because the root node is processed before its children. Pre-order traversal is useful for scenarios where you need to perform an operation on a node before exploring its subtrees.

For example, it is commonly employed in creating a prefix expression for expression trees or when copying a tree structure. The order in which nodes are visited is root-left-right, providing a different perspective on the structure of the binary tree.

Example: Part-1

```

#include <iostream>

using namespace std;

// Definition of a binary tree node
struct TreeNode {
    int data;
    TreeNode* left;
    TreeNode* right;

    TreeNode(int val) : data(val), left(nullptr), right(nullptr) {}
};

// Function to perform pre-order traversal of a binary tree
void preOrderTraversal(TreeNode* root) {
    if (root == nullptr) {
        return;
    }

    // Visit the current node
    cout << root->data << " ";

    // Recursively traverse the left subtree
    preOrderTraversal(root->left);

    // Recursively traverse the right subtree
    preOrderTraversal(root->right);
}
  
```

Part-2

```
int main() {
    // Example binary tree
    TreeNode* root = new TreeNode(1);
    root->left = new TreeNode(2);
    root->right = new TreeNode(3);
    root->left->left = new TreeNode(4);
    root->left->right = new TreeNode(5);
    root->right->left = new TreeNode(6);
    root->right->right = new TreeNode(7);

    cout << "Pre-order traversal: ";
    preOrderTraversal(root);

    // Clean up memory (deallocate the tree)
    // In a real-world scenario, you might use smart pointers or other memory management techniques.
    delete root->left->left;
    delete root->left->right;
    delete root->right->left;
    delete root->right->right;
    delete root->left;
    delete root->right;
    delete root;

    return 0;
}
```

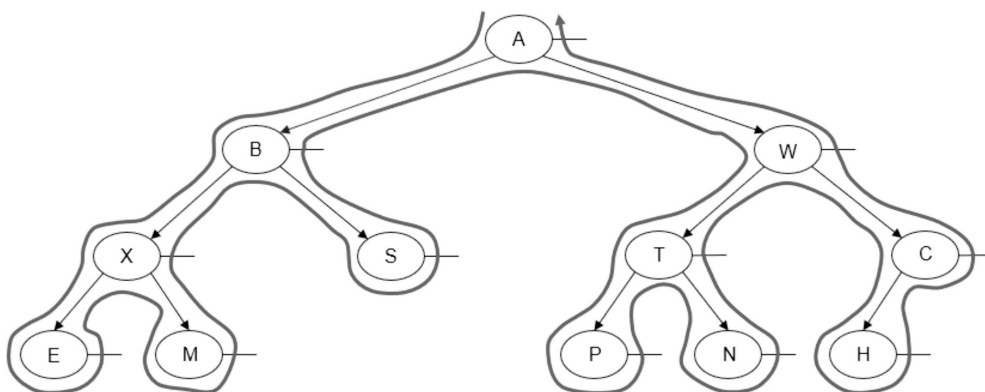
This program demonstrates the pre-order traversal of a binary tree. The preorder traversal function recursively visits each node in the tree in the order of "root, left, right." The example binary tree is created and the pre-order traversal result is printed. Finally, memory is deallocated to avoid memory leaks.

3. Post-Order Traversal

In Post-Order Traversal, each node in a binary tree is visited in a specific order, adhering to the following sequence.

Printing the value of each node as we visit it, we get the following output:

E M X S B P N T H C W A



- 1. Visit the Left Subtree:** Start by recursively traversing the left subtree of the current node.
- 2. Visit the Right Subtree:** After traversing the left subtree, recursively traverse the right subtree of the current node.
- 3. Process the Root Node:** Finally, process the current node or perform any desired operation on it.

This traversal method is called "post-order" because the root node is processed after its children. Post-order traversal is useful for scenarios where you want to perform operations on the subtrees before processing the root node.

It is commonly employed in tasks like deleting nodes from a tree, where it is generally safer to remove children before the parent. The order in which nodes are visited is left-right-root, providing a different perspective on the structure of the binary tree.

Example: Part-1

```
#include <iostream>

using namespace std;

// Definition of a binary tree node
struct TreeNode {
    int data;
    TreeNode* left;
    TreeNode* right;

    TreeNode(int val) : data(val), left(nullptr), right(nullptr) {}
};

// Function to perform post-order traversal of a binary tree
void postOrderTraversal(TreeNode* root) {
    if (root == nullptr) {
        return;
    }

    // Recursively traverse the left subtree
    postOrderTraversal(root->left);

    // Recursively traverse the right subtree
    postOrderTraversal(root->right);

    // Visit the current node
    cout << root->data << " ";
}

int main() {
    // Example binary tree
    TreeNode* root = new TreeNode(1);
    root->left = new TreeNode(2);
    root->right = new TreeNode(3);
    root->left->left = new TreeNode(4);
    root->left->right = new TreeNode(5);
    root->right->left = new TreeNode(6);
    root->right->right = new TreeNode(7);

    cout << "Post-order traversal: ";
    postOrderTraversal(root);

    // Clean up memory (deallocate the tree)
    // In a real-world scenario, you might use smart pointers or other memory management techniques.
    delete root->left->left;
    delete root->left->right;
    delete root->right->left;
    delete root->right->right;
    delete root->left;
    delete root->right;
    delete root;

    return 0;
}
```

This program demonstrates the post-order traversal of a binary tree. The postorder traversal function recursively visits each node in the tree in the order of "left, right, root." The example binary tree is created and the post-order traversal result is printed. Finally, memory is deallocated to avoid memory leaks.

5.4 Binary Search Tree (BST)

A Binary Search Tree (BST) is a binary tree data structure characterized by its ordering property. In a BST, each node has at most two children: a left child with values less than the node and a right child with values greater than the node.

This ordering facilitates efficient search operations, making it suitable for applications where quick retrieval and insertion of data are essential.

The key advantage of a BST lies in its ability to perform binary search operations, ensuring logarithmic time complexity for search, insertion, and deletion in well-balanced trees. The hierarchical structure and ordering property make in-order traversal of a BST result in a sorted sequence of elements.

To maintain efficiency, balanced variants of BSTs, such as AVL trees and red-black trees, automatically adjust their structure during insertions and deletions to prevent skewing and maintain optimal performance.

BST find extensive use in symbol tables, databases, and dynamic sets, providing an ordered and efficient means of organizing and accessing data.

However, it's crucial to manage balance to avoid worst-case scenarios, such as skewed trees, that could lead to degraded performance in terms of time complexity.

Summary

- A Binary Search Tree (BST) is a binary tree where each node has at most two children, and the values in the left subtree are less than the node, while values in the right subtree are greater.
- This ordering enables efficient search operations, with logarithmic time complexity for search, insertion, and deletion in well-balanced trees.
- In-order traversal of a BST produces a sorted sequence.
- Balanced variants like AVL and red-black trees automatically adjust to maintain optimal performance.
- BSTs are widely used in symbol tables, databases, and dynamic sets for organized and quick data retrieval, emphasizing the importance of managing balance for sustained efficiency.

Searching in BST

Searching in a Binary Search Tree (BST) involves efficiently locating a specific value by comparing it with nodes starting from the root. The process navigates through the tree's ordered structure, moving left or right based on comparisons until the target value is found or it's determined that the value is not present in the tree.

The effectiveness of BST search lies in its binary search algorithm, providing logarithmic time complexity for well-balanced trees and emphasizing the significance of tree balance for optimal performance.

STEPS

1. Start at the Root

Begin the search at the root of the tree.

2. Compare with Current Node

- Compare the target value with the value of the current node.
- If the target value is equal to the current node's value, the search is successful, and the node is found.

3. Navigate Left or Right

- If the target value is less than the current node's value, move to the left subtree.
- If the target value is greater, move to the right subtree.

4. Repeat Until Found

Repeat the process recursively or iteratively until the target value is found or until a leaf node is reached, indicating that the value is not present in the tree.

Due to the ordering property of the BST, this search strategy is akin to a binary search algorithm, systematically narrowing down the search space by traversing either the left or right subtree based on the comparisons.

The efficiency of this search is a key advantage of BSTs, providing a logarithmic time complexity ($O(\log n)$) in well-balanced trees, where n is the number of nodes.

It's important to note that in the worst-case scenario, a skewed tree (essentially a linked list) could result, leading to a degraded search performance with a time complexity of $O(n)$.

To mitigate this, balanced variants of BSTs, such as AVL trees or red-black trees, automatically adjust their structure during insertions and deletions to maintain balance and ensure optimal search efficiency.

Example:

```

#include <iostream>

using namespace std;

// Definition of a binary tree node
struct TreeNode {
    int data;
    TreeNode* left;
    TreeNode* right;

    TreeNode(int val) : data(val), left(nullptr), right(nullptr) {}
};

// Function to search for a key in a BST
TreeNode* searchBST(TreeNode* root, int key) {
    // Base case: If the tree is empty or the key is found
    if (root == nullptr || root->data == key) {
        return root;
    }

    // If the key is smaller than the root's key, search in the left subtree
    if (key < root->data) {
        return searchBST(root->left, key);
    }

    // If the key is larger than the root's key, search in the right subtree
    return searchBST(root->right, key);
}

```

```

int main() {
    // Example BST
    TreeNode* root = new TreeNode(8);
    root->left = new TreeNode(3);
    root->right = new TreeNode(10);
    root->left->left = new TreeNode(1);
    root->left->right = new TreeNode(6);
    root->left->right->left = new TreeNode(4);
    root->left->right->right = new TreeNode(7);
    root->right->right = new TreeNode(14);
    root->right->right->left = new TreeNode(13);

    // Key to search
    int key = 6;

    // Search for the key in the BST
    TreeNode* result = searchBST(root, key);

    // Check if the key is found or not
    if (result) {
        cout << "Key " << key << " is found in the BST.\n";
    } else {
        cout << "Key " << key << " is not found in the BST.\n";
    }

    // Clean up memory (deallocate the tree)
    // In a real-world scenario, you might use smart pointers or other memory management techniques.
    delete root->left->left;
    delete root->left->right->left;
    delete root->left->right->right;
    delete root->left->right;
    delete root->left;
    delete root->right->right->left;
    delete root->right->right;
    delete root->right;
    delete root;

    return 0;
}

```

This program demonstrates searching for a key in a Binary Search Tree (BST). The `searchBST` function recursively searches for the key in the BST. The example BST is created and the search result is printed. Finally, memory is deallocated to avoid memory leaks.

Summary

- Searching in a Binary Search Tree (BST) involves efficiently navigating the tree's ordered structure to locate a specific value.

- Beginning at the root, the algorithm compares the target value with the current node and moves left or right based on the comparison.
- This process continues recursively or iteratively until the target value is found or a leaf node is reached, indicating its absence in the tree.
- The binary search algorithm, inherent to BSTs, ensures logarithmic time complexity for well-balanced trees.
- However, in the worst-case scenario of a skewed tree, search performance may degrade to linear time complexity.
- Balanced variants like AVL trees or red-black trees mitigate this by automatically adjusting their structure to maintain balance, ensuring optimal search efficiency.

Inserting in BST

Inserting in a Binary Search Tree (BST) involves adding a new node with a specific value to the tree while maintaining its ordering property. The process starts at the root and follows a binary search algorithm comparing the value to be inserted with the current node's value and navigating left or right accordingly.

The goal is to find the appropriate position for the new node based on the ordering rules. Once a suitable spot is located, the new node is added, ensuring that the left subtree contains values less than the node, and the right subtree contains values greater.

Balancing actions may be required, especially in self-balancing trees like AVL or red-black trees, to maintain the tree's structural integrity and ordering property. Efficient insertion is essential for keeping the BST well-balanced and facilitating optimal search and retrieval operations.

Inserting a new node in a Binary Search Tree (BST) involves the following steps:

1. Start at the Root

Begin the insertion process at the root of the BST.

2. Compare with Current Node

Compare the value to be inserted with the value of the current node.

3. Navigate Left or Right

- If the value is less than the current node's value, move to the left subtree.
- If the value is greater, move to the right subtree.

4. Repeat Until Leaf Node Reached

Repeat the comparison and navigation process recursively or iteratively until a leaf node is reached. This leaf node will be the future parent of the new node.

5. Insert the New Node

Insert the new node as the left or right child of the leaf node, based on the final comparison.

6. Maintain Ordering Property

Ensure that the ordering property of the BST is maintained: values in the left subtree are less than the current node, and values in the right subtree are greater.

7. Balance the Tree (if needed)

If the BST is a self-balancing tree (e.g., AVL tree, red-black tree), perform necessary rotations or adjustments to maintain balance.

Following these steps guarantees the insertion of a new node into the BST while preserving its ordered structure. Efficient insertion is crucial for maintaining the overall balance of the tree, ensuring optimal performance for subsequent search and retrieval operations.

Example:

```
#include <iostream>

using namespace std;

// Definition of a binary tree node
struct TreeNode {
    int data;
    TreeNode* left;
    TreeNode* right;

    TreeNode(int val) : data(val), left(nullptr), right(nullptr) {}
};

// Function to insert a key into a BST
TreeNode* insertBST(TreeNode* root, int key) {
    // Base case: If the tree is empty, create a new node
    if (root == nullptr) {
        return new TreeNode(key);
    }

    // If the key is smaller than the root's key, insert in the left subtree
    if (key < root->data) {
        root->left = insertBST(root->left, key);
    }

    // If the key is larger than the root's key, insert in the right subtree
    else if (key > root->data) {
        root->right = insertBST(root->right, key);
    }

    return root;
}

// Function to perform in-order traversal of a BST
void inOrderTraversal(TreeNode* root) {
    if (root == nullptr) {
        return;
    }

    // Recursively traverse the left subtree
    inOrderTraversal(root->left);

    // Visit the current node
    cout << root->data << " ";

    // Recursively traverse the right subtree
    inOrderTraversal(root->right);
}
```

```

int main() {
    // Example BST
    TreeNode* root = nullptr;

    // Keys to insert into the BST
    int keys[] = {8, 3, 10, 1, 6, 4, 7, 14, 13};

    // Insert keys into the BST
    for (int key : keys) {
        root = insertBST(root, key);
    }

    // In-order traversal to verify the insertion
    cout << "In-order traversal after insertion: ";
    inOrderTraversal(root);
    cout << endl;

    // Clean up memory (deallocate the tree)
    // In a real-world scenario, you might use smart pointers or other memory management techniques.
    delete root->left->left;
    delete root->left->right->left;
    delete root->left->right->right;
    delete root->left->right;
    delete root->left;
    delete root->right->right->left;
    delete root->right->right;
    delete root->right;
    delete root;

    return 0;
}

```

This program demonstrates inserting keys into a Binary Search Tree (BST). The insertBST function recursively inserts a key into the BST. The example BST is created by inserting a set of keys, and an in-order traversal is performed to verify the insertion. Finally, memory is deallocated to avoid memory leaks.

Summary

- Inserting in a Binary Search Tree (BST) involves systematically adding a new node while adhering to the tree's ordering rules.
- The process begins at the root, where the binary search algorithm compares the value to be inserted with the current node's value, directing the insertion to the left or right subtree accordingly.
- The iteration continues recursively or iteratively until a leaf node is reached, which becomes the future parent of the new node.
- The new node is then inserted as the left or right child of the leaf node based on the final comparison.
- Crucially, this process ensures the maintenance of the ordering property: values in the left subtree are less than the current node, and values in the right subtree are greater.
- In self-balancing trees like AVL or red-black trees, additional balancing actions may be needed to uphold structural integrity.
- Efficient insertion is pivotal for sustaining the balance of the BST and facilitating optimal search and retrieval operations.

Deleting from BST

Deleting from a Binary Search Tree (BST) involves removing a node with a specific value from the tree while maintaining its structural integrity and the ordering property. The

process includes searching for the node to delete, considering various cases based on the number of children the node has (0, 1, or 2), and adjusting the tree accordingly.

If the node has no children, it is simply removed; if it has one child, it is replaced by that child; and if it has two children, its in-order successor or predecessor is used for replacement.

Balancing actions may be required, especially in self-balancing trees like AVL or red-black trees, to ensure that the ordering property is preserved and the tree remains balanced. Efficient deletion is crucial in maintaining the optimal performance of BSTs for search and retrieval operations.

STEPS:

1. Find the Node to Delete

Start by searching for the node containing the value to be deleted, following the same process as in searching. If the node is not found, the deletion operation concludes.

2. Node Found

- If the node is found, there are three possible cases to consider:

Case 1: Node has no Children (Leaf Node)

- Simply remove the node, and the tree remains balanced.

Case 2: Node has One Child

- Remove the node and replace it with its child, preserving the ordering property.

Case 3: Node has Two Children

Find the node's in-order successor (smallest value in its right subtree) or in-order predecessor (largest value in its left subtree).

- Replace the node's value with the in-order successor/predecessor.
- Recursively delete the in-order successor/predecessor.

3. Maintain Balance

- After deletion, it's crucial to ensure that the BST maintains its ordering property and remains balanced.
- For AVL trees or red-black trees, additional rotations or adjustments might be required to maintain balance.

The efficiency of deletion in a BST is influenced by the tree's balance.

Well-balanced trees, such as AVL trees or red-black trees, offer optimal time complexity for deletion operations ($O(\log n)$), ensuring that the tree remains balanced even after

repeated insertions and deletions.

Example:

```
#include <iostream>

using namespace std;

// Definition of a binary tree node
struct TreeNode {
    int data;
    TreeNode* left;
    TreeNode* right;

    TreeNode(int val) : data(val), left(nullptr), right(nullptr) {}
};

// Function to find the node with the minimum key value in a BST
TreeNode* findMinNode(TreeNode* node) {
    while (node->left != nullptr) {
        node = node->left;
    }
    return node;
}

// Function to delete a key from a BST
TreeNode* deleteNodeBST(TreeNode* root, int key) {
    // Base case: If the tree is empty
    if (root == nullptr) {
        return root;
    }

    // If the key to be deleted is smaller than the root's key, then it lies in the left subtree
    if (key < root->data) {
        root->left = deleteNodeBST(root->left, key);
    }
    // If the key to be deleted is larger than the root's key, then it lies in the right subtree
    else if (key > root->data) {
        root->right = deleteNodeBST(root->right, key);
    }
    // If the key to be deleted is equal to the root's key, then this is the node to be deleted
    else {
        // Node with only one child or no child
        if (root->left == nullptr) {
            TreeNode* temp = root->right;
            delete root;
            return temp;
        } else if (root->right == nullptr) {
            TreeNode* temp = root->left;
            delete root;
            return temp;
        }

        // Node with two children: Get the inorder successor (smallest in the right subtree)
        TreeNode* temp = findMinNode(root->right);

        // Copy the inorder successor's content to this node
        root->data = temp->data;

        // Delete the inorder successor
        root->right = deleteNodeBST(root->right, temp->data);
    }

    return root;
}
```

```

// Function to perform in-order traversal of a BST
void inOrderTraversal(TreeNode* root) {
    if (root == nullptr) {
        return;
    }

    // Recursively traverse the left subtree
    inOrderTraversal(root->left);

    // Visit the current node
    cout << root->data << " ";

    // Recursively traverse the right subtree
    inOrderTraversal(root->right);
}

int main() {
    // Example BST
    TreeNode* root = nullptr;

    // Keys to insert into the BST
    int keys[] = {8, 3, 10, 1, 6, 4, 7, 14, 13};

    // Insert keys into the BST
    for (int key : keys) {
        root = insertBST(root, key);
    }

    // In-order traversal before deletion
    cout << "In-order traversal before deletion: ";
    inOrderTraversal(root);
    cout << endl;

    // Key to delete from the BST
    int keyToDelete = 6;

    // Delete the key from the BST
    root = deleteNodeBST(root, keyToDelete);

    // In-order traversal after deletion
    cout << "In-order traversal after deletion: ";
    inOrderTraversal(root);
    cout << endl;

    // Clean up memory (deallocate the tree)
    // In a real-world scenario, you might use smart pointers or other memory management techniques.
    delete root->left->left;
    delete root->left->right->left;
    delete root->left->right->right;
    delete root->left->right;
    delete root->left;
    delete root->right->right->left;
    delete root->right->right;
    delete root->right;
    delete root;

    return 0;
}

```

This program demonstrates deleting a key from a Binary Search Tree (BST). The `deleteNodeBST` function recursively deletes a key from the BST. The example BST is created and an in-order traversal is performed before and after the deletion to verify the changes. Finally, memory is deallocated to avoid memory leaks.

Summary

- Deleting from a Binary Search Tree (BST) involves strategically removing a node with a specific value while upholding the tree's structural integrity and the ordering property.
- The process encompasses locating the node to delete, addressing different cases based on the node's number of children, and making appropriate adjustments to the tree.
- Three cases arise: when the node has no children when it has one child, and when it has two children.
- The last case involves finding the in-order successor or predecessor for replacement.

- Balancing actions, particularly in self-balancing trees like AVL or red-black trees, may be necessary to preserve the ordering property and balance.
- The deletion steps include finding the node, handling the three cases, and maintaining balance after deletion.
- The efficiency of deletion is enhanced in well-balanced trees, such as AVL or red-black trees, ensuring optimal time complexity ($O(\log n)$) for deletion operations, even with repeated insertions and deletions.

5.5 Heap

A heap is a specialized binary tree-based data structure that satisfies the heap property. In a max heap, for every node, the value of that node is greater than or equal to the values of its children.

This ensures that the maximum element is at the root. Conversely, in a min heap, the value of each node is less than or equal to the values of its children, placing the minimum element at the root.

Heaps are often implemented as arrays for efficient memory usage and support various operations, such as inserting elements, extracting the maximum or minimum, and heapifying – maintaining the heap property after insertions or deletions.

Heaps find applications in priority queues, sorting algorithms like heap sort, and graph algorithms like Dijkstra's shortest path algorithm.

Their ability to efficiently find and extract extremal elements makes them valuable in scenarios where prioritized access to elements is crucial.

Summary

- A heap is a binary tree-based data structure with a specialized order, where in a max heap, each node is greater than or equal to its children, and in a min heap, each node is less than or equal to its children.
- Heaps are commonly implemented as arrays and support operations like inserting, extracting extremal elements, and heapifying for maintaining the order.
- They are crucial in priority queues, sorting algorithms, and graph algorithms, offering efficient access to maximum or minimum elements.

Inserting Into a Heap

Inserting into a heap involves adding a new element to the next available position in the last level of the heap and then ensuring that the heap property is maintained.

For a max heap, this includes comparing the value of the newly inserted element with its parent and swapping them if the element's value is greater.

This process, known as "heapify up," is repeated until the element reaches a position where its value is less than or equal to its parent's value or becomes the root.

The goal is to preserve the heap property, where each node is greater than or equal to its children. Efficient insertion is essential for maintaining the structure of heaps, facilitating prioritized access to elements in algorithms like heap sort and priority queues.

Inserting into a heap involves the following steps:

1. Add Element

- Place the new element at the next available position in the last level of the heap, ensuring a left-to-right insertion.

2. Heapify Up

- Compare the value of the newly inserted element with its parent.
- If the value is greater (for a max heap) or smaller (for a min-heap) than the parent's value, swap the element with its parent.
- Continue this process iteratively until the element's value is less than or equal to its parent's value or until it becomes the root.

These steps ensure that the heap property is maintained after insertion.

The time complexity of the insertion operation is $O(\log n)$, where n is the number of elements in the heap, as it involves traversing the height of the heap. Efficient insertion is crucial for the proper functioning of heaps in various applications.

Example:

```

#include <iostream>
#include <vector>

using namespace std;

// Function to perform max heapify after inserting a new element
void heapifyUp(vector<int>& heap, int index) {
    while (index > 0) {
        int parent = (index - 1) / 2;

        if (heap[index] > heap[parent]) {
            swap(heap[index], heap[parent]);
            index = parent;
        } else {
            break;
        }
    }
}

// Function to insert a new element into a max heap
void insertIntoHeap(vector<int>& heap, int value) {
    heap.push_back(value);
    heapifyUp(heap, heap.size() - 1);
}

// Function to print the elements of a heap
void printHeap(const vector<int>& heap) {
    for (int value : heap) {
        cout << value << " ";
    }
    cout << endl;
}

int main() {
    // Example Max Heap
    vector<int> maxHeap = {90, 85, 80, 75, 70, 65, 60};

    cout << "Max Heap before insertion: ";
    printHeap(maxHeap);

    // Value to insert into the heap
    int valueToInsert = 82;

    // Insert the value into the Max Heap
    insertIntoHeap(maxHeap, valueToInsert);

    cout << "Max Heap after insertion: ";
    printHeap(maxHeap);

    return 0;
}

```

In this example, the `insertIntoHeap` function inserts a new element into a max heap represented by a vector. The `heapifyUp` function is called after insertion to maintain the heap property. The program demonstrates inserting an element (82 in this case) into a Max Heap and prints the heap before and after the insertion.

Summary

- Inserting into a heap involves placing a new element in the last available position in the heap and then maintaining the heap property.
- For a max heap, this entails comparing the new element's value with its parent and swapping them if the value is greater.
- Known as "heapify up," this process repeats until the element is appropriately positioned.
- The goal is to adhere to the heap property, ensuring each node is greater than or equal to its children.
- The insertion steps involve adding the element and iteratively heapifying up. The time complexity is $O(\log n)$, where n is the heap's size.
- Efficient insertion is vital for maintaining heap structure and supporting prioritized access in algorithms like heap sort and priority queues.

Deleting the root of the Heap

Deleting from the root of a heap involves removing the maximum (for a max heap) or minimum (for a min-heap) element from the top of the heap and restoring the heap properties.

The process begins by replacing the root with the last element in the heap, ensuring the structure remains complete.

Then, the replacement element is moved down the heap, swapping it with the larger (for a max heap) or smaller (for a min-heap) child as needed. This iterative process continues until the heap property is restored.

Deleting from the root is a crucial operation in heaps, often utilized in priority queues and sorting algorithms like heap sort. The time complexity is $O(\log n)$, where n is the number of elements, making it an efficient operation for maintaining the integrity of the heap.

The typical process is as follows:

1. Replace with Last Element

Swap the root, which is to be deleted, with the last element in the heap.

2. Heapify Down

Compare the new root with its children, and swap it with the larger (for a max heap) or smaller (for a min-heap) child if necessary.

Continue this process iteratively, moving down the heap, until the heap property is restored.

This process ensures that the heap remains complete and the ordering property is maintained. The time complexity of deleting from the root is $O(\log n)$, where n is the number of elements in the heap, as it involves traversing the height of the heap.

Deleting from the root is a fundamental operation in heaps and is crucial for maintaining their structure and efficiency in various algorithms and data structures.

Example:

```

#include <iostream>
#include <vector>

using namespace std;

// Function to perform max heapify after deleting the root element
void heapifyDown(vector<int>& heap, int index, int heapSize) {
    int largest = index;
    int leftChild = 2 * index + 1;
    int rightChild = 2 * index + 2;

    // Compare with left child
    if (leftChild < heapSize && heap[leftChild] > heap[largest]) {
        largest = leftChild;
    }

    // Compare with right child
    if (rightChild < heapSize && heap[rightChild] > heap[largest]) {
        largest = rightChild;
    }

    // If the largest value is not the current index, swap and recursively heapify down
    if (largest != index) {
        swap(heap[index], heap[largest]);
        heapifyDown(heap, largest, heapSize);
    }
}

```

```

// Function to delete the root element from a max heap
void deleteRoot(vector<int>& heap) {
    int heapSize = heap.size();

    if (heapSize == 0) {
        cout << "Heap is empty. Nothing to delete.\n";
        return;
    }

    // Swap the root with the last element
    swap(heap[0], heap[heapSize - 1]);

    // Reduce the heap size
    heapSize--;

    // Heapify down from the root
    heapifyDown(heap, 0, heapSize);

    // Remove the last element
    heap.pop_back();
}

// Function to print the elements of a heap
void printHeap(const vector<int>& heap) {
    for (int value : heap) {
        cout << value << " ";
    }
    cout << endl;
}

int main() {
    // Example Max Heap
    vector<int> maxHeap = {90, 85, 80, 75, 70, 65, 60};

    cout << "Max Heap before deletion: ";
    printHeap(maxHeap);

    // Delete the root element from the Max Heap
    deleteRoot(maxHeap);

    cout << "Max Heap after deletion: ";
    printHeap(maxHeap);

    return 0;
}

```

In this example, the `deleteRoot` function deletes the root element from a Max Heap represented by a vector. The `heapifyDown` function is called after deletion to maintain the heap property. The program demonstrates deleting the root element from a Max Heap and prints the heap before and after the deletion.

Summary

- Deleting the root of a heap involves removing the maximum or minimum element from the top and restoring the heap properties.
- The process begins by replacing the root with the last element, maintaining the heap's completeness.
- The replacement element is then moved down the heap, swapping with the appropriate child until the heap property is restored.
- This efficient operation is crucial for priority queues and heap sort algorithms, with a time complexity of $O(\log n)$, where n is the number of elements.
- The steps include replacing with the last element and iteratively heapifying down to ensure the heap remains structurally sound and adheres to the ordering property.

Heap sort

Heap Sort is a comparison-based sorting algorithm that efficiently sorts an array by constructing a max heap, where each parent is greater than or equal to its children.

The algorithm then repeatedly extracts the maximum element from the heap, placing it at the end of the array, and adjusts the heap accordingly. This process continues until the entire array is sorted.

Heap Sort has a consistent $O(n \log n)$ time complexity and is in place, making it suitable for scenarios where stable sorting and minimal additional memory usage are priorities.

The algorithm consists of two main phases:

1. Heap Construction

The initial phase involves building a max heap (for ascending order) or a min heap (for descending order) from the input array. This is done by treating the array as a binary tree and repeatedly qualifying it.

Starting from the last non-leaf node and moving upwards, the array is transformed into a heap such that the heap property is satisfied at every node: each node is greater than or equal to (for max heap) or less than or equal to (for min heap) its children.

2. Sorting

After constructing the heap, the maximum (for ascending order) or minimum (for descending order) element is at the root of the heap.

- Swap this root element with the last element in the array, effectively moving the maximum (or minimum) element to its sorted position at the end of the array.
- Reduce the size of the heap (considering the last element as sorted) and heapify the remaining elements to maintain the heap property.
- Repeat these steps until the entire array is sorted.

Heap Sort has a time complexity of $O(n \log n)$ for the worst, average, and best cases, making it a consistent performer.

While it is not as commonly used as quicksort or merge sort in practice, Heap Sort's in-place nature and stability under varying input conditions contribute to its relevance in certain scenarios.

Summary

- Heap Sort is a comparison-based sorting algorithm renowned for its efficiency in achieving $O(n \log n)$ time complexity.
 - The algorithm initially constructs a max heap, ensuring that each parent node is greater than or equal to its children.
 - This heap is formed by iteratively heapifying the array, treating it as a binary tree.
 - The sorting phase involves extracting the maximum element from the heap, placing it at the end of the array, and adjusting the heap accordingly.
 - This process is repeated until the entire array is sorted.
 - The in-place nature of Heap Sort, along with its stable performance, makes it suitable for scenarios prioritizing minimal additional memory usage and stable sorting.
 - While not as commonly used as other sorting algorithms, Heap Sort maintains its relevance in specific applications.
-



CHAPTER 6: GRAPH



6.1 Graph

A graph is a versatile and fundamental data structure in computer science that represents relationships between entities. It comprises a set of nodes, also known as vertices, and a set of edges connecting pairs of vertices.

Edges can be directed or undirected, indicating the nature of the relationship between nodes. Graphs can be weighted, with numerical values assigned to edges, reflecting attributes like distance or cost.

They find wide applications in modeling various real-world scenarios, from social networks and computer networks to transportation systems and recommendation algorithms.

The types of graphs are diverse, including directed and undirected graphs, weighted and unweighted graphs, cyclic and acyclic graphs, and connected and disconnected graphs, among others.

Each type serves different purposes and is applicable to specific problem domains. Graphs provide an intuitive representation of complex relationships and are crucial in algorithm design, enabling efficient solutions to problems such as pathfinding, network analysis, and recommendation systems.

The study and understanding of graphs play a pivotal role in addressing computational challenges across various domains.

Summary

- A graph is a foundational data structure representing relationships between nodes (vertices) through edges.
- These edges can be directed or undirected and may have weights, capturing diverse relationships.
- Graphs find applications in modeling real-world connections, from social networks to transportation systems.
- The various types of graphs, such as directed, weighted, and cyclic, cater to specific problem domains.

- Understanding and working with graphs are essential in algorithm design, enabling efficient solutions to problems like pathfinding and network analysis.

6.2 Graph Terminologies

In the realm of graph theory, various terminologies encapsulate essential concepts describing the structure and relationships within graphs. A vertex, also known as a node, represents a fundamental unit in a graph, embodying entities like cities, individuals, or web pages.

Edges denote connections between vertices, serving as links that can be either directed (with a specific origin and destination) or undirected, fostering mutual relationships. Edges may carry a weight, a numerical value representing factors such as distance, cost, or strength associated with the connection.

Weighted edges are particularly useful in scenarios like transportation networks where the intensity of the relationship matters.

Graphs themselves can take different forms, such as directed graphs (digraphs), where edges have a defined direction, or undirected graphs, where edges lack direction, symbolizing symmetric relationships. The degree of a vertex counts the number of edges connected to it, with directed graphs distinguishing between in-degree (incoming edges) and out-degree (outgoing edges).

Graphs can be categorized as cyclic if they contain cycles, forming closed loops, or acyclic if they lack such cycles. Acyclic graphs, like trees or directed acyclic graphs (DAGs), hold significance in algorithmic applications, such as topological sorting.

A path in a graph is a sequence of edges that connect a series of vertices, with the length of the path defined by the number of edges it comprises. Graphs can be classified as connected, meaning there exists a path between every pair of vertices, or disconnected, indicating the presence of isolated components.

These fundamental graph terminologies provide a comprehensive framework for expressing, analyzing, and solving problems related to relationships and connectivity within diverse applications.

Summary

- In graph theory, key terminologies elucidate fundamental concepts related to the structure and connections within graphs.
- Vertices, or nodes, represent entities, while edges signify connections between them, either directed or undirected.
- Weighted edges incorporate numerical values, such as distance or cost.
- Graphs can be directed or undirected, forming cyclic or acyclic structures, with degrees defining vertex connectivity.

- Paths represent sequences of connected edges and vertices, while graphs can be connected or disconnected.
- These terminologies provide a robust foundation for expressing and analyzing relationships and connectivity in diverse applications.

6.3 Nodes(Vertices)

In a graph, nodes play a crucial role in modeling real-world scenarios, such as social networks, transportation systems, or computational structures. The properties and attributes associated with nodes often determine the behavior and functionality of the overall system.

Nodes can represent diverse entities, such as individuals, cities, or devices, depending on the context of the application.

The arrangement and connections between nodes contribute to the graph's topology, influencing the efficiency of algorithms, the flow of information, and the study of patterns within the represented system.

Nodes in a graph can have additional characteristics, such as degrees, which indicate the number of edges connected to a node. High-degree nodes may act as central points of influence or critical components in certain applications.

Understanding the role and attributes of nodes is essential for effectively analyzing, modeling, and solving problems in graph theory, providing valuable insights into the structure and dynamics of interconnected systems.

Summary

- Nodes in a graph are fundamental components representing entities in various applications, such as social networks or transportation systems.
- Their properties influence the behavior of the overall system, and degrees, indicating edge connections, play a significant role.
- High-degree nodes can act as influential points.
- Understanding node attributes is crucial for analyzing and modeling graph structures, offering insights into the interconnected systems they represent.

6.4 Ares(Edge)

Edges in a graph are the connections or relationships between the nodes or vertices. Each edge represents a direct link between two vertices and can have additional information associated with it, such as weights or labels.

In a directed graph, edges have a specific direction, indicating a one-way connection from one vertex to another. Undirected graphs, on the other hand, have edges without a specific direction, signifying a mutual relationship between connected vertices.

The number of edges in a graph contributes to its structural characteristics and complexity. Graphs with fewer edges may form a sparse structure, while those with more edges may exhibit denser connectivity.

The presence and nature of edges play a vital role in various graph algorithms and analyses, impacting pathfinding, connectivity assessments, and network modeling.

Understanding and analyzing the edges in a graph are fundamental for gaining insights into relationships and patterns within the data it represents.

Summary

- Edges in a graph are the connections between nodes, representing relationships or links between vertices.
- They can be directional or undirected, indicating one-way or mutual connections.
- The quantity and nature of edges influence the graph's structure and complexity, impacting algorithms and analyses.
- Edges play a crucial role in tasks like pathfinding, connectivity assessments, and network modeling, providing insights into data relationships and patterns.
- Understanding and analyzing edges are essential for comprehending the overall structure and behavior of a graph.

6.5 Directed Graph

A directed graph, also known as a digraph, is a type of graph where edges have a direction, indicating a one-way connection between nodes.

In a directed graph, each edge connects two nodes, designating a starting node (tail) and an ending node (head). This directional aspect introduces a flow or orientation to the relationships between nodes.

Directed graphs find applications in various fields, such as representing dependencies, modeling relationships with specific directions, and analyzing processes where actions occur sequentially.

The presence and direction of edges in a directed graph provide a more nuanced representation of relationships, allowing for the modeling of scenarios where the order of connections is significant.

Summary

- In a directed graph, nodes represent entities, and directed edges signify relationships with a clear source and destination.
- These graphs can be acyclic or cyclic, depending on whether there are any directed cycles.

- Directed graphs are useful in various domains, such as representing dependencies in project management, expressing workflows, or modeling relationships with a natural direction.
- Analyzing directed graphs involves understanding connectivity patterns and considering the impact of directionality on traversal and information flow within the structure.

6.6 Undirected Graph

An undirected graph is a collection of nodes connected by edges without any specific direction. In this type of graph, edges represent symmetric relationships between nodes, meaning if there is an edge between node A and node B, it implies a mutual connection, and there is no inherent direction to the relationship.

Undirected graphs can be used to model various scenarios, such as social networks where friendships are reciprocal or road networks where the connectivity between locations is bidirectional.

The simplicity of undirected graphs makes them versatile for representing relationships in a wide range of fields.

Analyzing undirected graphs involves exploring connectivity patterns, identifying components, and understanding the overall structure of the relationships.

Additionally, algorithms for traversing, searching, and finding paths in undirected graphs play a fundamental role in solving problems across computer science, biology, and social sciences.

Summary

- An undirected graph is a type of graph in which nodes are connected by edges without a specified direction.
- These graphs represent symmetric relationships, where if there is a connection between two nodes, it is bidirectional.
- Undirected graphs find applications in modeling scenarios like social networks and road networks.
- Their simplicity makes them versatile for various fields, and analyzing them involves exploring connectivity, identifying components, and understanding overall relationship structures.
- Algorithms designed for traversing, searching, and finding paths in undirected graphs are fundamental tools with broad applications in computer science and other disciplines.

6.7 In-Degree and Out-Degree

In a directed graph, the terms "in-degree" and "out-degree" refer to the number of edges incident to a node concerning their direction.

1. In-Degree: The in-degree of a node is the count of incoming edges to that node. It represents the number of edges pointing towards the node.

Nodes with higher in-degrees often signify receiving more connections or dependencies in the context of various applications like web page ranking or network flow analysis.

2. Out-Degree: The out-degree of a node is the count of outgoing edges from that node. It represents the number of edges originating from the node.

Nodes with higher out-degrees might indicate more significant influence or contribution in scenarios such as social network analysis or flow networks.

Understanding in-degree and out-degree is crucial for analyzing the flow of information or resources in a directed graph.

These metrics provide insights into the relationships and dependencies among nodes, aiding in the interpretation and application of graph structures.

Summary

- In a directed graph, in-degree and out-degree metrics reveal valuable insights about node relationships.
- In-degree signifies the number of incoming edges, representing dependencies or connections pointing to a node,
- While out-degree indicates the count of outgoing edges, reflecting influences or contributions from a node.
- Analyzing these metrics aids in understanding information flow, dependencies, and the significance of nodes within the graph, crucial for applications such as web ranking, social network analysis, or flow networks.

6.8 Adjacent

Adjacent nodes in a graph are nodes that share a direct connection via an edge. This relationship plays a central role in graph representation and traversal.

In an undirected graph, adjacency is symmetric, meaning if node A is adjacent to node B, then node B is also adjacent to node A.

In a directed graph, adjacency indicates a directed connection from one node to another.

The adjacency concept is fundamental for tasks like pathfinding, network analysis, and understanding the structural relationships within a graph.

It forms the basis for various algorithms and operations that rely on the direct connections between nodes.

6.9 Successor

In the context of a graph, the term "successor" refers to a node that is directly reachable from another node through an outgoing edge.

In a directed graph, each node can have multiple successors, representing the nodes to which it has directed edges.

The concept is particularly relevant when analyzing paths, traversing graphs, or understanding the flow of relationships in directed networks.

For example, in a flow network or a process model, a node's successors might represent the subsequent steps or stages that follow it.

Successors provide a way to navigate through the directed connections in a graph, aiding in tasks such as route planning, process optimization, or identifying dependencies in systems.

6.10 Predecessor

In the context of a graph, the term "predecessor" refers to a node that has a directed edge leading to another specific node. It represents nodes that directly connect to a given node through incoming edges.

In a directed graph, each node may have multiple predecessors, illustrating the nodes from which it is reachable.

Predecessors are significant in understanding dependencies and relationships within a graph. They provide insights into the sources or origins of information, influence, or dependencies for a particular node.

For example, in a process model, a node's predecessors might signify the steps or stages that directly lead to it.

Analyzing predecessors is crucial for tasks such as tracing information flow, identifying root causes, or optimizing processes in systems represented by directed graphs.

6.11 Relation

In graph theory, a "relation" typically refers to the connection or association between nodes or vertices in a graph. A relation can exist between two nodes if there is an edge connecting them.

The nature of the relation depends on the type of graph. In a directed graph, the relation between nodes is directional, meaning it has a specific direction associated with the edges.

Each edge connects a source node to a target node, representing a one-way relationship. This directional aspect is crucial for understanding dependencies or influences between nodes.

In an undirected graph, relations are symmetric, as edges do not have a specific direction. The connection between nodes is mutual, indicating a two-way relationship. Undirected graphs are often used to represent symmetric relationships, such as friendships in a social network.

Analyzing relations in a graph provides insights into the connectivity, dependencies, and interactions between different elements.

It forms the basis for studying the structure and properties of graphs, allowing for a better understanding of complex relationships in various applications.

6.12 Path

In graph theory, a "path" is a sequence of edges that connects a sequence of vertices in a graph.

It represents a way to move from one vertex to another by traversing edges in the graph. The length of a path is given by the number of edges it contains.

1. Simple Path: A simple path is a path where no vertex is repeated, except possibly for the first and last vertices, which are the same in the case of a closed path or cycle.

2. Cycle: A cycle is a closed path where the first and last vertices are the same, and no other vertices are repeated.

3. Walk: A walk is a sequence of vertices and edges in a graph where consecutive vertices are connected by consecutive edges.

Paths are fundamental in understanding the connectivity and accessibility within a graph. They play a crucial role in algorithms and analyses related to network routing, optimization, and traversal.

The shortest path between two vertices is often of particular interest and is a key concept in various applications like transportation networks, computer networks, and social networks.

6.13 Sink

In graph theory, a "sink" refers to a vertex in a directed graph that has no outgoing edges.

Essentially, it is a node in the graph where information or influence flows in but doesn't flow out.

Sinks are sometimes also known as "absorbing vertices" because they absorb any flow or information that reaches them.

In the context of network analysis or directed flow models, identifying sinks is valuable for understanding how information or resources may accumulate or terminate in a system.

Sinks play a role in various applications, such as determining critical points in a network, analyzing the reachability of certain vertices, or identifying endpoints in processes like Markov chains.

6.14 Linear Representation of Graph

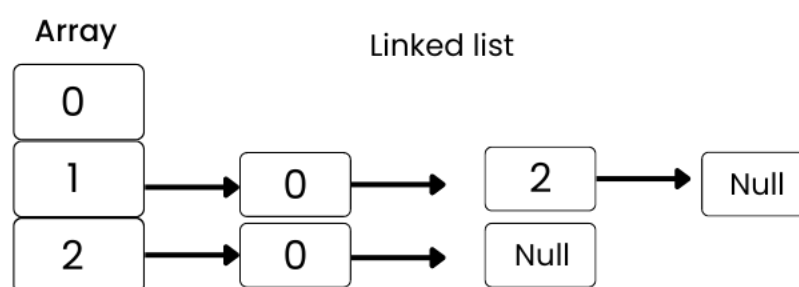
Linear representation of a graph involves converting the graph's structure into a linear data structure, facilitating efficient storage and traversal. Two common methods for linear representation are adjacency matrices and adjacency lists.

1. Adjacency Matrix

In this representation, a 2D array is used to depict the graph. Rows and columns correspond to vertices, and matrix entries signify the presence or absence of edges. For undirected graphs, the matrix is symmetric.

The matrix's (i, j) entry is non-zero if there is an edge between vertices i and j . This approach is efficient for dense graphs where many edges exist.

2. Adjacency List



Using arrays or linked lists, adjacency lists represent the graph by associating each vertex with a list of its adjacent vertices. For weighted graphs, the list may also include edge weights.

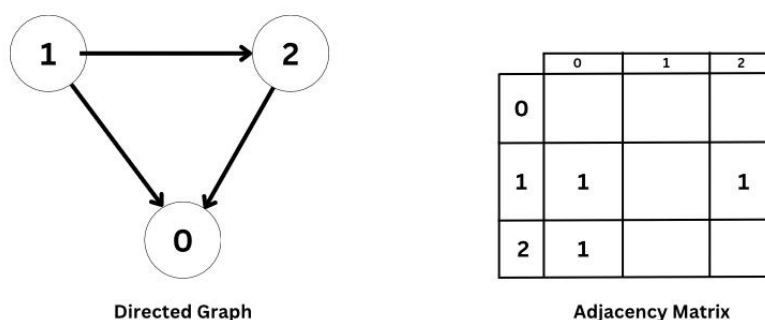
This representation is more memory-efficient, especially for sparse graphs where the number of edges is considerably less than the possible edges. It enables faster access to adjacent vertices during traversal.

The selection between adjacency matrices and lists depends on the graph's characteristics, with matrices being suitable for dense graphs and lists for sparse graphs, considering factors like memory usage and traversal efficiency.

Summary

- The linear representation of a graph involves transforming its structure into a linear data structure, streamlining storage, and traversal.
- Two widely used methods for linear representation are adjacency matrices and adjacency lists.
- Adjacency matrices employ a 2D array, where entries indicate the presence or absence of edges between vertices.
- This method is efficient for dense graphs with numerous edges.
- On the other hand, adjacency lists utilize arrays or linked lists to associate each vertex with a list of adjacent vertices, offering better memory efficiency and faster access to sparse graphs.
- The choice between matrices and lists depends on the graph's density, with matrices suited for dense graphs and lists for sparse ones, considering factors like memory usage and traversal efficiency.

6.15 Adjacency Matrix of Directed Graph



Graph Representation of Directed Graph to Adjacency Matrix

An adjacency matrix for a directed graph is a square matrix used to represent the graph's structure, where rows and columns correspond to vertices.

However, in this context, the matrix is not necessarily symmetric as it is in the case of undirected graphs. The matrix entries indicate the existence and direction of edges between vertices.

For a directed graph with n vertices, the adjacency matrix is an $n \times n$ matrix (let's call it A).

The entry $A[i][j]$ is used to represent the existence of an edge from vertex i to vertex j . Here's how the entries are interpreted:

- If there is an edge from vertex i to vertex j , then $A[i][j] = 1$.

- If there is no edge from vertex i to vertex j , then $A[i][j] = 0$.

This representation is effective for directed graphs because it clearly captures the directionality of edges.

It's worth noting that in the case of weighted directed graphs, the matrix entries could represent the weights of the edges instead of simply 0 or 1.

The adjacency matrix of a directed graph is a concise and straightforward way to represent the graph's connectivity and is particularly useful for certain graph algorithms that benefit from quick access to edge information.

However, it may not be the most memory-efficient choice for sparse graphs with many missing edges.

Summary

- An adjacency matrix for a directed graph is a square matrix where rows and columns correspond to vertices, indicating the existence and direction of edges.
- In this matrix, entries $A[i][j]$ are 1 if there is a directed edge from vertex i to vertex j and 0 otherwise.
- This representation is effective in capturing the directional relationships in directed graphs, making it useful for algorithms requiring quick access to edge information.
- However, it may not be optimal for sparse graphs due to potential memory inefficiency.

6.16 Warshall's Algorithm

The Warshall algorithm, also known as the Floyd-Warshall algorithm, is a dynamic programming approach used to find the transitive closure of a directed graph.

The transitive closure of a graph represents all pairs of vertices that are connected by some path, taking into account the direction of edges.

The algorithm iteratively computes the transitive closure matrix by considering all vertices as possible intermediate steps in connecting pairs of vertices.

It starts with an initial adjacency matrix and updates it in a way that reflects whether there exists a path between any two vertices, considering all vertices as potential intermediaries.

The Warshall algorithm is versatile, providing a straightforward way to compute the transitive closure for dense graphs.

However, its time complexity is cubic in the number of vertices, making it less efficient for sparse graphs compared to other algorithms like depth-first search.

The Warshall algorithm is a straightforward method for finding the transitive closure of a directed graph.

The steps involved are as follows:

- 1. Initialization:** Begin with the adjacency matrix representing the direct connections in the graph. For each pair of vertices (i, j) , if there is a direct edge from i to j , set the corresponding entry in the transitive closure matrix to 1; otherwise, set it to 0.
- 2. Iterative Update:** For each vertex k in the graph, iterate through all pairs of vertices (i, j) and update the transitive closure matrix entry (i, j) to 1 if there is a path from i to j passing through vertex k . This step involves updating the matrix based on whether the inclusion of vertex k creates a path between i and j .
- 3. Final Result:** After completing the iterations for all vertices, the transitive closure matrix reflects the connectivity between all pairs of vertices in the directed graph. The entry (i, j) is 1 if there exists a path from vertex i to vertex j , considering all possible paths in the graph.

The Warshall algorithm efficiently computes the transitive closure, capturing the reachability information in a directed graph.

Summary

- The Warshall algorithm, also known as the Floyd–Warshall algorithm, is a dynamic programming approach used to determine the transitive closure of a directed graph.
- This closure represents all pairs of vertices connected by some path, considering the direction of edges.
- The algorithm initializes with an adjacency matrix, progressively updating it to indicate the existence of paths between vertices by considering all possible intermediate vertices.
- While versatile, providing a straightforward solution for dense graphs, its cubic time complexity in the number of vertices makes it less efficient for sparse graphs compared to other methods.
- The algorithm involves initializing the transitive closure matrix, iteratively updating it for all possible intermediate vertices, and producing a final matrix that reflects the connectivity between all pairs of vertices in the directed graph.

6.17 Shortest paths

The shortest path problem in graph theory aims to find the most efficient path between two vertices in a graph, where the term "efficient" is often defined by the sum of weights associated with edges.

The problem can be formulated in various ways, but the most common scenarios involve finding the path with the minimum total edge weight.

1. Dijkstra's Algorithm

One popular algorithm for solving the single-source shortest path problem is Dijkstra's algorithm. Starting from a source vertex, it greedily selects the vertex with the smallest tentative distance and updates the distances to its neighboring vertices.

This process continues until the algorithm has explored all vertices, guaranteeing the shortest paths from the source to all other vertices.

Here are the steps of Dijkstra's algorithm for finding the shortest paths from a source vertex to all other vertices in a weighted graph:

1. Initialization

Set the initial distance to the source vertex as 0, and distances to all other vertices as infinity.

- Mark all vertices as unvisited.
- Create a priority queue (or min-heap) to store vertices with their tentative distances.

2. Start with the Source

Enqueue the source vertex with its tentative distance (0) into the priority queue.

3. Process Each Vertex

- While the priority queue is not empty:
- Dequeue the vertex with the smallest tentative distance.
- Mark the vertex as visited.

4. Update Neighbors

- For each unvisited neighbor of the current vertex:
- Calculate the tentative distance from the source to the neighbor through the current vertex.
- If this tentative distance is less than the recorded distance for the neighbor, update the neighbor's distance.

5. Repeat

Repeat steps 3 and 4 until all vertices are visited or the priority queue is empty.

6. Result

The final recorded distances represent the shortest paths from the source to all other vertices.

Dijkstra's algorithm guarantees finding the shortest paths in non-negative weighted graphs. The priority queue helps in selecting the vertex with the smallest tentative

distance efficiently.

Example: Part-1

```
#include <iostream>
#include <vector>
#include <queue>
#include <limits>

using namespace std;

// Structure to represent an edge in the graph
struct Edge {
    int destination, weight;
};

// Comparison operator for priority queue
struct Compare {
    bool operator()(const pair<int, int>& a, const pair<int, int>& b) {
        return a.second > b.second;
    }
};

// Function to perform Dijkstra's algorithm
void dijkstra(int vertices, const vector<vector<Edge>>& graph, int start) {
    // Initialize distance array with maximum values
    vector<int> distance(vertices, numeric_limits<int>::max());
    distance[start] = 0;

    // Priority queue to store vertices and their distances
    priority_queue<pair<int, int>, vector<pair<int, int>>, Compare> pq;
    pq.push({start, 0});

    while (!pq.empty()) {
        int current = pq.top().first;
        int currentDistance = pq.top().second;
        pq.pop();

        // Update distances for adjacent vertices
        for (const Edge& edge : graph[current]) {
            int newDistance = currentDistance + edge.weight;
            if (newDistance < distance[edge.destination]) {
                distance[edge.destination] = newDistance;
                pq.push({edge.destination, newDistance});
            }
        }
    }

    // Print the shortest distances
    cout << "Shortest distances from node " << start << ":\n";
    for (int i = 0; i < vertices; ++i) {
        cout << "To node " << i << ": " << distance[i] << endl;
    }
}
```

Part-2

```
int main() {
    // Example graph represented as an adjacency list of edges
    vector<vector<Edge>> graph = {
        {{1, 5}, {2, 4}},
        {{3, 3}},
        {{1, -6}, {3, 8}, {4, 7}},
        {{4, 9}},
        {{3, 6}}
    };

    int vertices = graph.size(); // Number of vertices in the graph
    int startNode = 0; // Starting node for Dijkstra's algorithm

    dijkstra(vertices, graph, startNode);

    return 0;
}
```

This program demonstrates Dijkstra's algorithm using a simple directed graph with weighted edges.

The Dijkstra function calculates the shortest paths and prints the results. The priority queue is used to efficiently select the vertex with the minimum distance at each step.

Summary

- Dijkstra's algorithm is a methodical approach for finding the shortest paths from a source vertex to all other vertices in a weighted graph.
- It starts by initializing distances, marking vertices as unvisited, and employing a priority queue.
- The algorithm systematically explores vertices, updating distances and selecting the nearest unvisited vertex until all vertices are visited or the queue is empty.
- The final recorded distances provide the shortest paths from the source vertex to every other vertex.
- Dijkstra's algorithm is efficient and particularly useful for navigating networks where edges have non-negative weights.

2. Bellman-Ford Algorithm

Another widely used algorithm is the Bellman-Ford algorithm, which handles graphs with negative edge weights.

It iteratively relaxes the edges, gradually improving the estimated shortest distances. After a certain number of iterations, the algorithm ensures that it has found the shortest paths.

Here are the steps of the Bellman-Ford algorithm, which is used to find the shortest paths from a single source vertex to all other vertices in a weighted graph, including graphs with negative weight edges:

1. Initialization

- Set the distance to the source vertex as 0, and distances to all other vertices as infinity.
- Mark all vertices as unvisited.

2. Relaxation

- Repeat the relaxation step for each edge in the graph $(V-1)$ times, where (V) is the number of vertices.
- For each edge (u, v) with weight (w) , if the distance to vertex (u) plus (w) is less than the current distance to vertex (v) , update the distance to (v) with the new value.

3. Check for Negative Cycles

After $(V-1)$ iterations, if there are further updates, it indicates the presence of a negative weight cycle in the graph. The algorithm can be halted, and the negative cycle can be identified using additional iterations.

4. Result

The final recorded distances represent the shortest paths from the source to all other vertices.

The Bellman-Ford algorithm is versatile and can handle graphs with negative weights, making it useful in scenarios where Dijkstra's algorithm may not be applicable.

However, it is less efficient than Dijkstra's algorithm for non-negative weight graphs.

Example: Part-1

```
#include <iostream>
#include <vector>
#include <limits>

using namespace std;

// Structure to represent an edge in the graph
struct Edge {
    int source, destination, weight;
};

// Function to perform Bellman-Ford algorithm
void bellmanFord(int vertices, int edges, const vector<Edge>& graph, int start) {
    // Initialize distance array with maximum values
    vector<int> distance(vertices, numeric_limits<int>::max());
    distance[start] = 0;

    // Relax edges repeatedly (V-1 times)
    for (int i = 1; i < vertices; ++i) {
        for (const Edge& edge : graph) {
            if (distance[edge.source] != numeric_limits<int>::max() &&
                distance[edge.source] + edge.weight < distance[edge.destination]) {
                distance[edge.destination] = distance[edge.source] + edge.weight;
            }
        }
    }

    // Check for negative weight cycles
    for (const Edge& edge : graph) {
        if (distance[edge.source] != numeric_limits<int>::max() &&
            distance[edge.source] + edge.weight < distance[edge.destination]) {
            cout << "Graph contains a negative weight cycle!\n";
            return;
        }
    }

    // Print the shortest distances
    cout << "Shortest distances from node " << start << ":\n";
    for (int i = 0; i < vertices; ++i) {
        cout << "To node " << i << ": " << distance[i] << endl;
    }
}
```

Part-2

```

int main() {
    // Example graph represented as a list of edges
    vector<Edge> graph = {
        {0, 1, 5},
        {0, 2, 4},
        {1, 3, 3},
        {2, 1, -6},
        {2, 3, 8},
        {2, 4, 7},
        {3, 4, 9},
        {4, 3, 6}
    };

    int vertices = 5; // Number of vertices in the graph
    int edges = graph.size(); // Number of edges in the graph
    int startNode = 0; // Starting node for Bellman-Ford algorithm

    bellmanFord(vertices, edges, graph, startNode);

    return 0;
}

```

This program demonstrates the Bellman-Ford algorithm using a simple directed graph with weighted edges. The `bellmanFord` function calculates the shortest paths and prints the results. The algorithm also checks for the presence of negative weight cycles in the graph.

Summary

- The Bellman-Ford algorithm, employed for finding the shortest paths from a source vertex to all other vertices in a weighted graph, follows a multi-step process.
- It initializes distances, marking the source as 0 and all others as infinity. The algorithm iteratively relaxes edges, updating distances if a shorter path is discovered.
- After $(V-1)$ iterations (V being the number of vertices), the recorded distances represent the shortest paths.
- Detection of further updates indicates a negative weight cycle.
- Despite being capable of handling negative weights, the Bellman-Ford algorithm is less efficient than Dijkstra's algorithm for non-negative weight graphs.

3. Floyd-Warshall Algorithm

The Floyd-Warshall algorithm is suitable for finding the shortest paths between all pairs of vertices in a graph.

It uses dynamic programming to iteratively update the matrix of shortest path distances, considering all possible intermediate vertices. This algorithm is versatile but can be less efficient than Dijkstra's or Bellman-Ford for specific scenarios.

The choice of algorithm depends on factors such as the graph's characteristics, the presence of negative weights, and the specific requirements of the application.

These algorithms play a crucial role in network routing, transportation planning, and various real-world scenarios where finding the most efficient paths is essential.

The Floyd–Warshall algorithm is employed for finding the shortest paths between all pairs of vertices in a weighted graph.

The algorithm follows these steps:

1. Initialization

- Create a distance matrix, where ($\text{dist}[i][j]$) represents the shortest distance between vertices (i) and (j).
- Initialize the matrix with direct edge weights between vertices, setting ($\text{dist}[i][j]$) to the weight of the edge connecting (i) and (j), and ($\text{dist}[i][i]$) to 0.

2. Iterative Updates

- For each vertex (k) in the graph:
- For each pair of vertices (i) and (j):
- If the path from (i) to (j) through (k) is shorter than the current ($\text{dist}[i][j]$), update ($\text{dist}[i][j]$) with the shorter path.

3. Result

- After completing all iterations, the (dist) matrix contains the shortest paths between all pairs of vertices.

The Floyd–Warshall algorithm has a time complexity of ($O(V^3)$) and is suitable for dense graphs or those with negative weights. However, its efficiency may be surpassed by other algorithms for sparse graphs.

Example: Part-1

```
#include <iostream>
#include <vector>
#include <limits>

using namespace std;

// Function to perform Floyd-Warshall algorithm
void floydWarshall(int vertices, vector<vector<int>>& distance) {
    // Update distances for all pairs of vertices
    for (int k = 0; k < vertices; ++k) {
        for (int i = 0; i < vertices; ++i) {
            for (int j = 0; j < vertices; ++j) {
                if (distance[i][k] != numeric_limits<int>::max() &&
                    distance[k][j] != numeric_limits<int>::max() &&
                    distance[i][k] + distance[k][j] < distance[i][j]) {
                    distance[i][j] = distance[i][k] + distance[k][j];
                }
            }
        }
    }

    // Print the shortest distances
    cout << "Shortest distances between all pairs of nodes:\n";
    for (int i = 0; i < vertices; ++i) {
        for (int j = 0; j < vertices; ++j) {
            if (distance[i][j] == numeric_limits<int>::max()) {
                cout << "INF\t";
            } else {
                cout << distance[i][j] << "\t";
            }
        }
        cout << endl;
    }
}
```

Part-2

```
int main() {
    // Example graph represented as a matrix of distances
    vector<vector<int>> distance = {
        {0, 5, 4, numeric_limits<int>::max(), numeric_limits<int>::max()},
        {numeric_limits<int>::max(), 0, numeric_limits<int>::max(), 3, numeric_limits<int>::max()},
        {numeric_limits<int>::max(), numeric_limits<int>::max(), 0, numeric_limits<int>::max(), 7},
        {2, numeric_limits<int>::max(), numeric_limits<int>::max(), 0, numeric_limits<int>::max()},
        {numeric_limits<int>::max(), numeric_limits<int>::max(), numeric_limits<int>::max(), 6, 0}
    };

    int vertices = distance.size(); // Number of vertices in the graph
    floydWarshall(vertices, distance);

    return 0;
}
```

This program demonstrates the Floyd–Warshall algorithm using a simple directed graph with weighted edges. The `floydWarshall` function calculates the shortest paths between all pairs of vertices and prints the results. The `numeric_limits<int>::max()` is used to represent infinity in the distances matrix.

Summary

- The Floyd–Warshall algorithm is a method for determining the shortest paths between all pairs of vertices in a weighted graph.
- The algorithm initializes a distance matrix and iteratively updates it by considering all possible intermediate vertices.
- After completion, the matrix provides the shortest paths between all pairs of vertices.
- While suitable for graphs with negative weights, the algorithm has a time complexity of $O(V^3)$, making it less efficient for sparse graphs compared to other algorithms.

6.18 Traversing a Graph (BFS, DFS)

Traversing a graph involves systematically visiting all the vertices and edges in the graph to access or process information associated with each element.

The traversal process facilitates the exploration of the graph's structure and is crucial for various graph-related algorithms and applications.

Two primary methods for graph traversal are

1. Depth–First Search (DFS)
2. Breadth–First Search (BFS).

1. Depth–First Search (DFS)

In DFS, the traversal explores as far as possible along each branch before backtracking. It starts at a source vertex, visits an adjacent vertex, and continues the exploration deeper

until reaching the end of a branch. The process then backtracks to explore other branches.

- DFS can be implemented using recursion or a stack data structure to keep track of vertices and their exploration order.
- Applications of DFS include topological sorting, cycle detection, and connectivity analysis in graphs.

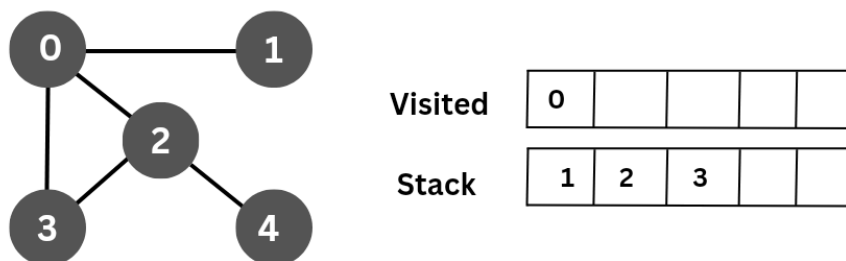
Depth-First Search (DFS) is a graph traversal algorithm that systematically explores the vertices and edges of a graph by prioritizing depth over breadth.

Here's how DFS works:

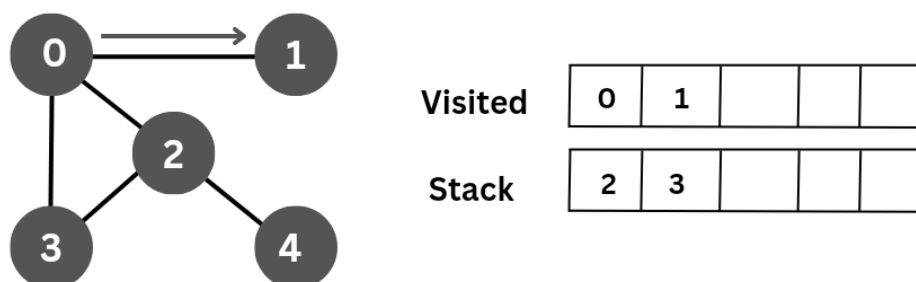
Step 1: Initially stack and visited arrays are empty.



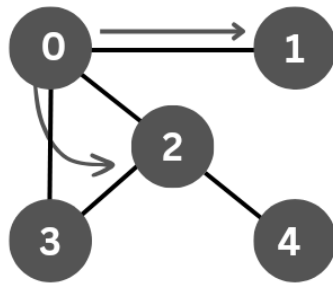
Step 2: Visit 0 and put its adjacent nodes which are not visited yet into the stack.



Step 3: Now, Node 1 is at the top of the stack, so visit node 1 pop it from the stack, and put all of its adjacent nodes that are not visited in the stack.

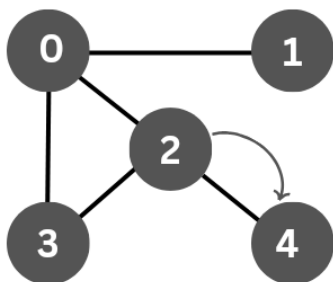


Step 4: Now, Node 2 is at the top of the stack, so visit Node 2 and pop it from the stack, and put all of its adjacent nodes that are not visited (i.e., 3, 4) in the stack.



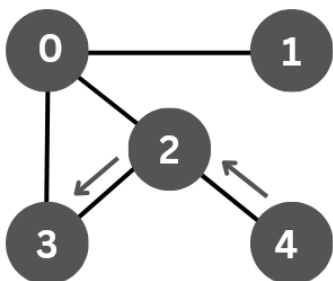
Visited	0	1	2		
Stack	4	3			

Step 5: Now, Node 4 is at the top of the stack, so visit Node 4 pop it from the stack, and put all of its adjacent nodes which are not visited in the stack.



Visited	0	1	2	4	
Stack	3				

Step 6: Now, Node 3 is at the top of the stack, so visit Node 3 pop it from the stack, and put all of its adjacent nodes that are not visited in the stack.



Visited	0	1	2	3	4
Stack					

1. Starting Point

- DFS starts from a chosen source vertex. This vertex is marked as visited.

2. Exploration

The algorithm explores as far as possible along each branch before backtracking. It selects an unvisited neighbor of the current vertex and moves deeper into the graph.

3. Recursion or Stack

DFS can be implemented using recursion or a stack data structure. In the recursive approach, the algorithm recursively calls itself for each unvisited neighbor. In the stack approach, the algorithm uses a stack to keep track of vertices to be visited.

4. Backtracking

If a dead end is reached (a vertex with no unvisited neighbors), the algorithm backtracks to the previous vertex and explores other unvisited branches.

5. Completion

The process continues until all vertices are visited or until the algorithm is manually terminated.

DFS explores the graph by diving as deep as possible, which leads to a thorough exploration of individual branches before moving on to the next. This traversal strategy is particularly useful for tasks such as topological sorting, cycle detection, and connectivity analysis in graphs.

Overall, DFS provides a comprehensive understanding of the structure and relationships within a graph.

Example:

```

#include <iostream>
#include <vector>
#include <stack>

using namespace std;

// Function to perform Depth-First Search on a graph
void dfs(int start, const vector<vector<int>>& graph, vector<bool>& visited) {
    stack<int> s;
    s.push(start);

    while (!s.empty()) {
        int current = s.top();
        s.pop();

        if (!visited[current]) {
            cout << current << " ";
            visited[current] = true;

            // Traverse adjacent nodes
            for (int neighbor : graph[current]) {
                if (!visited[neighbor]) {
                    s.push(neighbor);
                }
            }
        }
    }
}

int main() {
    // Example graph represented as an adjacency list
    vector<vector<int>> graph = {
        {1, 2}, // Node 0 is connected to nodes 1 and 2
        {0, 3, 4}, // Node 1 is connected to nodes 0, 3, and 4
        {0, 5}, // Node 2 is connected to nodes 0 and 5
        {1}, // Node 3 is connected to node 1
        {1}, // Node 4 is connected to node 1
        {2} // Node 5 is connected to node 2
    };

    int startNode = 0; // Starting node for DFS

    // Initialize visited array
    vector<bool> visited(graph.size(), false);

    cout << "DFS starting from node " << startNode << ": ";
    dfs(startNode, graph, visited);

    return 0;
}

```

This program demonstrates a simple DFS implementation using a stack to keep track of nodes to visit. The graph is represented as an adjacency list, and the dfs function performs the DFS traversal starting from a specified node. The visited array is used to keep track of visited nodes to avoid revisiting them.

Summary

- Traversing a graph involves systematically visiting its vertices and edges, and two fundamental methods for this process are Depth-First Search (DFS) and Breadth-First Search (BFS).
- DFS explores the graph by prioritizing depth, starting at a source vertex and traversing as far as possible along each branch before backtracking.
- It can be implemented using recursion or a stack data structure. DFS is valuable for tasks such as topological sorting, cycle detection, and connectivity analysis in graphs.
- The algorithm provides a comprehensive understanding of the graph's structure by delving deeply into individual branches before moving on to others.

2. Breadth-First Search (BFS)

In BFS, the traversal explores vertices level by level, starting from the source vertex. It visits all the neighbors at the current level before moving on to the next level.

- This process continues until all vertices are visited.
- BFS is commonly implemented using a queue data structure, ensuring that vertices are visited in the order they were discovered.
- BFS is useful for finding the shortest path between two vertices, connected component analysis, and network analysis.

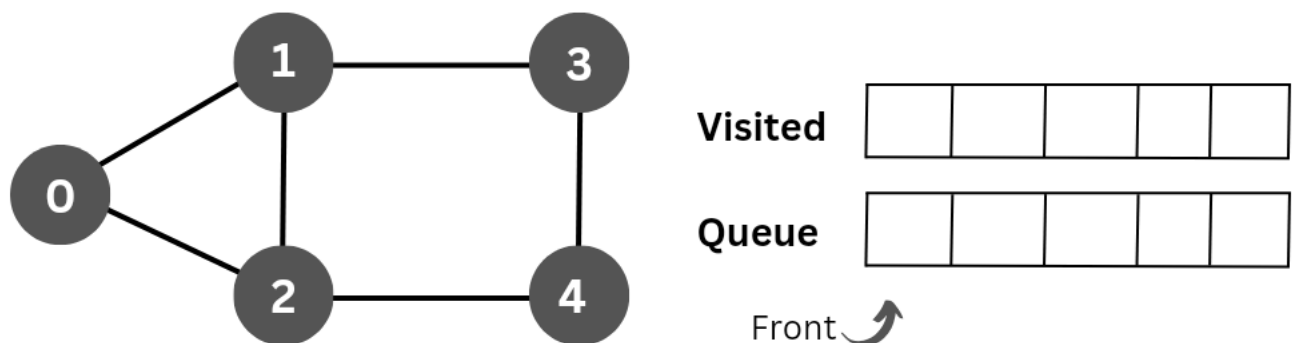
Traversing a graph provides insights into its connectivity, structure, and relationships between vertices.

The choice between DFS and BFS depends on the specific requirements of the application or algorithm. Efficient graph traversal is fundamental for solving various graph-related problems. Summarize this in one paragraph.

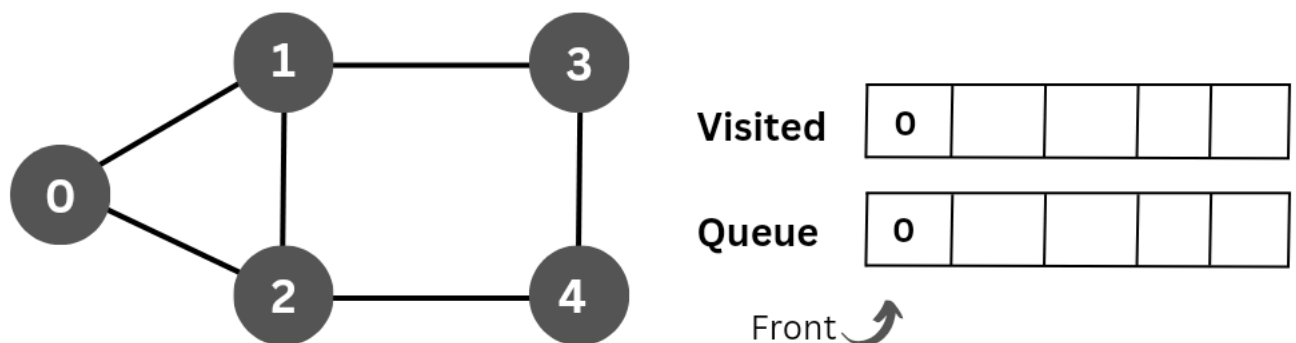
Breadth-First Search (BFS) is a graph traversal algorithm that systematically explores the vertices and edges of a graph in layers, prioritizing breadth over depth

Here's how BFS works:

Step 1: Initially queue and visited arrays are empty.

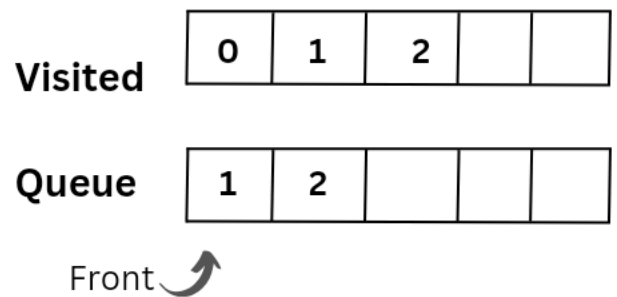
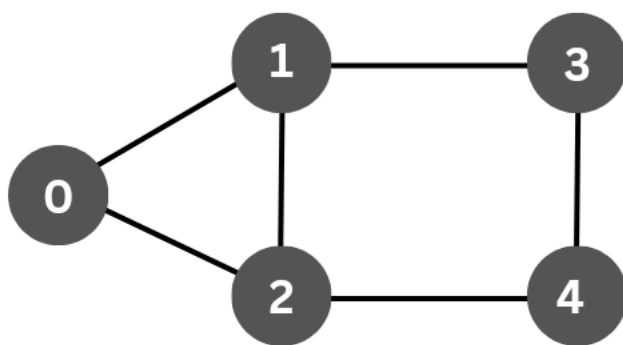


Step 2: Push node 0 into a queue and mark it visited.

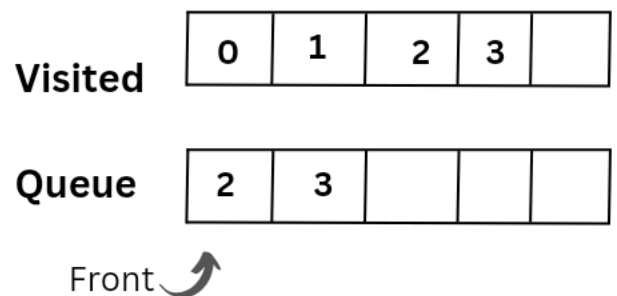
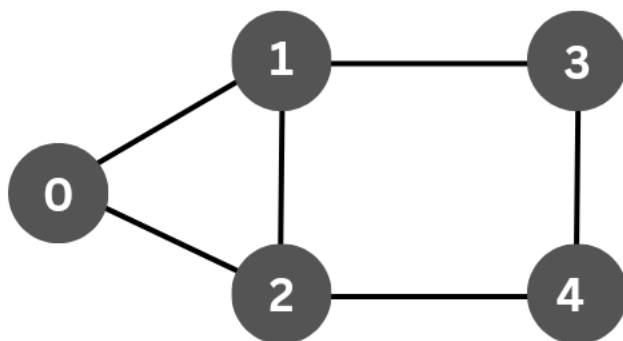


Step 3: Remove node 0 from the front of the queue visit the unvisited neighbors and push

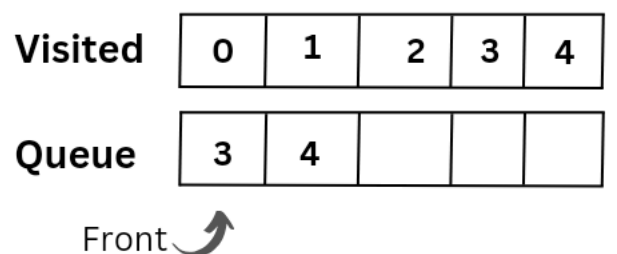
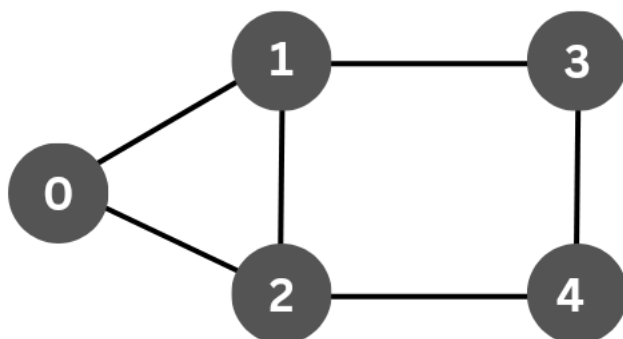
them into the queue.



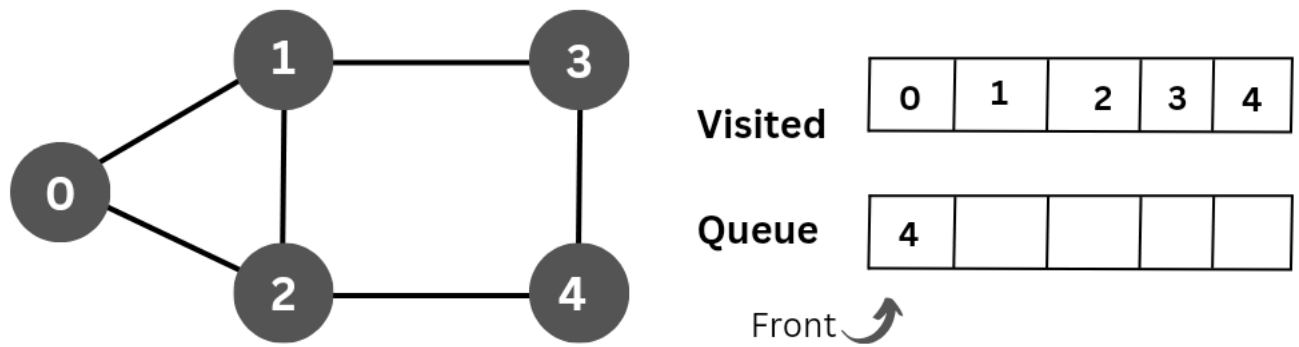
Step 4: Remove node 1 from the front of the queue and visit the unvisited neighbors and push them into the queue.



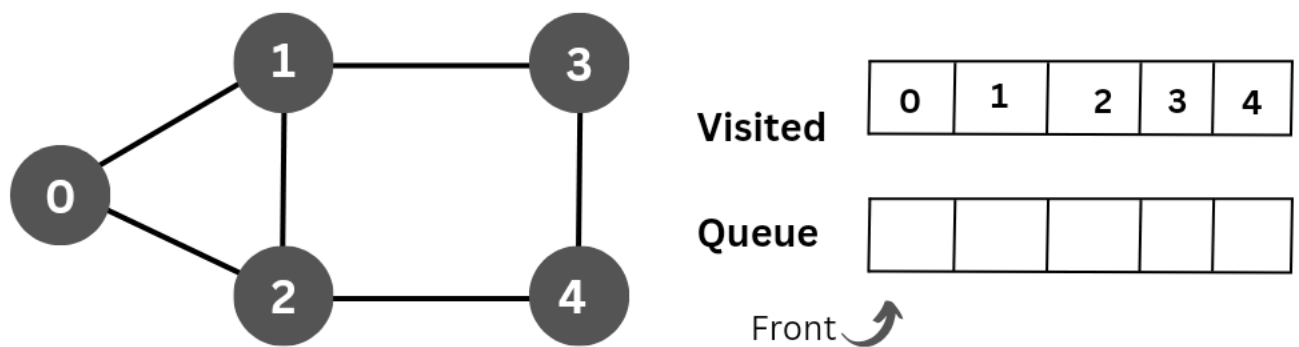
Step 5: Remove node 2 from the front of the queue visit the unvisited neighbors and push them into the queue.



Step 6: Remove node 3 from the front of the queue visit the unvisited neighbors and push them into the queue.



Step 7: Remove node 4 from the front of the queue visit the unvisited neighbors and push them into the queue.



1. Starting Point

BFS starts from a chosen source vertex. This vertex is marked as visited and becomes the starting point of the traversal.

2. Exploration in Layers

BFS explores all the neighbors of the current vertex before moving on to their neighbors. This results in a layer-by-layer exploration, starting from the source vertex.

3. Queue

A queue data structure is typically used to keep track of vertices to be visited. The source vertex is enqueued, and then, while the queue is not empty, vertices are dequeued for exploration.

4. Visit and Enqueue

For each dequeued vertex, BFS visits its unvisited neighbors, marks them as visited, and enqueues them. This ensures that vertices are visited in order of their distance from the source.

5. Completion

The process continues until all vertices are visited or until the algorithm is manually terminated.

BFS systematically explores the graph in layers, ensuring that all vertices at a certain distance from the source are visited before moving on to vertices at a greater distance.

This traversal strategy is valuable for tasks such as finding the shortest path in unweighted graphs, connected component analysis, and level-order tree traversals.

Example:

```
#include <iostream>
#include <vector>
#include <queue>

using namespace std;

// Function to perform Breadth-First Search on a graph
void bfs(int start, const vector<vector<int>>& graph, vector<bool>& visited) {
    queue<int> q;
    q.push(start);

    while (!q.empty()) {
        int current = q.front();
        q.pop();

        if (!visited[current]) {
            cout << current << " ";
            visited[current] = true;

            // Traverse adjacent nodes
            for (int neighbor : graph[current]) {
                if (!visited[neighbor]) {
                    q.push(neighbor);
                }
            }
        }
    }
}

int main() {
    // Example graph represented as an adjacency list
    vector<vector<int>> graph = {
        {1, 2}, // Node 0 is connected to nodes 1 and 2
        {0, 3, 4}, // Node 1 is connected to nodes 0, 3, and 4
        {0, 5}, // Node 2 is connected to nodes 0 and 5
        {1}, // Node 3 is connected to node 1
        {1}, // Node 4 is connected to node 1
        {2} // Node 5 is connected to node 2
    };

    int startNode = 0; // Starting node for BFS

    // Initialize visited array
    vector<bool> visited(graph.size(), false);

    cout << "BFS starting from node " << startNode << ": ";
    bfs(startNode, graph, visited);

    return 0;
}
```

This program demonstrates a simple BFS implementation using a queue to keep track of nodes to visit.

The graph is represented as an adjacency list, and the BFS function performs the BFS traversal starting from a specified node. The visited array is used to keep track of visited nodes to avoid revisiting them.

Summary

- Breadth-First Search (BFS) is a graph traversal algorithm that systematically explores the vertices and edges of a graph in layers, prioritizing breadth over depth.
- It starts from a chosen source vertex, marking it as visited and becoming the starting point of the traversal.
- BFS explores all the neighbors of the current vertex before moving on to their neighbors, resulting in a layer-by-layer exploration.
- A queue data structure is commonly used to manage the vertices to be visited, ensuring a systematic order.
- This traversal strategy is valuable for tasks such as finding the shortest path in unweighted graphs, connected component analysis, and level-order tree traversals.
- Efficient graph traversal is fundamental for solving various graph-related problems, and the choice between DFS and BFS depends on the specific requirements of the application or algorithm.

6.19 Application of Graph

Graphs have a wide range of applications across various domains due to their ability to model relationships and connections.

Here are some key applications of graphs:

1. Social Networks

Graphs represent social connections on platforms like Facebook, LinkedIn, and Twitter. Nodes can represent individuals, and edges represent relationships or connections between them. Analyzing the structure of social networks helps understand influence, trends, and community formations.

2. Transportation Networks

Graphs model transportation systems such as road networks, flight routes, and subway systems. Nodes represent locations, and edges represent connections or routes. Graph algorithms assist in optimizing routes, scheduling, and transportation planning.

3. Internet and Web

The Internet and the World Wide Web are modeled as graphs. Web pages are nodes, and hyperlinks are edges. Search algorithms like Google's PageRank use graph-based structures to rank web pages based on their connectivity and relevance.

4. Recommendation Systems

Graphs are employed in recommendation systems. Nodes can represent users or items, and edges represent interactions or preferences. Analyzing the graph structure enables recommending relevant products, movies, or content.

5. Biological Networks

Graphs model biological systems, including protein-protein interaction networks, gene regulatory networks, and ecological interactions. Understanding these networks aids in studying diseases, genetics, and ecosystems.

6. Circuit Design

Electrical circuits can be represented as graphs, where components are nodes, and connections are edges. Graph algorithms help optimize circuit design, analyze connectivity, and troubleshoot issues.

7. Game Theory

In game theory, graphs model strategic interactions between players. Nodes can represent players, and edges represent possible interactions or strategies. Graph-based models facilitate the analysis of strategic choices and outcomes.

8. Computer Networks

Networks, including the internet, are modeled using graphs. Nodes represent devices (computers, routers), and edges represent connections. Graph algorithms aid in network optimization, routing, and fault tolerance.

9. Supply Chain Management

Graph model supply chain relationships between suppliers, manufacturers, and distributors. Analyzing the supply chain graph helps optimize logistics, reduce costs, and improve efficiency.

10. Graph Databases

Graph databases like Neo4j leverage graph structures to represent and query relational data. They are particularly effective for scenarios where relationships between entities are as important as the entities themselves.

Graphs provide a powerful and versatile framework for solving complex problems in diverse fields, making them a fundamental tool in computer science and beyond.

Summary

- Graphs serve as a versatile and powerful framework across various domains, addressing complex challenges in computer science and beyond.
- In social networks like Facebook and Twitter, graphs represent individuals and their connections, aiding in the analysis of the influence and community formations.
- Transportation systems, internet, and recommendation systems leverage graph models to optimize routes understand web structures and enhance user experiences.
- Biological networks help study protein interactions and ecological systems, while in circuit design, electrical components are nodes connected by edges.

- Game theory, computer networks, supply chain management, and graph databases further showcase the broad applicability of graphs in strategic modeling, network optimization, logistics, and relational data representation.
- The adaptability of graphs makes them an indispensable tool for solving intricate problems and gaining insights in diverse fields.