

1. Define Time complexity. Why is it Important in Algorithm Analysis?

Answer - Time Complexity is a measure used in computer science to describe the amount of time an algorithm takes to execute in relation to the size of its input data. Instead of calculating the exact running time in seconds or milliseconds, time complexity focuses on the growth rate execution time as the input size increases. This growth is usually represented using mathematical functions such as constant, linear, logarithmic, quadratic, or exponential. By studying time complexity, we can understand how efficiently an algorithm performs when handling small as well as very large datasets.

Time complexity is generally expressed using asymptotic notation, especially Big O notation, which represents the worst-case running time of an algorithm. These notations help programmers compare algorithms independently of hardware or programming language, making analysis more universal and reliable.

The importance of time complexity in algorithm analysis is very high. First, it helps developers choose the most efficient algorithm among multiple possible solutions for the same problems, which is crucial in real-world applications such as search engines. Analyzing time complexity, a program might work correctly for small inputs but become extremely slow or unusable when the data size grows.

Another key reason time complexity is important is that it guides optimization and improvement. When programmers know the time complexity of an algorithm, they can identify performance bottlenecks and replace inefficient approaches with better ones. Thus, time complexity plays a major role in writing high-

performance and cost-effective software.

In conclusion, time complexity is a fundamental concept that measures how the execution time of an algorithm grows with input size. It is essential for comparing algorithms, ensuring scalability, improving efficiency, and designing reliable software system. Understanding time complexity enables developers to create programs that are not correct but also fast, efficient and suitable for real-world large scale computing environments.

Q:-2. What is Big O Notation? Provide an Example.

Answer:- Big O Notation is a mathematical representation used in computer science to describe the upper bound of an algorithm's running time or space requirement as the size of the input increases. It expresses how the performance of an algorithm grows rather than measuring the exact execution time in seconds. Because actual running time depends on factors like hardware speed, compiler optimization and programming language. Big O Notation provides a standard, machine-independent way to compare algorithm based on their efficiency.

Big O mainly focuses on the worst-case scenario, which means the maximum time an algorithm could take for a given input size. Since, they become less significant as the input size grows very large. For example an expression like $3n+5$ is simplified to $O(n)$ because the linear term n dominates the growth rate.

- There are several common Big O complexities used to classify algorithm.
- $O(1)$ - Constant Time:- Execution time does not change with input size, such as accessing an element in an array by index.
- $O(\log n)$ - Logarithmic Time:- Time increases slowly as input grows, as seen in binary search.
- $O(n)$ - Linear Time:- Execution time grows directly with input size, such as linear search.

- $O(n \log n)$:- Found in efficient sorting algorithm like merge sort and heap sort.
- $O(n^2)$:- Quadratic time:- occurs in simple sorting algorithm like bubble sort or selection sort, nested loops are used.
- $O(2^n)$ or factorial time:- Extremely slow growth, usually seen in recursive brute-force algorithms.

Big O notation is extremely important in algorithm design and analysis because it helps programmers compare different solutions, predicts scalability and choose efficient approaches for real-world applications such as databases, networking, artificial intelligence, and large-scale data processing. Without Big O analysis programs that appears fast for small inputs may become unusable for large datasets.

In conclusion, Big O notation is a fundamental tool that describes the growth rate of an algorithm's complexity in the worst case. It enables developers to design efficient, scalable, and high performance software system by focusing on how algorithms behave as input size become very large.

3. What is the Difference Between a Stack and a Queue

Answer:- A stack and a queue are two fundamental linear data structures used in computer science to store and manage data in an organized manner. Both structures allow insertion and deletion of elements. Understanding the difference between a stack and a queue is essential for designing efficient algorithms and software system.

A stack follows the principle of LIFO (Last In First Out). This means that the last element inserted into the stack is the first one to be removed. The main operations in a stack are push (to insert an element), pop (to remove the top element), and peek/top (to view the top element without removing it). All insertions and deletions occur at one end only, called the top of the stack.

Stacks are commonly used in situations that require reversal of order or tracking of function calls, such as undo/redo operations in text editors, expression evaluation, syntax parsing, and recursion handling in programming language.

In contrast, a queue follows the principle of FIFO (First In, First Out). This means the first element inserted into the queue is the first one to be removed. A queue has two ends; the rear (where insertion takes place, called enqueue) and the front (where deletion takes place, called dequeue). (4)

The key differences between a stack and a queue lie in their processing order and access points. A stack allows access to only the most recently added element, while a queue processes elements in the same order they were inserted. This difference makes stack suitable for backtracking and recursive process, whereas queues are better for sequential processing and scheduling tasks.

In conclusion, both stacks and queues are essential data structures that support efficient data handling in algorithms and systems.

A stack operates on the LIFO principle and is useful for universal and nested operations, while a queue operates on the FIFO principle and is ideal for scheduling and orderly processing. Making their understanding crucial in algorithm design and software development.

Q:-4 what is a Priority Queue and How is it Implemented ?

Answer :- A priority queue is a special type of abstract data structure in which each element is associated with a priority value, & elements are removed from the queue based on their priority order rather than the order in which they were inserted. This makes priority queues extremely useful in applications where importance or urgency determines processing order rather than arrival time.

The basic operations of a priority queue include insertion of an element with a priority, deletion of the highest priority element. Efficient handling of these operations is essential for maintaining good performance and scalability, especially when dealing with large datasets or real-time systems.

There are several ways to implement a priority queue. One simple approach is using an unsorted array or linked list. Another method is using a sorted array or sorted linked list, where deletion becomes fast $O(1)$ but insertion requires shifting or traversal giving $O(n)$ time complexity.

The most efficient and commonly used implementation of a priority queue is the binary heap, which is a complete binary tree that satisfies the heap property. Using binary heap, both insertion and deletion operations can be performed in $O(\log n)$ time. Binary heaps are usually implemented using arrays for memory efficiency and fast index calculations. ⑤

Priority queues are widely used in real-world computing problems. For example, they are used in CPU scheduling, where processes with higher priority must execute first; in Dijkstra's shortest path algorithm and Prim's minimum spanning trees algorithm in graph theory; Their ability to quickly access the most critical element makes them essential in many high-performance systems.

In conclusion, a priority queue is an advanced queue structure that processes elements based on priority instead of arrival order. Priority queues play a crucial role in algorithm design, scheduling, graph processing, and real-time computing, making them an important concept in data structures and computer science.

5) Explain the Concept of Dynamic Memory Allocation with an Example.

Answer - Dynamic memory allocation is a method in computer programming through which memory is allocated during the runtime of a program instead of at compile time. This provides greater flexibility, efficient memory utilization and scalability, especially in application where the amount of data is not known beforehand.

Dynamic memory allocation is typically performed in the heap memory area, which is managed by the operating system or runtime environment. Programmers use specific functions or operators to allocate and deallocate memory.

For example, in the C programming language, functions such as `malloc()`, `calloc()`, `realloc()` and `free()` are used.

⑥

- `malloc()` allocates a specified number of bytes.
- `calloc()` allocates memory for multiple elements and initializes them to zero.
- `realloc()` resizes previously allocated memory.
- `free()` release the allocated memory back to the system.

An example in C is shown below:

```
#include <stdlib.h>
```

```
int *ptr;
```

```
ptr = (int *) malloc (5 * sizeof(int)); // allocate memory for 5 integers
```

In this example, memory for five integer values is allocated during program execution. The pointer `ptr` stores the address of the allocated memory. After the memory is no longer needed, it should be released using;

```
free (ptr);
```

Releasing memory is important to prevent memory leaks, which occurs when allocated memory is not returned to the system. Memory leaks can reduce program performance and system stability, especially in long-running applications.

The major advantages of dynamic memory allocation include efficient use of memory, ability to handle variable data size, and improved program flexibility. However it also requires careful management, because improper use may cause memory leaks, dangling pointers, or fragmentations.

In conclusion, dynamic memory allocation is a powerful technique that allows programs to allocate and manage memory at runtime according to their needs. making it an essential concept in data structures, system programming and modern application development.

6. What is the time complexity of Inserting an Element at the Beginning of a Singly Linked List?

Answer:- In a singly linked list, elements are stored in the form of nodes, where each node contains two parts; the data field and a pointer (reference) to the next node in the sequence. Unlike arrays, linked lists do not store elements in contiguous memory locations, which allows dynamic memory usage and flexible insertion or deletion of elements without shifting existing data. Because of this structure, inserting an element at different positions in a linked list may have different time complexities depending on the amount of traversal required.

When inserting an element at the beginning (head) of a singly linked list, the process is very simple and different.

First a new node is created & memory is allocated for it. Second the next pointer of new node is made to point to the current head node of the list.

Finally, the head pointer is updated to point to this new node. These steps involve only a constant number of operations and no traversal of the list is required, regardless of how many elements are already present.

Because the number of operations does not depend on the size of the list (n), the time complexity of inserting an element at the beginning of a singly linked list is $O(1)$. This efficiency is one of the major advantages of linked lists over arrays, where inserting at the beginning requires shifting all existing elements, resulting in $O(n)$ time complexity.

However, while insertion at the beginning is very efficient, linked lists also have some limitations. Accessing an arbitrary element requires sequential traversal, which takes $O(n)$ time. Despite these drawbacks, the constant-time insertion at the beginning remains a key reason why singly linked lists are widely used in data structure design and system programming.

In conclusion, inserting an element at the beginning of a singly linked list requires only pointer adjustment and node creation without any traversal of existing elements. Therefore its time complexity is $O(1)$, making it an extremely efficient operation for dynamic data management and an important concept in algorithm analysis and data structures.

Q:-7 Demonstrate AVL Rotations when the Nodes with keys 10, 20, 32, 35, 25, 27 are inserted in an AVL Tree.

Answer:- An AVL tree is a type of self-balancing binary search tree (BST) in which the difference between the heights of the left and right subtrees of any node, known as the balance factor, is always maintained between -1, 0 and +1. Whenever an insertion causes the tree to become unbalanced, rotations are applied to restore balance.

Let us insert the keys 10, 20, 32, 35, 25 and 27 one by one and observe the required AVL rotations.

Step-1 Insert 10

The tree contains only one node, so it is perfectly balanced. No rotation is needed.

Step-2 Insert 20

Node 20 becomes the right child of 10.

Balance factor of 10 = -1, which is within the allowed range.

The tree remains balanced, so no rotation is required.

Step 3: Insert 32.

Node 32 becomes the right child of 20, forming a right-right (RR) imbalance at node 10 because the height difference becomes -2.

To fix this, a single left rotation is performed at node 10.

After rotation, 20 becomes the root, with 10 as its left child and 32 as its right child.

The tree is now balanced again.

Step-4 Insert 35

9

Node 35 is inserted as the right child of 32.
The balance factors of all nodes remains within -1 to $+1$, so no rotation is needed.

Step-5 Insert 25

Node 25 is inserted as the left child of 32.
This creates a right-left (RL) imbalance at node 20 because the right subtree is heavier, but the new node is the left subtree of the right child.

To resolve this, a double rotation is applied.

1. Right rotation at node 32
2. Left rotation at node 20

After these rotations the AVL tree becomes balanced again.

Step-6 Insert 27

Node 27 is inserted in the subtree such that a left-right (LR) imbalance occurs at node 32 (or the affected subtree root).

To fix this, another double-rotation is performed.

1. Left rotation on the left child
2. Right rotation on the unbalanced node.

After performing these rotations, the final AVL tree satisfies the balance factor condition at every nodes.

In conclusion, inserting the keys 10, 20, 32, 35, 25, 27 in an AVL tree demonstrates how RR, RL and LR rotations are applied to maintain balance after each insertion. This self-balancing property ensures that AVL trees remain efficient, stable and suitable for high-performance searching applications in computer science.

8) What is the Maximum Possible Number of Nodes in a Binary Tree at Level 6? (10)

Answer A binary tree is a hierarchical data structure in which each node can have at most two children, commonly referred to as the left child and the right child. Binary trees are widely used in data. One important theoretical property of binary trees is determining the maximum number of nodes present at a particular level, which helps in analysing tree height, storage requirements and algorithm efficiency.

In a binary tree, the number of nodes that can exist at each level follows a geometric progression.

- At level 0 (root level), the maximum number of nodes is 1.
- At level 1, the maximum number of nodes becomes 2.
- At level 2, the maximum number increases to 4.

The pattern continues such that the maximum number of nodes at level L of a binary tree is given by the formula.

$$\text{Maximum nodes at level } L = 2^L \quad \text{Maximum nodes at level } L = 2^L$$

Understanding this property is important in algorithm design and analysis, particularly in operations involving tree traversal, searching and storage estimation. The exponential growth of nodes per level also highlights why tree-based algorithms scale efficiently compared to linear structures.

In conclusion, the maximum number of nodes at any level L of a binary tree is 2^L . This mathematical property plays a crucial role in understanding tree structure, efficiency, and performance in computer science and data structure applications.

In conclusion a hash function is a crucial tool in computer science that maps keys to fixed-size indices for fast data storage and retrieval. It is widely used because it provides efficient average-case performance, scalability and practical usefulness in databases, security systems, and modern software applications. Understanding hash functions is fundamental for designing high-performance data structures and algorithms.

10) Distinguish Between Depth-First Search (DFS) and Breadth-First Search (BFS)

Answer:- Depth-First Search (DFS) and Breadth-First Search (BFS) are two fundamental graph traversal algorithms used to explore the nodes and edges of a graph or tree data structure. Understanding the distinction b/w DFS and BFS is essential for solving problem related to path finding, connectivity, searching and network analysis in computer science.

Depth-First Search (DFS) explores a graph by going as deep as possible along one branch before backtracking. It follows a recursive or stack-based approach, where the most recently discovered node is processed first. DFS is closely related to the stack data structure and is often implemented using recursion.

In contrast Breadth-First Search (BFS) explores the graph level by level, visiting all neighbors of the current node before moving to the next depth level. BFS uses queue data structure, the first time a node is reached is through the shortest path from the source in an unweighted graph.

Q. What is a Hash Function? Why is it used? (12)

Answer: A hash function is a mathematical or computational function that converts a given input value (key), which is then used to store or retrieve data efficiently in a hash table. Because of this efficiency, hashing is widely used in data structures, databases, cryptography, caching systems, and indexing mechanisms.

When data is stored using a hash table, the hash function processes the key and generates a corresponding array index where the value should be placed. For example, a simple hash function may compute:

$$\text{index} = \text{key} \bmod \text{table size} \quad \text{index} = \text{key} \bmod \text{table size}$$

This calculation ensures that the key is mapped to a valid position inside the table. This makes hashing significantly faster than linear or binary search in many practical situations.

A good hash function must satisfy several important properties.

It should be fast to compute so that lookup operation remains efficient.

It should distribute keys uniformly across the hash table to avoid clustering.

Hash functions are extremely important in many real-world applications. They are used in database indexing to speed up record retrieval, in modern programming language also use hashing internally for data structures like Hash Map, Hash Set, and dictionaries, which provide near constant-time performance for common operations.

Despite their advantages, hash tables may experience performance degradation if the load factor becomes too high or if the hash function distributes keys poorly. Therefore, choosing an efficient hash function and appropriate table size is essential for maintaining optimal performance.

There are several important differences between DFS and BFS. (13)
DFS follows a deep traversal strategy and uses a stack or recursion while BFS follows a level-order traversal using a queue.

In terms of time complexity, both DFS and BFS typically run in $O(V+E)$, where V is the number of vertices and E is the number of edges. The main distinction lies in memory consumption and traversal order rather than execution time.

In conclusion, DFS and BFS are essential graph traversal techniques with different strengths and use cases. DFS is efficient for deep exploration, recursion-based problems and structural analysis, choosing the appropriate algorithm depends on the nature of the problem, memory constraints and desired outcome, making both DFS and BFS fundamental tools in data structures and algorithms design.