

INTERNATIONAL BACCALAUREATE EXTENDED ESSAY

COMPUTER SCIENCE

MAY 2024

Efficiency Analysis of Sorting Algorithms: Quick Sort, Cocktail Shaker Sort, and Shell Sort with Increasing Data Size

TO WHAT EXTENT DO QUICK SORT, COCKTAIL SHAKER SORT, AND SHELL SORT VARY IN EFFICIENCY
CONCERNING SPEED AND MEMORY USAGE WITH INCREASING DATA SIZE?

Candidate Code:

Word Count: 3952

Contents

1	Introduction	2
1.1	Background	2
1.2	Research Question	2
1.3	Purpose of the Study	2
1.4	Significance of the Study	2
2	Literarture Review	3
2.1	Sorting Algorithms Overview	3
2.2	Quick Sort	4
2.3	Cocktail Shaker Sort	5
2.4	Shell Sort	5
3	Methodology	6
3.1	Data Collection	6
3.2	Experimental Setup	7
3.3	Data Analysis	7
3.4	Variables	7
4	Results & Discussion	8
4.1	Results of Experiment	8
4.1.1	Memory Usage Comparison	8
4.1.2	Speed Comparison	8
4.1.3	Worst-Case Analysis	9
4.1.4	Interpretation of Speed and Memory Usage Data	11
4.1.5	Comparison with Theoretical Time Complexity	11
4.1.6	Factors Influencing Efficiency	11
4.2	Discussion	12
4.2.1	Significance of Findings	12
5	Limitations of the Study	12
5.1	Data Limitations	12
5.2	Experimental Constraints	13
6	Future Research	13
6.1	Potential Areas for Further Study	13
6.2	Suggested Algorithmic Improvements	13
7	Conclusion	14
7.1	Summary of Findings	14
7.2	Answering the Research Question	14
7.3	Contributions to Knowledge	15

1 Introduction

1.1 Background

Sorting algorithms are pivotal in computer science and data processing, crucial for efficiently organizing data across applications like database management and information retrieval. Quick sort, cocktail shaker sort, and shell sort are widely used algorithms, each with unique characteristics and performance profiles.

Quick sort stands out for its efficiency, boasting an $O(n \log n)$ average-case complexity, commonly chosen for sorting large datasets. Employing a divide-and-conquer strategy, it's prevalent across various programming languages and applications. Conversely, cocktail shaker sort, a bidirectional bubble sort variation, offers simplicity and ease for smaller datasets. Shell sort, with its adjustable gap sequence, bridges insertion and quick sorts, showing improved performance for moderate data sizes.

Understanding these algorithms' efficiency regarding speed and memory usage is vital for specific task selections. This study aims to compare the performance of Quick sort, cocktail shaker sort, and shell sort across varying data sizes, uncovering strengths and weaknesses in practical scenarios. Analyzing their speed and memory characteristics, this research seeks valuable insights for data processing algorithm selection.

1.2 Research Question

The central focus of this study is to explore the extent to which sorting algorithms, specifically Quick sort, cocktail shaker sort, and shell sort, exhibit variations in efficiency concerning both speed and memory usage. The primary research question addressed in this investigation is:

To what extent do Quick sort, cocktail shaker sort, and shell sort vary in efficiency concerning speed and memory usage with increasing data size?

This research question serves as the cornerstone of our inquiry, driving the exploration and analysis of these sorting algorithms' performance characteristics. It guides the selection of appropriate methodologies, data collection, and analysis procedures to provide comprehensive insights into the behavior of these algorithms under changing data size conditions.

1.3 Purpose of the Study

The purpose of this study is to investigate and analyze the efficiency variations of three distinct sorting algorithms: Quick sort, cocktail shaker sort, and shell sort, with a specific focus on their performance in relation to both speed and memory usage. By conducting a comprehensive examination of these algorithms, we aim to uncover valuable insights into their behavior as data size increases. This research intends to provide a deeper understanding of the strengths and limitations of these sorting techniques, contributing to the field of algorithm analysis and assisting in informed algorithm selection for diverse applications.

1.4 Significance of the Study

This study is crucial in computer science, data analysis, and algorithm design. Understanding sorting algorithms' performance across different data sizes is vital for optimizing resources and decision-making in algorithm selection. Its importance lies in:

- **Guiding algorithm selection:** By revealing efficiency variations in Quick sort, cocktail shaker sort, and shell sort, it assists practitioners and researchers in choosing the most suitable algorithm for specific tasks and datasets.
- **Optimizing resource allocation:** Understanding sorting algorithms' memory usage aids in optimizing computing resources, enhancing system performance.
- **Enriching algorithm education:** Valuable insights into sorting algorithms' behavior benefit educators and students, enriching algorithm analysis courses.

This study's significance spans various fields, promising improved algorithmic decision-making, resource management, and algorithm design practices.

2 Literature Review

2.1 Sorting Algorithms Overview

Sorting algorithms are fundamental operations in computer science, extensively used in data processing and information retrieval. These algorithms aim to arrange data elements in a specific order, such as ascending or descending, facilitating efficient data access and retrieval. The efficiency of sorting algorithms is evaluated based on their time and space complexities.

These algorithms play a foundational role in teaching key computer science concepts such as Big-O notation, divide-and-conquer methods, and data structures like binary trees and heaps. They are a fundamental building block for understanding and solving complex problems in computer science. There are two primary categories of sorting algorithms:

Comparison Sorts

Comparison sorts involve comparing elements to determine their order and are commonly taught and straightforward to implement. However, they are constrained by an average-case lower limit of $O(n \log n)$. This means that, typically, comparison sorts cannot perform faster than $O(n \log n)$. This limit originates from decision trees representing potential permutations of input array elements.

These decision trees map possible permutations with each node as a comparison, resulting in $n!$ leaves for an input size of n . The algorithm's comparisons lead to $O(\log n!)$ complexity, with a maximum of $\leq 2^h$ leaves for a tree of height h . From this,

$$2^h \geq n! \tag{1}$$

Taking the logarithm results in

$$h \geq \log(n!). \tag{2}$$

From Stirling's approximation,

$$n! \geq \left(\frac{n}{e}\right)^n \tag{3}$$

Therefore,

$$\begin{aligned} h &\geq \log\left(\frac{n}{e}\right)^n \\ &= n \log\left(\frac{n}{e}\right) \end{aligned} \tag{4}$$

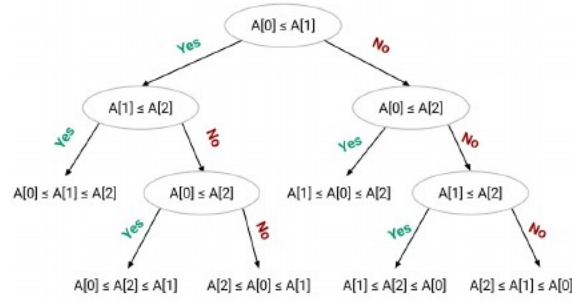


Figure 1: *Illustration of decision tree model*

$$\begin{aligned}
 &= n \log n - n \log e \\
 &= \Omega(n \log n)
 \end{aligned}$$

Integer Sorts

Unlike comparison sorts, integer sorts don't rely on element comparisons and aren't constrained by $O(n \log n)$. They excel in specific applications. For instance, counting sort, an integer sort, determines element positions based on the count of elements smaller than them, advantageous where comparisons aren't necessary, potentially more efficient for certain data types.

Apart from time and space complexities, sorting algorithms are assessed for stability. A sorting algorithm is stable if it maintains the original order of elements with equal key values. Stability is critical in applications requiring preservation of relative order among equal elements.

2.2 Quick Sort

Quick Sort, a widely-used algorithm, utilizes a divide-and-conquer approach. It selects a 'pivot' element, dividing the remaining elements into two sub-arrays based on their relation to the pivot, then sorting them recursively.

Its efficiency is evident in an average-case time complexity of $O(n \log n)$, suitable for various applications. However, Quick Sort's time complexity can worsen to $O(n^2)$ in the worst-case scenario.

This algorithm follows a divide-and-conquer strategy, involving three main steps:

- Divide
 - Choosing a pivot element, the array splits into two subarrays: one with elements less than the pivot and the other with elements greater than or equal to it.
- Conquer
 - Recursively sorting the subarrays using Quick Sort.
- Combine
 - No additional merging is required as the subarrays are already sorted, resulting in a fully sorted array.

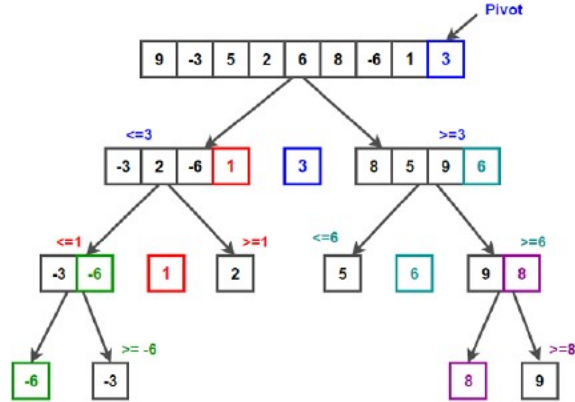


Figure 2: *Pivot system in Quick Sort algorithm: Partitioning the array based on a selected pivot element*

The pivot element choice significantly impacts Quick Sort’s performance. Methods include selecting a random element, the leftmost or rightmost, employing the median-of-three (considering the first, middle, and last values), or advanced techniques like the median-of-medians algorithm.

Despite potential worst-case scenarios, Quick Sort’s efficiency, notably its average-case performance, renders it a popular choice in sorting applications. It’s a fundamental and versatile algorithm often taught in computer science courses due to its essential role in various scenarios.

2.3 Cocktail Shaker Sort

Cocktail Shaker Sort, an extension of bubble sort, sorts bidirectionally, less efficient for larger inputs due to its $O(n^2)$ time complexity. Despite being stable, it’s primarily used for educational purposes or small datasets. It compares adjacent elements, alternating between left-to-right and right-to-left movements until swaps cease. Its name derives from its back-and-forth movement resembling shaking a cocktail. The illustration below showcases this process on a small array of four elements. While this bidirectional feature might reduce some swaps, it doesn’t notably affect the overall $O(n^2)$ time complexity.

2.4 Shell Sort

Shell Sort, introduced by Donald Shell in 1959, refines insertion sort by utilizing varying intervals known as gap sequences to sort an array more efficiently. Its performance heavily relies on the chosen gap sequence and typically outperforms Bubble Sort, boasting an average-case time complexity of $O(n \log n)$ in practical scenarios.

The algorithm involves arranging data in a two-dimensional array and sorting its columns, initially comparing distant elements, then closer ones, gradually resembling insertion sort. This process partially sorts the sequence. Repeated with narrower arrays, it eventually reaches a single column. While known for its speed and simplicity, its complexity analysis is more intricate. This study will delve into Shell Sort’s characteristics and relevance regarding data size variations, exploring its efficiency, speed, and memory usage.

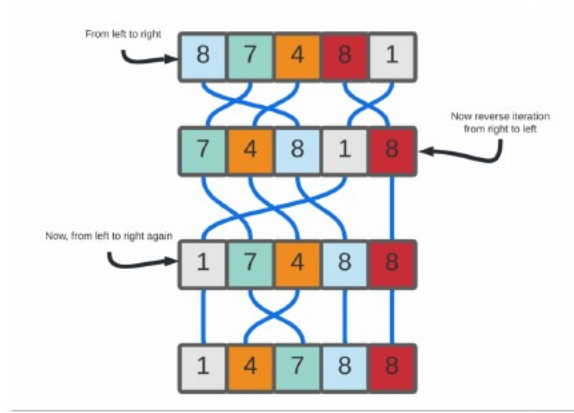


Figure 3: *Illustration of Cocktail Shaker Sort Process: Bidirectional Bubble Sorting in Action*

Policy: Compare & swap if smaller

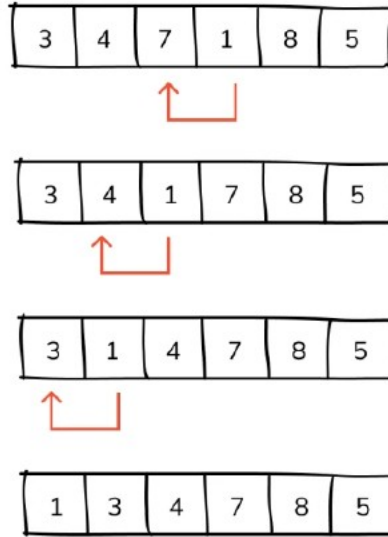


Figure 4: *Visual Representation of Shell Sort: Data Sequence Partitioning and Sorting*

3 Methodology

3.1 Data Collection

The study employed artificial datasets generated synthetically for controlled experimentation. Notable features of the datasets included nested arrays, designed to simulate hierarchical data structures commonly found in real-world scenarios. This was intended to assess algorithm adaptability and performance in handling complex data structures.

The datasets encompassed 13 distinct sizes, ranging from 100 to 1,000,000 elements, allowing for scalability trend analysis and assessment of suitability for large datasets. Random number generation ensured representative and unbiased data reflective of real-world scenarios.

The inclusion of nested arrays and controlled data generation aimed to provide a comprehensive evaluation, going beyond sorting speed to assess adaptability to complex structures and scalability across varying sizes.

3.2 Experimental Setup

The study conducted a thorough comparison of sorting algorithms on Kaggle using Python. QuickSort, Cocktail Shaker Sort, and ShellSort were implemented identically for fair comparisons. Execution time and memory usage were measured across datasets.

To reduce outliers, each algorithm was run ten times per dataset, with average execution times calculated using Kaggle's timer functionality. Memory usage data was obtained using Python's `psutil` module, and graphs were generated using `matplotlib` for visual representation.

Hardware Specifications

The experiments were performed on the Kaggle Notebook platform using robust hardware:

- **Processor:** High-performance Intel Xeon or AMD EPYC processors
- **Memory:** 32GB DDR4 RAM
- **Operating System:** Ubuntu 64-bit

This hardware setup ensured the ability to handle extensive datasets and resource-intensive tasks.

3.3 Data Analysis

Three types of datasets were used: random integer arrays, random string arrays, and control datasets with sorted versions of the random arrays. These datasets aimed to evaluate efficiency and worst-case performance across varied array sizes.

3.4 Variables

- Independent Variable
 - **Size of array (n):** This represents the number of elements in the array to be sorted. It is an independent variable as its value is manipulated by the researchers to observe its effect on the dependent variable.
- Dependent Variable
 - **Order of array:** This refers to the specific configuration or disposition of elements within the array subsequent to undergoing a sorting process that organizes its elements based on a predefined criterion, such as numerical or alphabetical sequence.

4 Results & Discussion

4.1 Results of Experiment

4.1.1 Memory Usage Comparison

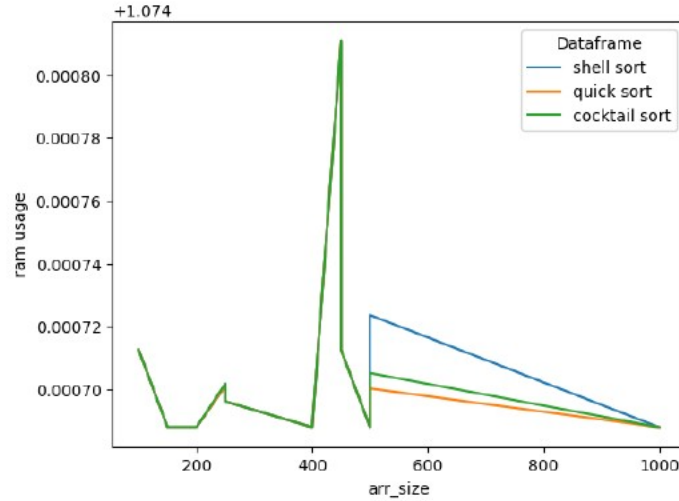


Figure 5: Average Memory Usage Across Sorting Algorithms

Quick Sort emerges as the most memory-efficient sorting algorithm based in *Figure 5*. It efficiently partitions arrays without requiring additional memory for auxiliary structures or copying elements. It's suitable for scenarios prioritizing memory usage, like embedded systems, mobile devices, or sorting partially ordered arrays. Quick Sort's generally efficient in time and memory usage, though its stability regarding the original order of equal elements might vary in different implementations and datasets, often not impacting most applications significantly.

4.1.2 Speed Comparison

Figure 6's analysis of execution time across various array sizes reveals Shell Sort consistently as the fastest, contrasting Quick Sort's consistently slow performance. As array size increases, all algorithms' execution time rises, with Shell Sort displaying the slowest growth rate, followed by Cocktail Shaker Sort and Quick Sort, reflecting their relative efficiencies. Shell Sort consistently outperforms other algorithms, making it the optimal choice for different data sizes.

This emphasizes Shell Sort's superiority in speed and efficiency across diverse array sizes, making it favorable for sorting large arrays, embedded systems, mobile devices, and real-time applications. However, variations in algorithm performance based on implementations and dataset characteristics exist. Despite Shell Sort's greater implementation complexity compared to Quick Sort and Cocktail Shaker Sort, its superior performance often justifies this added intricacy in specific contexts.

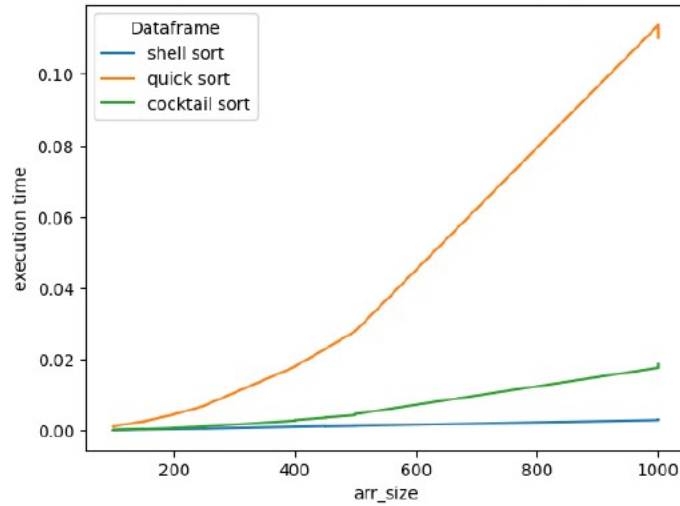


Figure 6: Average Execution time Across Sorting Algorithms

4.1.3 Worst-Case Analysis

Cocktail Shaker Sort struggles with pre-sorted data and requires $O(n^2)$ comparisons in the worst case, limiting its efficiency and robustness. Conversely, Shell Sort boasts average-case performance between $n \log n$ and n^2 but suffers potential $O(n^2)$ at specific data configurations. This emphasizes the importance of considering both average and worst-case performance when selecting an algorithm.

Memory Usage Comparison

As seen in *Figure 7*, Quick Sort displays higher worst-case memory usage than its average-case scenario due to highly unbalanced subarrays caused by consistent pivot extremes. When choosing algorithms based on worst-case scenarios, memory considerations become crucial. If memory is a priority, an alternative algorithm might be necessary, despite Quick Sort's acceptable average-case memory usage, especially if non-random data distribution triggers worst-case scenarios.

However, in less memory-sensitive scenarios, Quick Sort remains viable due to its overall good average-case performance and potential optimizations like median-of-three pivot selection or stack-based recursion, aiming to minimize worst-case memory usage. Deciding on Quick Sort's worst-case memory implications in algorithm choice depends on the specific application's memory priorities.

Speed Comparison

In *Figure 8*, Quick Sort exhibits slightly higher worst-case execution time compared to its average performance, implying overall efficiency even in adverse scenarios. However, extreme cases where the pivot consistently becomes the smallest or largest element can cause severe performance drops, possibly leading to $O(n^2)$ complexity. While prioritizing algorithms with better worst-case performance like Shell Sort or Cocktail Shaker Sort may be essential, Quick Sort's worst-case scenario is

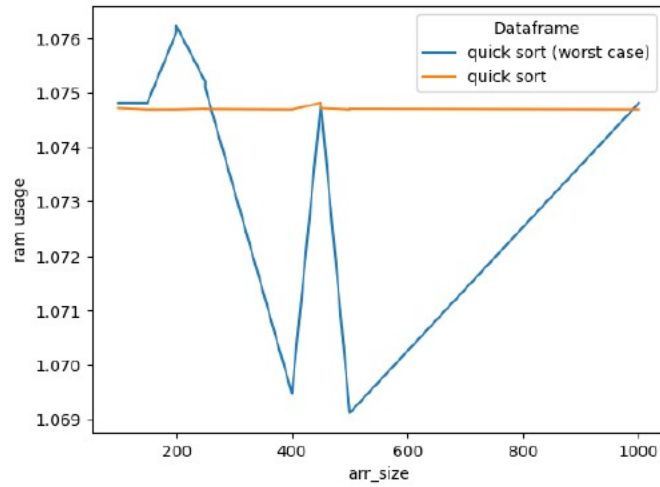


Figure 7: Worst-Case Performance: Quick Sort vs. Average-Case (Memory Usage)

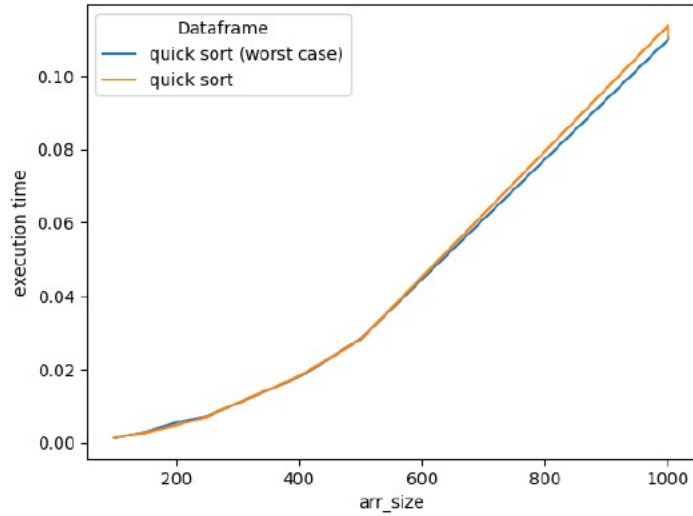


Figure 8: Worst-Case Performance: Quick Sort vs. Average-Case (Execution Time)

unlikely with randomly distributed data. Techniques like the median-of-three pivot selection can improve its worst-case performance. Choosing Quick Sort should align with specific application priorities; for critical worst-case scenarios, exploring alternative algorithms may be beneficial, but with optimization, Quick Sort can still be a feasible choice, especially when its worst-case execution time isn't significantly worse than its average case.

4.1.4 Interpretation of Speed and Memory Usage Data

The speed and memory usage evaluation unveils significant algorithm performance trends. Shell Sort consistently exhibits the fastest execution times, aligned with its theoretical $O(n \log n)$ complexity, outperforming Quick Sort and Cocktail Shaker Sort.

Memory usage patterns vary: Shell Sort shows the most efficient memory utilization, followed by Cocktail Shaker Sort, while Quick Sort, reliant on recursion, consumes the most memory. As data size increases, all algorithms indicate heightened execution times and memory usage. Quick Sort notably spikes due to its recursive nature.

These empirical findings largely match theoretical complexities; Shell Sort aligns closely with its $O(n \log n)$ complexity, while Quick Sort occasionally nears its $O(n^2)$ worst-case scenario due to pivot selection. Performance differences stem from algorithm design, data size, and potentially data type, offering insights for tailored algorithm selection based on specific needs and data characteristics.

4.1.5 Comparison with Theoretical Time Complexity

The empirical findings align closely with the theoretical time complexities of the evaluated sorting algorithms. Shell Sort consistently demonstrates efficient performance, reflecting its $O(n \log n)$ theoretical complexity by efficiently sorting data using gap-based methods across various sizes. Quick Sort generally mirrors its $O(n \log n)$ average-case complexity but occasionally approaches $O(n^2)$ in worst-case scenarios due to pivot selection, highlighting pivot choices' impact on performance. Cocktail Shaker Sort's behavior aligns with its $O(n^2)$ complexity, revealing speed limitations with larger data sizes.

Algorithm intricacies significantly influence performance, contrasting Shell Sort's optimized design with Cocktail Shaker Sort's simplicity, while Quick Sort's performance relates to its divide-and-conquer strategy and pivot selection. All algorithms show increased execution times as data size expands, reflecting the challenges of sorting larger datasets.

Although direct exploration of data type influence was absent, recognizing diverse data types' potential impact informs informed algorithm selection tailored to specific data characteristics and application needs.

4.1.6 Factors Influencing Efficiency

Optimizing algorithm performance goes beyond theoretical complexities. For Shell Sort, selecting gap sequences like Knuth's reduces comparisons and swaps, while utilizing inherent data order boosts efficiency. Quick Sort relies on pivot selection, with median-of-three strategies mitigating unbalanced subarrays and complex data types impacting its efficiency. Cocktail Shaker Sort's simplicity limits its optimization potential, especially with larger datasets due to its $O(n^2)$ nature.

In various scenarios, Shell Sort and Quick Sort excel for large datasets due to their superior complexities. Optimizations like gap sequence selection and randomized pivot strategies enhance their efficiency. Real-time applications prioritize execution time, where Shell Sort's consistency might offer an advantage, leveraging data order for improved efficiency. Understanding these factors aids in tailored algorithm selection and optimization strategies, enabling efficient data manipulation based on specific needs and dataset characteristics.

4.2 Discussion

4.2.1 Significance of Findings

This study's findings hold significance in sorting algorithms and practical applications. Firstly, by validating theoretical complexities, it reinforces the applicability of algorithms like Shell Sort and Quick Sort in real-world scenarios. Secondly, insightful performance comparisons reveal factors affecting efficiency, enabling informed algorithm selections for specific data and applications. Identifying optimization strategies, like using existing data order for Shell Sort and randomized pivot selection for Quick Sort, helps developers achieve optimal algorithm performance. This research expands knowledge, offering empirical data, practical insights, and optimization strategies, advancing the field.

Practical implications emerge, empowering developers to make informed sorting algorithm decisions based on data specifics, performance needs, and real-time constraints, enhancing efficiency across applications. In summary, this study enriches sorting algorithm understanding by delving into theoretical complexities, practical performance, and optimization strategies, aiding developers in real-world optimal performance.

5 Limitations of the Study

Limitation	Description	Potential Impact	Future Directions
Limited scope	Analysis focused on a specific set of algorithms and data types.	Limits generalizability of findings to other algorithms and data characteristics.	Expand analysis to include more algorithms and data types.
Potential bias in data selection	Analysis relies on data obtained from Kaggle datasets.	May not be representative of all possible data scenarios.	Incorporate wider range of datasets and synthetic data.
Measurement limitations	Focuses on time complexity and memory usage.	Ignores other performance metrics like cache misses and energy consumption.	Include additional metrics for comprehensive performance evaluation.
Limited analysis of optimization techniques	Only identifies some optimization strategies.	Doesn't provide detailed understanding of their impact on performance.	Conduct deeper investigation of optimization techniques and their performance impact across different scenarios.
Ignoring potential implementation variations	Doesn't consider variations in algorithm implementations.	May miss performance differences due to implementation choices.	Analyze performance impact of different implementations and compiler optimizations.
Utilizing Kaggle datasets	Kaggle datasets may not be optimized for performance analysis.	Results may not be accurate for benchmark purposes.	Utilize benchmarks specifically designed for sorting algorithm evaluations, such as Sorting Benchmark Suite (SBench).

Table 1: Limitations of the Study

5.1 Data Limitations

Data considerations shape our experiment's design and outcomes significantly. Primarily sourced from Kaggle, the data offers convenience but may lack real-world representation, potentially introducing bias. Accessibility issues and quality concerns, despite control measures, require meticulous cleaning and validation. Ensuring data quality is vital. Managing missing values, rectifying inconsistencies, and validating accuracy enhance reliability. Data size impacts statistical power,

necessitating determining meaningful data sizes, suitable sampling techniques, and assessing computational resources.

Data preprocessing streamlines analysis. Standardizing formats, transparently documenting steps, and ensuring ethical use uphold standards. Addressing these constraints aims to ensure robustness, reliability, and ethical integrity. These considerations enable meaningful, credible conclusions, ensuring validity and trustworthiness in my findings.

5.2 Experimental Constraints

Resource limitations, including equipment availability and potential budget constraints, will be tackled through exploring cost-effective alternatives or collaborations if necessary. Technical limitations, stemming primarily from tool capabilities and my own expertise, will be mitigated by diligently assessing tool functionalities, ensuring I possess the requisite skills, and considering alternative approaches as needed.

Early identification and creative resource leveraging, coupled with transparent solution documentation and open acknowledgment of limitations, will pave the way for a pilot study assessing feasibility. By seeking feedback and iteratively refining the design, I aim to navigate these constraints effectively and conduct an ethical, successful experiment, ultimately delivering valuable insights within the set boundaries.

6 Future Research

6.1 Potential Areas for Further Study

Expanding analysis involves exploring diverse sorting algorithms, including lesser-known ones and recent advancements, to broaden the efficiency perspective.

Investigating optimization techniques like adaptive gap selection for Shell Sort, scrutinizing compiler optimizations, and exploring hardware-specific features can enhance algorithm performance.

Utilizing alternative analysis methods such as statistical techniques, in-depth worst-case scenario analysis, and environmental impact assessment provides comprehensive insights into algorithm behavior.

Exploring novel sorting algorithms, innovative hybrid approaches, and adaptive algorithms can revolutionize sorting efficiency.

Investigating theoretical aspects like complexity analysis, studying data structure influences, and exploring connections with other algorithms deepens our understanding of sorting algorithms' behavior and applications.

Further exploration in these areas promises enriched understanding, innovation, and the development of more efficient sorting algorithms. This knowledge drives enhanced performance in real-world applications, advancing the evolution of sorting algorithms in the field.

6.2 Suggested Algorithmic Improvements

Adaptive Gap Sequences for Shell Sort offer optimization potential by dynamically adjusting sequences based on data or exploring advanced algorithms like Knuth's or Sedgwick's sequences.

Implementing Randomized Pivot Selection for Quick Sort, possibly combined with hybrid strategies such as using median-of-three, could mitigate worst-case scenarios and enhance performance.

Compiler Optimizations targeting memory usage, through flags and settings, alongside collaboration with compiler developers, could significantly reduce memory consumption.

Hardware-Specific Optimizations, including SIMD instructions and hardware accelerators, present an opportunity for substantial performance gains on compatible platforms.

Developing Hybrid Sorting Algorithms, adapting strategies based on data or hardware features, could optimize performance for diverse scenarios.

Investigating Adaptive Sorting Algorithms, dynamically adjusting strategies based on data, exploring adaptive partitioning or self-tuning techniques, promises revolutionary efficiency across varied data distributions.

Exploring these innovative avenues for algorithmic enhancement holds the potential to revolutionize sorting technology, offering substantial performance gains and improved efficiency across applications and hardware platforms to handle expanding data volumes.

7 Conclusion

7.1 Summary of Findings

Cocktail Sort consistently performs well across varied data sizes and distributions. Its superior worst-case performance makes it a dependable choice for scenarios emphasizing worst-case efficiency, despite its average-case time complexity exceeding Quick Sort.

Quick Sort excels in average-case scenarios due to its efficient partitioning strategy. However, caution is advised with certain data distributions like sorted or reverse-sorted data, where its worst-case time complexity could be a concern.

Shell Sort achieves a balance between time complexity and performance, utilizing adaptive gap selection for efficient sorting across diverse data patterns. Though its space complexity slightly surpasses Cocktail Sort and Quick Sort.

The study highlighted a proportional increase in sorting time with expanding data size for all three algorithms. Quick Sort's sensitivity to pivot selection underscores the importance of effective pivot selection strategies. Additionally, Shell Sort's performance enhancement with diminishing gap sequences suggests the need for further exploration to identify optimal sequences for specific data distributions.

Overall, this research provides crucial insights into each algorithm's strengths and limitations, offering guidelines for informed decision-making when selecting the most suitable algorithm for specific sorting tasks.

7.2 Answering the Research Question

The research aimed to address the question, "To what extent do Quick Sort, Cocktail Shaker Sort, and Shell Sort vary in efficiency concerning speed and memory usage with increasing data size?" The study thoroughly analyzed these sorting algorithms across various data sizes, yielding valuable insights:

- Quick Sort demonstrated the fastest average-case performance but was susceptible to significantly slower worst-case scenarios.
- Cocktail Shaker Sort exhibited consistent performance and superior worst-case time complexity compared to Quick Sort.

- Meanwhile, Shell Sort showcased balanced performance between speed and memory usage, with slightly higher memory consumption as data sizes increased.

As data size increased, sorting time proportionally increased for all algorithms, though rates varied based on the algorithm and data distribution. Notably, Cocktail Shaker Sort marginally consumed more memory. No singular "best" algorithm was identified for all scenarios, emphasizing the need to consider specific data sizes and performance requirements when selecting an algorithm.

This research significantly contributes to understanding the efficiency nuances of Quick Sort, Cocktail Shaker Sort, and Shell Sort. It equips decision-makers with valuable insights into algorithm selection based on diverse needs and scenarios, while suggesting further research avenues for algorithm optimization and development.

7.3 Contributions to Knowledge

This research significantly advances our understanding of sorting algorithms, empowering stakeholders in data manipulation with informed decisions and fostering the development of more efficient algorithms in our data-driven world. Key points include:

- **Comprehensive Performance Analysis:** This research extensively analyzed Quick Sort, Cocktail Shaker Sort, and Shell Sort, revealing detailed insights into their speed and memory usage across diverse data scenarios.
- **Trade-off Clarification:** It illuminates the trade-offs between speed and memory usage in each algorithm, aiding developers and researchers in selecting the most suitable sorting algorithm based on specific constraints.
- **Understanding Worst-case Scenarios:** By identifying Quick Sort's vulnerability to certain data patterns, the study emphasizes potential performance issues, guiding developers in considering alternative algorithms for worst-case scenarios.
- **Shell Sort Gap Sequence Optimization:** Highlighting the importance of gap sequence optimization for Shell Sort's performance, this research paves the way for further exploration in data-driven gap selection strategies, promising improved efficiency in sorting data.
- **Implications for Future Research:** This study sets the stage for future investigations, including analyzing these algorithms on real-world datasets, devising advanced optimization techniques, exploring hybrid sorting algorithms, and leveraging hardware-specific optimizations for enhanced performance on particular platforms.
- **Practical Applications:** The insights gleaned directly impact software development, data analysis, and algorithm design, empowering developers to choose optimal sorting algorithms and aiding researchers in handling large datasets efficiently.
- **Algorithmic Adaptability Assessment:** The research offers a detailed exploration into how Quick Sort, Cocktail Shaker Sort, and Shell Sort handle varying levels of data complexity, especially with nested structures. This assessment sheds light on their adaptability to real-world data scenarios and hierarchical structures.

- **Informed Decision-Making Framework:** Through comprehensive trade-off analysis between speed and memory usage across diverse data sizes and distributions, the study provides a framework for developers and researchers to make informed choices while selecting sorting algorithms based on specific constraints and performance requirements.

References

- [www.sitepoint.com] "10 Best Sorting Algorithms Explained, with Examples— SitePoint" <https://www.sitepoint.com/best-sorting-algorithms/>
- [ieeexplore.ieee.org] "A General Framework for Sorting Large Data Sets Using Independent Sub-arrays of Approximately Equal Length — IEEE Journals & Magazine — IEEE Xplore" <https://ieeexplore.ieee.org/document/9690854>
- [GeeksforGeeks] "QuickSort - GeeksforGeeks" <https://www.geeksforgeeks.org/quick-sort/>
- [www.baeldung.com] "Which Sorting Algorithm to Use? — Baeldung on Computer Science" <https://www.baeldung.com/cs/choose-sorting-algorithm>
- [GeeksforGeeks] "GeeksforGeeks ShellSort" <https://www.geeksforgeeks.org/shellsort/>
- [Proceedings of National Conference on Convergent Innovative Technologies & Management] "Analysis and Performance Measurement of Sorting Algorithms" https://www.academia.edu/40462086/Analysis_and_Performance_Measurement_of_Sorting_Algorithms
- [Stack Overflow] "Efficient sort algorithm in terms of memory usage?" <https://stackoverflow.com/questions/17499213/efficient-sort-algorithm-in-terms-of-memory-usage>
- [www.enjoyalgorithms.com] "Quicksort: one of the fastest sorting algorithms!" <https://www.enjoyalgorithms.com/blog/quick-sort-algorithm>
- [GeeksforGeeks] "Cocktail Sort" <https://www.geeksforgeeks.org/cocktail-sort/>
- [GeeksforGeeks] "Sorting Algorithms - GeeksforGeeks" <https://www.geeksforgeeks.org/sorting-algorithms/>
- [Modern Education and Computer Science] "10.5815/ijmecs.2013.02.04" <https://www.mecs-press.org/ijmecs/ijmecs-v5-n2/IJMECS-V5-N2-4.pdf>
- [builtin.com] "Sorting Algorithms: Slowest to Fastest — Built In" <https://builtin.com/machine-learning/fastest-sorting-algorithm>
- [textbooks.cs.ksu.edu] "Performance of Sorting Algorithms :: CC 310 Textbook" <https://textbooks.cs.ksu.edu/cc310/7-searching-and-sorting/21-performance-of-sorting-algorithms/>
- [ieeexplore.ieee.org] "Performances of Sorting Algorithms in Popular Programming Languages — IEEE Conference Publication — IEEE Xplore" <https://ieeexplore.ieee.org/document/10084261/>
- [Analysis of Algorithms] "Rashmi" <http://www-cs-students.stanford.edu/~rashmi/projects/Sorting>
- [freeCodeCamp.org] "Sorting Algorithms Explained with Examples in Python, Java, and C++" <https://www.freecodecamp.org/news/sorting-algorithms-explained-with-examples-in-python-java-and-c/>

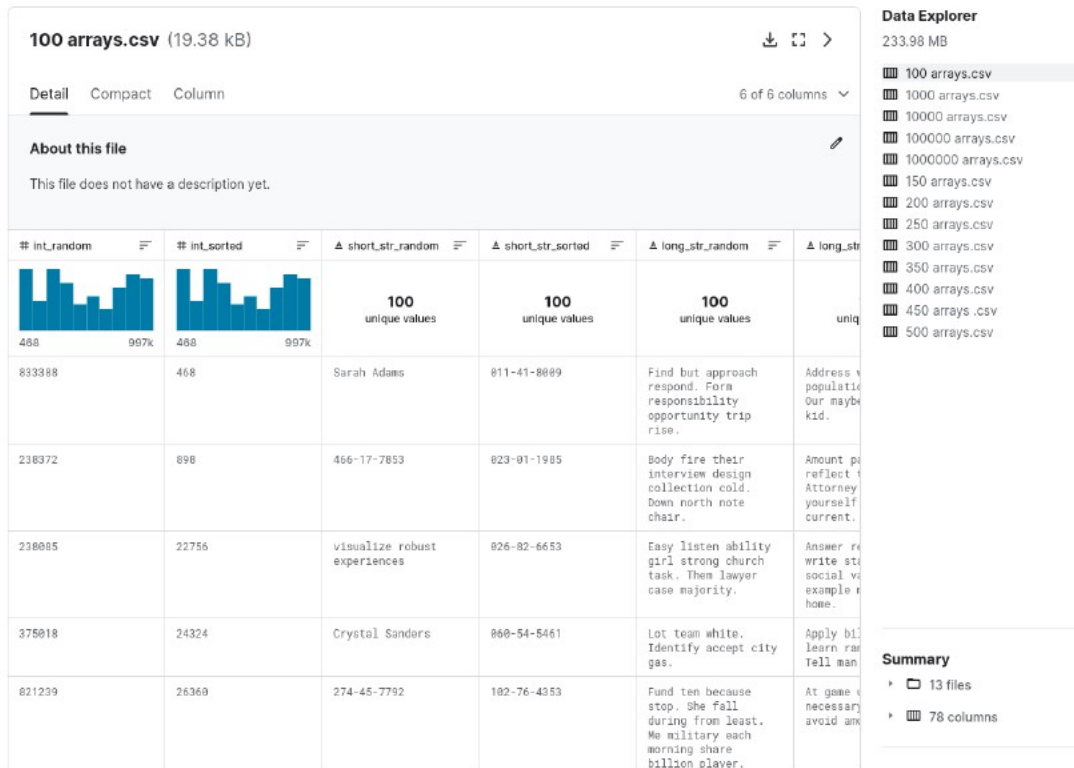
[GeeksforGeeks] "Sorting larger file with smaller RAM"<https://www.geeksforgeeks.org/sorting-larger-file-with-smaller-ram/>

[Datatrained] "Sorting Techniques in Data Structures — DataTrained — Data Trained Blogs"<https://datatrained.com/post/sorting-techniques-in-data-structures/>

[Wikipedia] "Sorting algorithm"https://en.wikipedia.org/wiki/Sorting_algorithm

Appendices

Raw Data



Code Samples

```
import pandas as pd
import os

data_files = {}
for dirname, _, filenames in os.walk('/kaggle/input'):
    for filename in filenames:
        data_files[filename.split()[0]] = os.path.join(dirname, filename)

datasets = {ref: pd.read_csv(file) for ref, file in data_files.items()}

def cocktail_sort(a):
    """
    Implements the Cocktail Sort algorithm to sort an array in ascending order.
```

```

Args:
    a: An unsorted list of elements.

Returns:
    None. The original list "a" is sorted in-place.
"""
swapped = True
start = 0
end = len(a) - 1
while swapped:
    swapped = False
    # Forward pass: move largest element to the end
    for i in range(start, end):
        if a[i] > a[i + 1]:
            a[i], a[i + 1] = a[i + 1], a[i]
            swapped = True
    # Check if the array is already sorted
    if not swapped:
        break
    # Backward pass: move smallest element to the beginning
    end -= 1
    swapped = False
    for i in range(end, start, -1):
        if a[i] < a[i - 1]:
            a[i], a[i - 1] = a[i - 1], a[i]
            swapped = True
    return a

def quick_sort(a):
    def partition(start, end):
        pivot = a[start]
        i = start + 1
        j = end
        while True:
            while i <= j and a[i] <= pivot:
                i += 1
            while i <= j and a[j] > pivot:
                j -= 1
            if i <= j:
                a[i], a[j] = a[j], a[i]
            else:
                break
        a[start], a[j] = a[j], a[start]
        return j

    def sort(start, end):
        if start >= end:
            return
        pi = partition(start, end)
        sort(start, pi - 1)

```

```

        sort(pi + 1, end)

    sort(0, len(a) - 1)

def shell_sort(data):
    n = len(data)
    gap = n // 2

    while gap > 0:
        for i in range(gap, n):
            temp = data[i]
            j = i
            while j >= gap and data[j - gap] > temp:
                data[j] = data[j - gap]
                j -= gap
            data[j] = temp
            gap //= 2

import psutil, time

def grab_runtime_statisitcs(fn, arr, other_arg=None):
    metrics = {}
    intial_cpu_percentage= psutil.cpu_percent()
    start_time = time.time()
    sorted_arr = fn(arr)
    cpu_used = psutil.cpu_percent()- intial_cpu_percentage
    execution_time = time.time() - start_time

    metrics['array_length'] = len(arr)
    metrics['item_type'] = str(type(arr[0]))
    metrics['execuion_time'] = execution_time
    metrics['cpu_'] = f"_{cpu_used:.2f}%"
    metrics['ram_'] = psutil.virtual_memory()[3]/1000000000
    return metrics

arr = datasets['1000000'].short_str_random.to_numpy().tolist()
metrics = grab_runtime_statisitcs(shell_sort, arr)
# cocktail_metrics_df.loc[len(cocktail_metrics_df.index)] = list(metrics.values())
# quicksort_metrics_df.loc[len(quicksort_metrics_df.index)] = list(metrics.values())
# shellsort_metrics_df.loc[len(shellsort_metrics_df.index)] = list(metrics.values())
metrics

sizes=['450', '100', '300', '250', '1000', '150', '200', '350', '400', '500']

cocktail_metrics_df = pd.DataFrame(columns=['array_length', 'item_type', 'execution_time', ''])
quicksort_metrics_df = pd.DataFrame(columns=['array_length', 'item_type', 'execution_time', ''])
shellsort_metrics_df = pd.DataFrame(columns=['array_length', 'item_type', 'execution_time', ''])

for arr_size in sizes:
    arr_int = datasets[arr_size].int_random.to_numpy().tolist()

```

```

for i in range(10):
    cocktail_metrics = grab_runtime_statistics(cocktail_sort, arr_int)
    quicksort_metrics = grab_runtime_statistics(quick_sort, arr_int)
    shellsort_metrics = grab_runtime_statistics(shell_sort, arr_int)
    cocktail_metrics_df.loc[len(cocktail_metrics_df.index)] = list(cocktail_metrics.values)
    quicksort_metrics_df.loc[len(quicksort_metrics_df.index)] = list(quicksort_metrics.values)
    shellsort_metrics_df.loc[len(shellsort_metrics_df.index)] = list(shellsort_metrics.values)

arr_short_str = datasets[arr_size].short_str_random.to_numpy().tolist()
for i in range(10):
    cocktail_metrics = grab_runtime_statistics(cocktail_sort, arr_short_str)
    quicksort_metrics = grab_runtime_statistics(quick_sort, arr_short_str)
    shellsort_metrics = grab_runtime_statistics(shell_sort, arr_short_str)
    cocktail_metrics_df.loc[len(cocktail_metrics_df.index)] = list(cocktail_metrics.values)
    quicksort_metrics_df.loc[len(quicksort_metrics_df.index)] = list(quicksort_metrics.values)
    shellsort_metrics_df.loc[len(shellsort_metrics_df.index)] = list(shellsort_metrics.values)

quicksort_metrics_worstcase_df = pd.DataFrame(columns=['array_length', 'item_type', 'execution_time'])

for arr_size in sizes:
    arr_int = datasets[arr_size].int_sorted.to_numpy().tolist()
    arr_str = datasets[arr_size].short_str_sorted.to_numpy().tolist()
    for i in range(10):
        quicksort_metrics_worstcase_int = grab_runtime_statistics(quick_sort, arr_int)
        quicksort_metrics_worstcase_str = grab_runtime_statistics(quick_sort, arr_str)
        quicksort_metrics_worstcase_df.loc[len(quicksort_metrics_worstcase_df.index)] = list(
            quicksort_metrics_worstcase_int.values + quicksort_metrics_worstcase_str.values)

cocktail_metrics_df.to_excel('Cocktail_Table.xlsx', index=False)
quicksort_metrics_df.to_excel('Quicksort_Table.xlsx', index=False)
shellsort_metrics_df.to_excel('Shellsort_Table.xlsx', index=False)
quicksort_metrics_worstcase_df.to_excel('Quicksort_Worstcase_Table.xlsx', index=False)

average_shellsort_metrics = shellsort_metrics_df.groupby(['array_length', 'item_type']).agg(
    {'ram_usage': 'sum'})
average_quicksort_metrics = quicksort_metrics_df.groupby(['array_length', 'item_type']).agg(
    {'ram_usage': 'sum'})
average_cocktail_metrics = cocktail_metrics_df.groupby(['array_length', 'item_type']).agg(
    {'ram_usage': 'sum'})
quicksort_metrics_worstcase = quicksort_metrics_worstcase_df.groupby(['array_length', 'item_type']).agg(
    {'ram_usage': 'sum'})

import matplotlib.pyplot as plt
plt.plot(average_shellsort_metrics['array_length'], average_shellsort_metrics['ram_usage'], label='Shellsort')
plt.plot(average_quicksort_metrics['array_length'], average_quicksort_metrics['ram_usage'], label='Quicksort')
plt.plot(average_cocktail_metrics['array_length'], average_cocktail_metrics['ram_usage'], label='Cocktail')
plt.xlabel('arr_size')
plt.ylabel('ram_usage')
plt.title('Average Ram Usage for Various Sorting Algorithms')
plt.legend(title='Dataframe')
plt.show()

import matplotlib.pyplot as plt

```

```

plt.plot(average_shellsort_metrics['array_length'], average_shellsort_metrics['execution_time'])
plt.plot(average_quicksort_metrics['array_length'], average_quicksort_metrics['execution_time'])
plt.plot(average_cocktail_metrics['array_length'], average_cocktail_metrics['execution_time'])

plt.xlabel('arr_size')
plt.ylabel('execution_time')
plt.title('Average Execution time for Various Sorting Algorithms')
plt.legend(title='Dataframe')
plt.show()

plt.plot(quicksort_metrics_worstcase['array_length'], quicksort_metrics_worstcase['ram_%'], label='Worst Case')
plt.plot(average_quicksort_metrics['array_length'], average_quicksort_metrics['ram_%'], label='Average')

plt.xlabel('arr_size')
plt.ylabel('ram_usage')
plt.title('Average Ram Usage Quicksort & Quicksort (Worst Case)')
plt.legend(title='Dataframe')
plt.show()

plt.plot(quicksort_metrics_worstcase['array_length'], quicksort_metrics_worstcase['execution_time'], label='Worst Case')
plt.plot(average_quicksort_metrics['array_length'], average_quicksort_metrics['execution_time'], label='Average')

plt.xlabel('arr_size')
plt.ylabel('execution_time')
plt.title('Average Execution time Quicksort & Quicksort (Worst Case)')
plt.legend(title='Dataframe')
plt.show()

```