

A Recovery-Aware Optimization Framework for Petabyte-Scale Distributed Data Processing: Balancing Performance, Reliability, and Infrastructure Cost

¹Chalapathi Koneni, ²Ashok Kumar

¹*Independent Researcher*

²*Walmart global tech*

ARTICLE INFO

Received: 01 March 2025

Revised: 17 Apr 2025

Accepted: 29 Apr 2025

ABSTRACT

A batch processing performance problem is not a single objective since throughput, recovery behavior and cost of infrastructure are highly interrelated at a petabyte scale. Over-parallelism can enhance performance but also raise the shuffling overhead, sensitivity to failures and operational cost, and conservative choices can decrease performance and serve to slow down data delivery. There is a major challenge as petabyte scale distributed processing systems have a rising failure probability, recovery cost and infrastructure cost. In this study, a new recovery-conscious optimization model has been suggested, which considers the decision made in checkpointing and also materialization, depending on workload features, failure requirements, re-computation cost, and system efficiency. A compatible model workload scenario that comprises 500,000 petabytes of workload cases is constructed in order to broaden synthetic data to simulate diverse distributed settings. The four recovery strategies No Checkpoint, Periodic Checkpoint, Adaptive Checkpoint and Selective Materialization have been compared by utilizing a cost-based optimization model and Pareto analysis. Findings indicate that Adaptive Checkpointing offers better recovery performance through proportional reliability, throughput and cost of operation.

Keywords: Distributed Data Processing, Petabyte Computing, Big Data Analytics, Fault Tolerance, Recovery Optimization, Cost-Aware Computing

I. INTRODUCTION

Background

The speed of the growth of the digital information has led to the necessity of petabyte-scale distributed data processing systems that can efficiently handle complex analytic workloads. Current systems like Apache Spark, Hadoop, and processing solutions based on clouds are based on mass, parallel implementation to obtain substantial volume [1]. Nonetheless, the mere speed of processing is not sufficient as failure, skew in the data, resource consumption, and-recovery overhead directly affect operational performance and cost [2]. Thus, designing distributed pipelines necessitates careful consideration of throughput, reliability and infrastructure cost [3]. In order to efficiently and affordably process and process data in a large enterprise setting, a recovery-conscious approach is crucial to ensuring predictable, scalable, and cost-effective data processing.

Research Problem

Current moves in the field of distributed data processing optimization predominantly revolve around enhanced performance in terms of speed and resource usage with little consideration to reliabilities and recovery expenses [4].

Aggressive parallelism has the potential of enhancing throughput, but it might cause an increase in shuffle overhead, failure probability, retries and infrastructure costs [5]. Equally, over check pointing enhances fault recovery at the expense of introducing more storage and putational expenses.

Research Aim and Objectives

Aim

The aim of the research is to create an analytical model recovery-aware that measures the performance of distributed data processing systems by operational cost, performance and reliability under operational load and failure requirements.

Objectives

- To develop a petabyte-scale workload simulation for distributed processing environments.
- To analyze the impact of workload characteristics, failures, recovery time, and processing cost.
- To compare recovery strategies based on throughput, reliability, and operational cost.
- To develop a recovery-aware optimization model for cost-effective strategy selection.

Research Contributions

The paper suggests a recovery-conscious optimization model of petabytes of distributed processing infrastructures, that considers the workload, fault analysis, and recovery, and computing costs, and strategy. In contrast to traditional performance-based models, the suggested model considers the costs of execution, costs of infrastructure, overhead during a checkpoint and the cost of recomputation to amend effective recovery selections. A workload simulation of synthetic petabyte volume is created where a variety of distributed processing conditions, such as different volumes of data, failures, shuffle intensity, and resource requirements, are modeled. The framework analyses strategies of No Checkpoint, Periodic Checkpoint, Adaptive Checkpoint and Selective Materialisation based on cost analysis. More so, Pareto optimisation achieves ideal trade-offs among throughput, reliability and cost of running the machine whereas machine learning is implemented as a complementary approach to automated recovery strategies recommendation.

II. RELATED WORK

Frameworks of Distributed Data Processing of Large-Scale Analytics

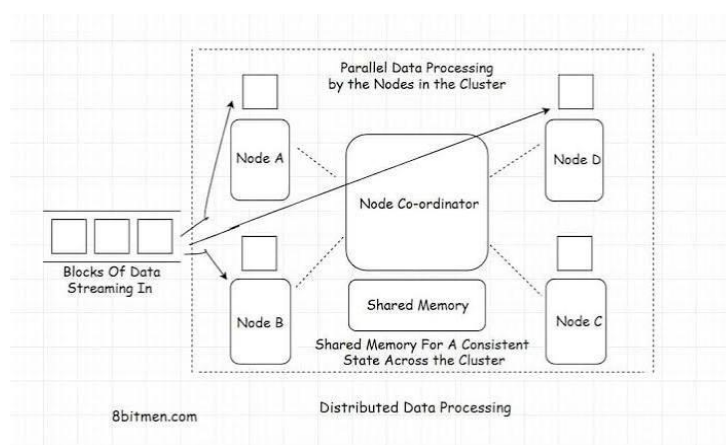


Fig. 1. Distributed Data Processing of Large-Scale Analytics

Large-scale distributed processing systems have become central to the operation of petabyte level data analytics, by separating computational loads into a set of processing nodes [6]. Hadoop MapReduce and Apache Spark technologies equipped scalable models of execution with parallelization of work, distributed storage and fault-consciousness in task-scheduling [7]. Those frameworks enhance the processing capacity and provide more

parallelism in computation but large clusters present issues with resource coordination, network communication, and consistency of executing. MapReduce-based frameworks make use of storage of intermediate data to aid in recovery whereas In-Memory Frameworks like Spark are more efficient in their execution but demand efficient memory and recovery management [8]. Distributed architectures should take the workload properties, data traffic and behavior of a cluster to attain scalable performance [9].

Fault Tolerance and Recovery in Distributed Systems

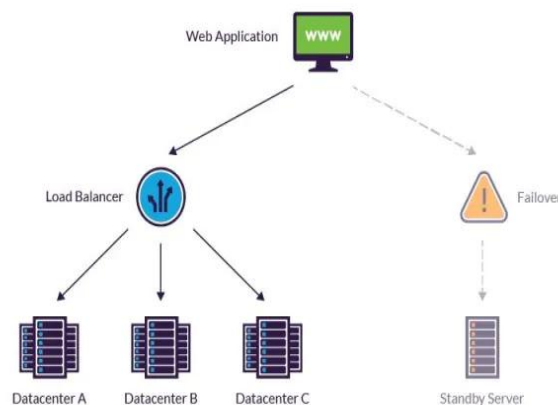


Fig. 2. Fault Tolerance and Recovery in Distributed Systems

Fault tolerance is extremely important to distributed data processing, as the size of a cluster directly proportional to the likelihood of hardware, network and software failures [10]. The studies [11] have examined various recovery methods, such as check pointing, replication, task recompilation and intermediate data materialization. Checkpoint-based methods minimize recovery time by caching execution states which means that failed processes can restart at previously saved states instead of executing an entire computation again [12]. The frequent check pointing however creates extra overhead in storage and processing. Research on cost-based fault tolerance emphasizes the idea that intermediate results should be selectively materialized at the expense of minimizing failure recovery cost through various tradeoffs between re-computation and storage [13]. Adaptive check pointing approaches take into account failure probability and workload behavior to enhance efficiency in the system [14].

Big Data Pipelines Performance Optimization and Resource Management

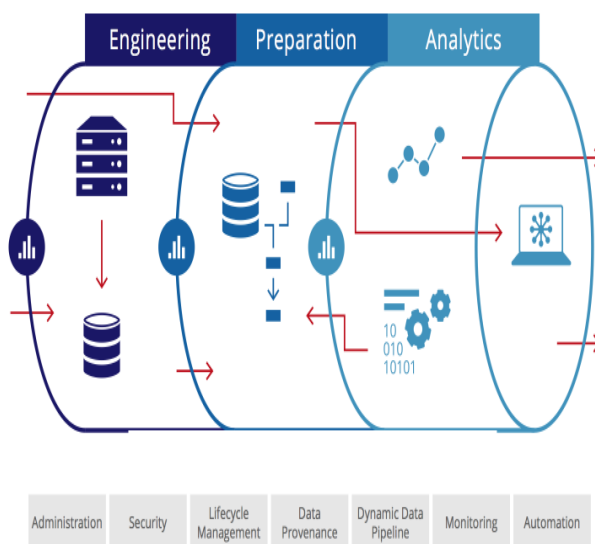


Fig. 3. Big Data Pipelines Performance Optimization and Resource Management

The aim of the Distributed processing systems is to optimize performance by ensuring better throughput, lowering latency, and efficient use of computing resources [15]. The current literature has seen the methods of how to schedule workloads, optimize partitions, localize data, allocate resources and isolate workloads [16]. Perfect partitioning leads to a less imbalanced processing whereas cunning scheduling leads to less wastage of resources. Nevertheless, more parallelism does not necessarily lead to higher performance since over-allocation of tasks may amplify overheads in the communication pathways, complexity of shuffling as well as resource usage [17]. Data skew and move of intermediate data are of particular concern to large-scale workloads in analytical and data processing like joins and aggregations [18]. Thus, the contemporary methods of optimization need the aspect of both performance and stability of the system [19].

Multi-Objective Processing Strategies and Cost-Aware Optimization

The need to optimize cost-wisely has been enhanced by cloud-based distributed data processing due to the contribution of computational resources, storage and data transport towards operation costs [20]. The literature [21] explores ways to lower the cost of infrastructure by using elastic allocation of resources, schedule workloads, and selective data persistence. The cost-based fault tolerance solutions show that the middle level materialization decisions must be dependent on the failure probability, recovery needs, and storage overhead instead of using any fixed strategies [22]. Likewise, new cloud systems promote dynamic scaled up to meet workload needs and prevent unjustified resource usage. Nonetheless, minimizing cost to an undesirable level leads to deteriorating reliability and processing schedules, whereas maximization of throughput can lead to higher costs incurred on infrastructural developments [23]. Consequently, the trio of performance, reliability, and cost variables need to be mutually considered in the distributed processing approach.

Research Gap

The available literature offers major value additions in distributed processing optimization, fault tolerance implementation and cost cutting policies. The majority of research methods, however, explores the improvement of throughput, efficiency of recovery, or cost optimization as independent objectives and do not examine the concerted effect of these factors. Conventional performance-oriented approaches might implement faster execution rates and augment failure-sensitivity and operations costs where the reliability-oriented approaches might impose extra overheads in terms of storage and computation.

III. PROPOSED RECOVERY-AWARE FRAMEWORK

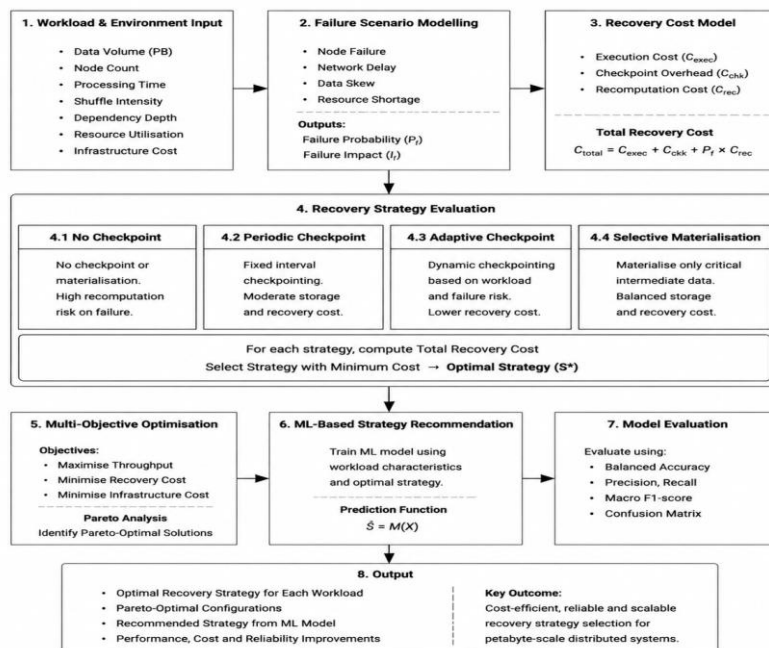


Fig. 4. Proposed Recovery-Aware Framework

The suggested recovery-conscious framework suggests a combined optimization framework that delivers an all-encompassing method of distributing data processing at petabyte scale by collectively addressing throughput, reliability, and costs. Workload generation starts with the framework simulating various processing scenarios like scan, join, shuffle, and aggregation workloads based on workload characteristics, resource parameters, and failure conditions. These workloads that are generated are the basis for determining the large-scale behavior of pipelines [24].

The second step is devoted to distributed pipeline modelling and execution in which workloads are modeled as dependency graphs and dealt with in a simulated distributed computing environment [25].

The third phase brings on board failure injection and recovery management to simulate realistic distributed failures, such as node loss, network delays, and data skew. Various recovery schemes are compared in terms of the recovery time, retries overhead, extra cost, and success rate. Lastly is the analytics and optimization layer, this utilizes machine learning models as a supportive analysis to find effective configurations [26].

IV. EXPERIMENTAL METHODOLOGY

Research Design

The framework followed to avoid petabytes of data processing environments to a recovery-aware optimization framework is quantitative simulation-based experimental. This research paper aims to study the trade-off between the execution efficiency, recovery reliability, the overhead of checkpoints and the cost of infrastructure. The suggested methodology models the large-scale distributed workloads and checks at which recovery strategies will be optimal in both different failure conditions.

Dataset of Petabyte-Scale Workload

Synthetic workload data, which consists of 500,000 distributed situations of processing, is created to model petabyte-scale computing environments. Every workload instance is a potential execution state as opposed to a single data record. The dataset models feature workloads of 1 PB to 500 PB at different cluster sizes, processing demands, failure probability, complexity of dependencies and data movement features.

Attributes relating to workloads are the volume of data, processing nodes, duration of execution, intensity of shuffles, depth of dependency, level of skew in data, size of intermediate data and usage of infrastructure [27]. These parameters have been chosen since they directly affect the decisions of checkpointing and recovery performance of distributed systems.

Failure Scenario Modelling

In order to test recovery behavior, various failure conditions were included in the simulative setting. The methodology takes into consideration node failures, delays over the network, skew in data and scarcity of resources as some of the key issues in terms of reliability in distributed processing systems [28]. Workload characteristics generated are failure probability, the number of failed nodes, network latency, lost packets, CPU used, and memory pressure. The expected failure impact is calculated as:

$$Ef = Pf \times If$$

where:

- Ef = expected failure impact,
- Pf = probability of failure occurrence,
- If = failure intensity.

For node failure scenarios, the failed node impact is estimated as:

$$Nf = N \times Pn$$

where:

- Nf = expected failed nodes,

- N = total cluster nodes,
- P_n = node failure probability.

Network degradation is represented as:

$$L_{net} = L_0 + \Delta L$$

where:

- L_{net} = effective network latency,
- L_0 = normal latency,
- ΔL = additional failure-induced delay.

The anticipated recomputation need and recovery time depends on each failure situation. This allows the means of evaluating recovery strategies when they are run in realistic conditions as opposed to them assessing only normal execution performance [29].

Development of Recovery Cost Model

The design that is created as a central element of the proposed framework is a recovery cost model. This model aims to identify the most effective recovery strategy with minimal cost of normal execution, checkpoint overhead, and anticipated recomputation cost. The sum of the recovery cost is achieved:

$$C_{total} = C_{execution} + C_{checkpoint} + P_f \times C_{recompute}$$

where:

- C_{total} = total recovery cost,
- $C_{execution}$ = normal execution cost,
- $C_{checkpoint}$ = materialization overhead,
- P_f = failure probability,
- $C_{recompute}$ = recomputation cost after failure.

The recomputation cost is estimated as:

$$C_{recompute} = T_r \times D_l \times S_i$$

where:

- T_r = recovery processing time,
- D_l = lost intermediate data volume,
- S_i = shuffle intensity.

Checkpoint overhead is calculated as:

$$C_{checkpoint} = S_c + T_c$$

where:

- S_c = storage overhead,
- T_c = checkpoint execution time.

This model allows comparing recovery strategies using operational cost and not using prediction accuracy.

Recovery Strategy Evaluation

There are 4 strategies of recovery assessed:

No Checkpoint:

Expenses traditional implementation without extra recovery overhead. This will be a very efficient strategy in case of low-risk workloads where failures are not frequent.

Periodic Checkpoint:

Adds constant state preservation to minimize recomputations in case of failures and adds storage overhead.

Adaptive Checkpoint:

Such dynamically redirects checkpoint behaviour based on the workload properties, including failure probability, complexity of its computation, and data-wire movement intensity.

Selective Materialization:

The stores only store critical intermediate states but of high recompute cost, minimizing undesirable checkpoint overhead.

The recovery cost of all four strategies is determined with respect to each workload scenario and the one having the minimum cost was chosen as allowed.

Recovery Optimization and Pareto Analysis

For each workload scenario, the cost function for every strategy is calculated:

$$S^* = \arg \min(C1, C2, C3, C4)$$

where:

- S^* = optimal recovery strategy,
- $C1$ = cost of No Checkpoint,
- $C2$ = cost of Periodic Checkpoint,
- $C3$ = cost of Adaptive Checkpoint,
- $C4$ = cost of Selective Materialization.

The model also uses multi-objective optimization to determine Pareto-optimal settings. The optimization takes into account three conflicting objectives which include: maximizing processing throughput, minimizing infrastructure cost, and minimizing recovery overhead. Pareto analysis determines what workload configurations are practical to improve one objective without compromising another providing practical advice to the administrators of distributed systems.

Machine learning is introduced as a complementary service of automated recommendation of the strategy. The recovery cost model was used to train a Random Forest classifier that used the workload characteristics to forecast the best recovery strategy. To overcome the issue of class imbalance, performance was measured with the help of balanced accuracy, macro F1-score, precision, recall, and confusion matrices.

Pseudo code

1. Load workload scenarios.
2. For each workload:
 Extract data volume, failure probability, re-computation cost, and infrastructure cost.
3. Calculate recovery cost for:
 - No Checkpoint
 - Periodic Checkpoint
 - Adaptive Checkpoint
 - Selective Materialization

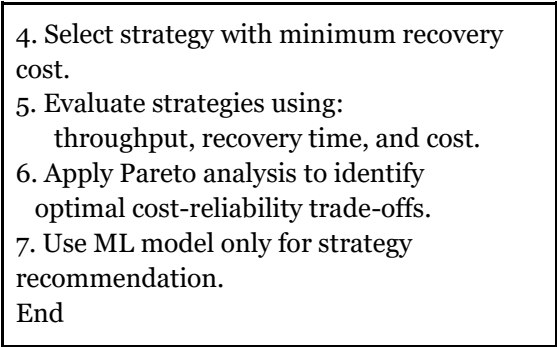


Fig. 5. Pseudo code

The pseudocode is the recovery-aware optimization process whereby workload properties are optimized, failure conditions are modelled and recovery costs are estimated on various strategies. The strategy picks the low-cost recovery strategy, measures performance with optimization metrics and uses machine learning as a supplementary tool to recommend strategies automatically.

Methodology Flow Diagram

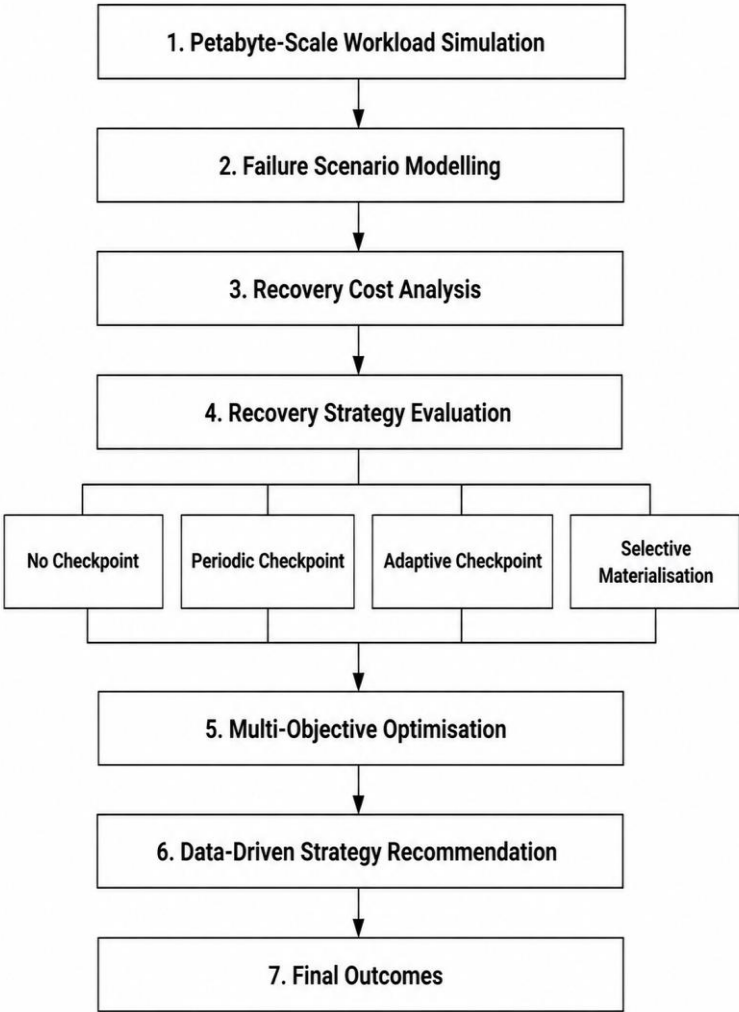


Fig. 6. Flow Diagram

The presented framework presents the entire workflow on the optimization of recovery, which starts with the simulation of workloads on the petabyte level and models of failures. It considers the cost aspects in recovery of various checkpoint strategies and then enforces multi-objective optimization. The framework discovers cost-effective and trustworthy settings, and the information-based advice is helpful in automated choice of optimal recovery plans in distributed processing settings.

V. RESULT AND ANALYSIS

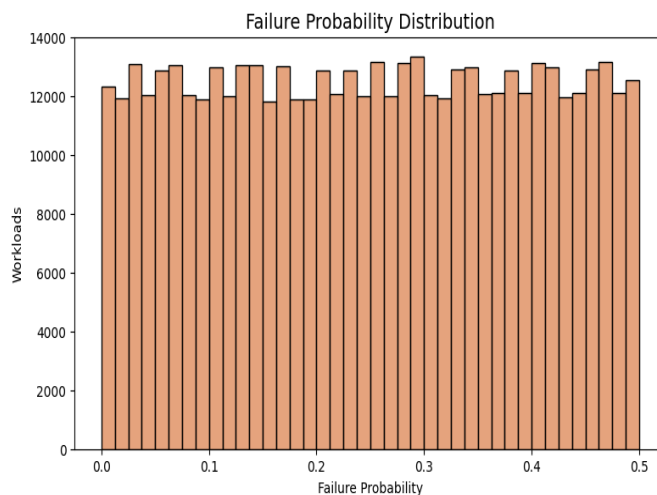


Fig. 7. Distribution of Failure Probability

The failure probability distribution shows that the reliability conditions of workloads are uniformly simulated in 500,000 scenarios of petabytes. The outcomes of probability are between 0 and 0.5, which are not parameters to cover only one set of failure conditions. This equal distribution facilitates the objective assessment of recovery plans in the various levels of operational risks.

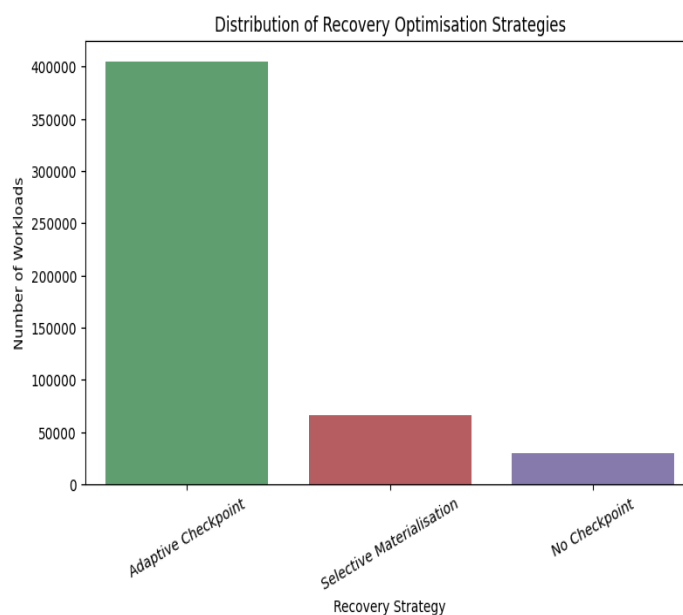


Fig. 8. Distribution of Recovery Optimisation Strategies

The distribution of the recovery strategies means that Adaptive Checkpointing is chosen as the dominating optimization technique because of its balancing checkpoint overhead and recovery efficiency. No Checkpoint and Selective Materialization strategies are selected based on certain conditions of workload where they are more efficient in terms of low failure or storage cost.

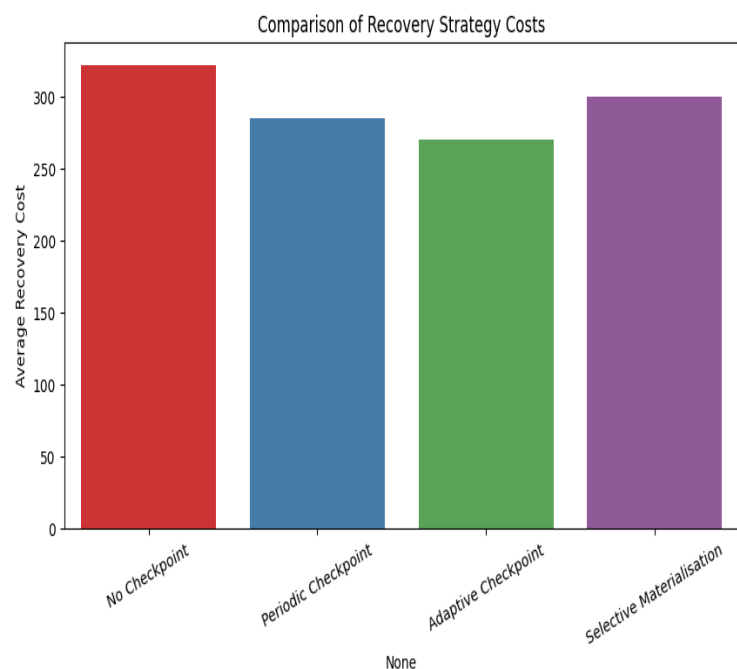


Fig. 9. Comparison of Costs of Recovery Strategies

Comparison of recovery costs reveals that Adaptive Checkpointing has the lowest average recovery cost when compared to calculated strategies. None of the Checkpoints will generate the highest cost since failures incur a high cost of re-computation whereas Periodic Checkpoint and Selective Materialization offer intermediate performance, decreasing the cost of recovery by investing more in storage.

	adaptive_improvement	selective_improvement
count	500000.000000	500000.000000
mean	13.202574	3.647096
std	13.234952	16.610990
min	-710.353732	-485.756728
25%	5.831878	-7.636967
50%	13.849444	3.951059
75%	21.542697	15.636666
max	34.356479	38.883095

Fig. 10. Analysis of Improvement in recovery

The analysis of the improvement shows that Adaptive Checkpointing offers a mean recovery cost cut of 13.2% over No Checkpoint execution. Selective Materialization yields smaller average improvement of 3.65 implying that adaptive strategies are more cost effective in workload which have varied attributes.

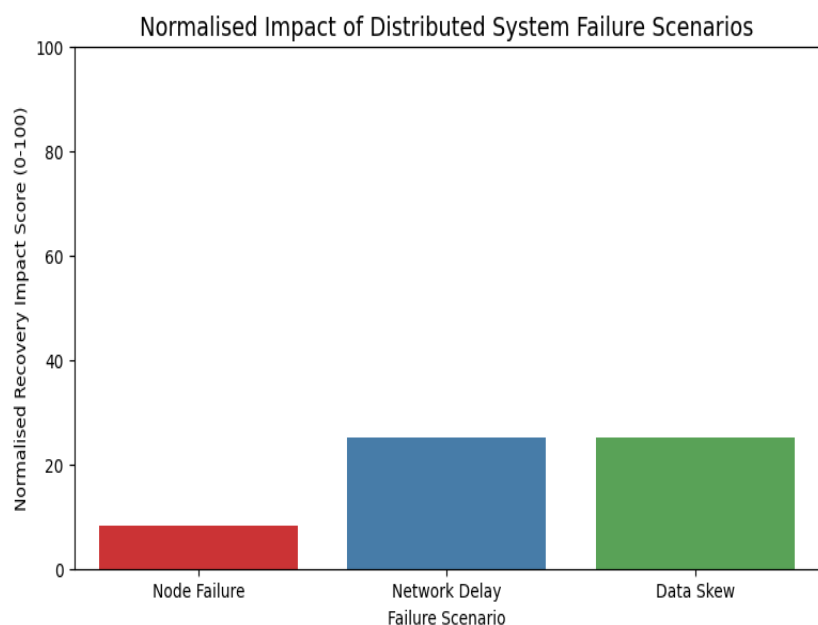


Fig. 11. Normalised Effect of Failure Scenarios of a Distributed System

The normalized failure impact analysis is a comparison of node failure, network delay, and data skew effects in a standard 0 -100 recovery impact scale. Network delay, data skew varies more in operation in comparison to node failure, indicating that analyzing various dimensions of failure is important in recovery optimization.

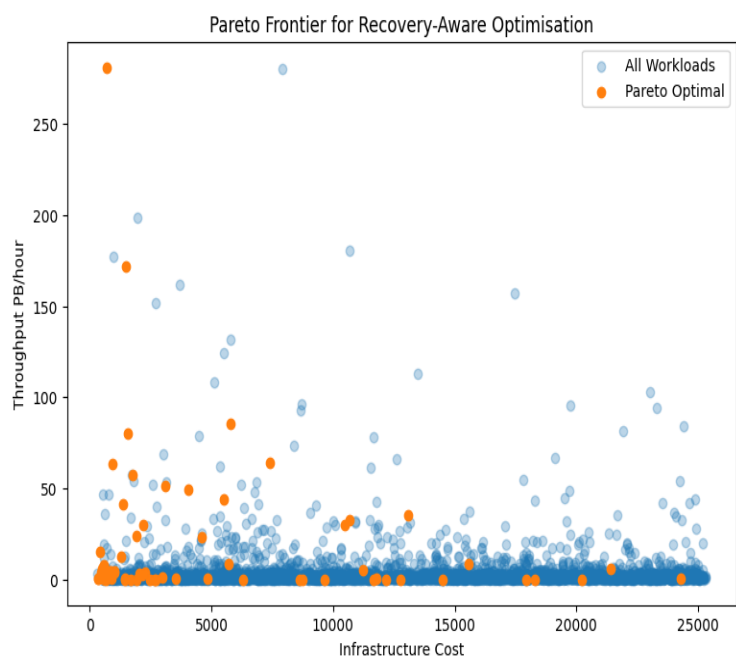


Fig. 12. Pareto Frontier Recovery-Aware Optimization

The Pareto frontier analysis determines those workload configurations that provide the best trade-offs between the cost of infrastructure and processing throughput. It can be seen that the Pareto-optimal points highlighted indicate situations in which it is hard to increase or decrease throughput or lower cost without influencing another objective, which helps to effectively choose a recovery strategy.

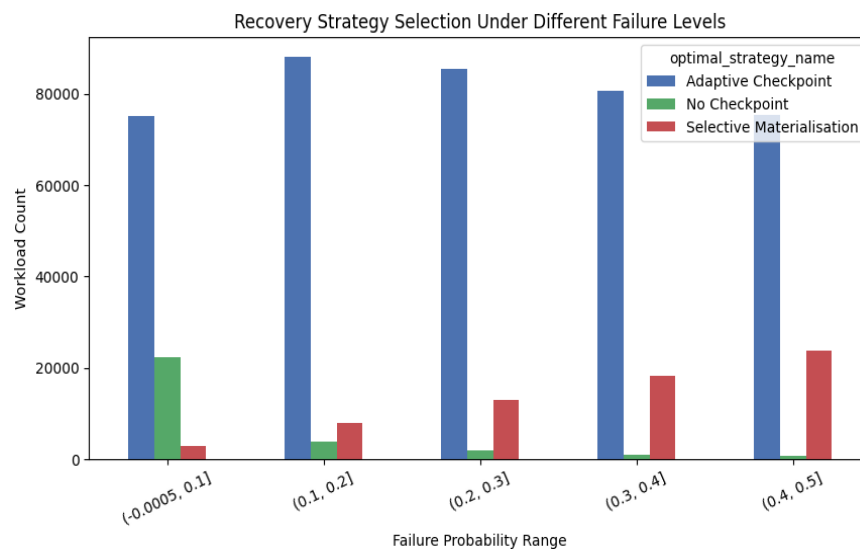


Fig. 13. Choice of Recovery Strategy at various Levels of failure

The failure level analysis shows that Adaptive Checkpointing can be used in various failure level ranges. Selective Materialization is increasingly important in the context of more difficulty in computation, due to the increasing probability on the left that the current critical states will remain unchanged by invalidation. This affirms recovery strategies ought to dynamically vary with workload risk conditions.

Discussion

TABLE I. Comparative Analysis of Recovery Strategies

Strategy	Runtime	Recovery Time	Storage Cost	Reliability
No checkpoint	High after failure	Very high	Low	Low
Periodic	Medium	Medium	Medium	Medium
Adaptive	Low	Low	Optimized	High
Selective	Low	Lowest	Targeted	Highest

The results indicate recovery-conscious optimization is a superior method to exploit in comparison to the use of predictive classification. Adaptive Checkpointing had the lowest recovery cost due to its inability to give checkpoint overhead and precomputation reduction balance to various workload situations with and without petabytes of data. The configurations with the Pareto analysis are those that maximize the throughput, reliability and cost of infrastructure. Simulation of failures ensured that recovery decisions should factor in the various operational risks such as network degradation and workload imbalance. The main contribution is system-level cost optimization, although machine learning can be used to aid automated strategy recommendation.

Limitation

→ The suggested framework is based on simulations of workloads at the scale of petabytes, instead of implementing the smoking on real distributed clusters, which can hamper direct tests on realistic production conditions with chaotic system behaviors.

→ The recovery cost model makes use of workload and failure parameters that have been established; hence, future research should test the framework with actual work traces of platforms like Spark, Hadoop, or cloud-based distributed systems.

VI. FUTURE RESEARCH

The suggested framework can be further tested by future studies through practical work operation datasets, production workload traces, and the extensive benchmark environments to test its feasibility within the contexts of the practical distributed computing in practical settings with Spark, Hadoop, or cloud-based distributed systems. Continuous advancements may involve adaptive recovery decision making with reinforcement learning, real-time failure prediction, and auto allocation of resources [30]. It would be practical to consider applications and reliability testing with the integration with large-scale production environments.

VII. CONCLUSION

This study introduced a recovery-conscious optimization model of a petabyte-scale distributed processing infrastructure through the examination of the trade-off among the execution performance, recovery reliability and infrastructure cost. Four recovery strategies: No Checkpoint, Periodic Checkpoint, Adaptive Checkpoint, Selective Materialization were evaluated with the aids of a synthetic workload simulation with a variety of failure conditions. Findings have shown that Adaptive Checkpointing was more cost efficient as it minimized needless re-computations without diminishing recovery. Pareto analysis also determined the best workload settings which struck the balance terms of throughput, cost and reliability. Machine learning was placed as a supportive recommendation system, with the main input being an optimization model of cost-driven recovery in scale distributed systems.

REFERENCE

- [1] Gupta, A., Agarwal, D., Tan, D., Kulesza, J., Pathak, R., Stefani, S. and Srinivasan, V., 2015, May. Amazon redshift and the case for simpler data warehouses. In Proceedings of the 2015 ACM SIGMOD international conference on management of data (pp. 1917-1923).
- [2] Hider Jr, R., Kleissas, D., Gion, T., Xenos, D., Matelsky, J., Pryor, D., Rodriguez, L., Johnson, E.C., Gray-Roncal, W. and Wester, B., 2022. The brain observatory storage service and database (bossdb): A cloud-native approach for petascale neuroscience discovery. *Frontiers in neuroinformatics*, 16, p.828787.
- [3] Gaddapuri, N.S., 2021. Big data storage observation system. *Power System Protection and Control*, 49(2), pp.7-19.
- [4] Adelusi, B.S., Ojika, F.U. and Uzoka, A.C., 2022. A Conceptual Model for Cost-Efficient Data Warehouse Management in AWS, GCP, and Azure Environments. *International Journal of Multidisciplinary Research and Growth Evaluation*, 3(2), pp.831-846.
- [5] Popescu, C., Merugu, D., Nguyen, G. and Shivakumar, S., 2019. WarpFlow: exploring petabytes of space-time data. *arXiv preprint arXiv:1902.03338*.
- [6] Qiao, S., Nicoara, A., Sun, J., Friedman, M., Patel, H. and Ekanayake, J., 2019. Hyper dimension shuffle: Efficient data repartition at petabyte scale in scope. *Proceedings of the VLDB Endowment*, 12(10), pp.1113-1125.

- [7] Keshireddy, S.R. and Kavuluri, H.V.R., 2021. Scalable ETL Frameworks for High Volume Transactional Systems in Distributed Data Warehouses. *The SIJ Transactions on Computer Science Engineering & its Applications*, 9(1), pp.24-28.
- [8] Zhang, Y., Cao, T., Li, S., Tian, X., Yuan, L., Jia, H. and Vasilakos, A.V., 2016. Parallel processing systems for big data: a survey. *Proceedings of the IEEE*, 104(11), pp.2114-2136.
- [9] Magnoni, L., Suthakar, U., Cordeiro, C., Georgiou, M., Andreeva, J., Khan, A. and Smith, D.R., 2015, December. Monitoring WLCG with lambda-architecture: a new scalable data store and analytics platform for monitoring at petabyte scale. In *Journal of Physics: Conference Series* (Vol. 664, No. 5, p. 052023). IOP Publishing.
- [10] Siddiqua, A., Karim, A. and Gani, A., 2017. Big data storage technologies: a survey. *Frontiers of Information Technology & Electronic Engineering*, 18(8), pp.1040-1070.
- [11] Dai, H., Wang, Y., Kent, K.B., Zeng, L. and Xu, C., 2022. The state of the art of metadata managements in large-scale distributed file systems—scalability, performance and availability. *IEEE Transactions on Parallel and Distributed Systems*, 33(12), pp.3850-3869.
- [12] Parthasarathi, S.H.K., Sivakrishnan, N., Ladkat, P. and Strom, N., 2019. Realizing petabyte scale acoustic modeling. *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*, 9(2), pp.422-432.
- [13] Moysiadis, V., Sarigiannidis, P. and Moscholios, I., 2018. Towards distributed data management in fog computing. *Wireless Communications and Mobile Computing*, 2018(1), p.7597686.
- [14] Shah, M. and Gogineni, A., 2022. Distributed Query Optimization for Petabyte-Scale Databases. *Int. J. Recent Innov. Trends Comput. Commun*, 10(10), pp.223-231.
- [15] Tariq, U. and Akram, N., 2022. Advanced Data Analytics and Big Data Frameworks for Intelligent Decision-Making Systems. *International Research Journal of Advanced Science*, 3(2), pp.43-52.
- [16] Luo, X., Gao, X., Tan, Z., Liu, J., Yang, X. and Chen, G., 2018, July. D2-tree: A distributed double-layer namespace tree partition scheme for metadata management in large-scale storage systems. In *2018 IEEE 38th International Conference on Distributed Computing Systems (ICDCS)* (pp. 110-119). IEEE.
- [17] Zhang, J., 2015. Research on Improving Reliability, Energy Efficiency and Scalability in Distributed and Parallel File Systems.
- [18] Gao, Y., Gao, X., Yang, X., Liu, J. and Chen, G., 2019. An efficient ring-based metadata management policy for large-scale distributed file systems. *IEEE Transactions on Parallel and Distributed Systems*, 30(9), pp.1962-1974.
- [19] Mahmud, M.S., Huang, J.Z., Salloum, S., Emara, T.Z. and Sadatdiynov, K., 2020. A survey of data partitioning and sampling methods to support big data analysis. *Big Data Mining and Analytics*, 3(2), pp.85-101.
- [20] Andersen, M.P., Kumar, S., Brooks, C., von Meier, A. and Culler, D.E., 2015, November. DISTIL: Design and implementation of a scalable synchrophasor data processing system. In *2015 IEEE International Conference on Smart Grid Communications (SmartGridComm)* (pp. 271-277). IEEE.
- [21] Ather, S., Muslin-Ud-Din, H., Nabeel, M., Ahsan, M. and Hassan, B., 2019. Several Adaptive Replica Synchronization Approaches for Distributed file System. *VAWKUM Transactions on Computer Sciences*, 7(1), pp.1-8.
- [22] Sikeridis, D., Papapanagiotou, I., Rimal, B.P. and Devetsikiotis, M., 2017. A Comparative taxonomy and survey of public cloud infrastructure vendors. *arXiv preprint arXiv:1710.01476*.
- [23] Ahmed, E., Yaqoob, I., Hashem, I.A.T., Khan, I., Ahmed, A.I.A., Imran, M. and Vasilakos, A.V., 2017. The role of big data analytics in Internet of Things. *Computer Networks*, 129, pp.459-471.
- [24] Fazul, R.W.A. and Barcelos, P.P., 2022, April. An event-driven strategy for reactive replica balancing on apache hadoop distributed file system. In *Proceedings of the 37th ACM/SIGAPP symposium on applied computing* (pp. 255-263).

- [25] Arafa, Y., Barai, A., Zheng, M. and Badawy, A.H.A., 2018, September. Fault tolerance performance evaluation of large-scale distributed storage systems HDFS and Ceph case study. In 2018 IEEE High Performance extreme Computing Conference (HPEC) (pp. 1-7). IEEE.
- [26] Wang, F., Melesse Vergara, V., Leverman, D.B. and Oral, S., 2016. FCP: A Fast and Scalable Data Copy Tool for High Performance Parallel File Systems. Oak Ridge National Laboratory (ORNL), Oak Ridge, TN (United States).
- [27] Ben-Nun, T. and Hoefler, T., 2019. Demystifying parallel and distributed deep learning: An in-depth concurrency analysis. ACM Computing Surveys (CSUR), 52(4), pp.1-43.
- [28] Elshater, Y., 2017. Optimizing data locality in analytic workloads over distributed computing environments. Queen's University (Canada).
- [29] Du, J., 2019. Online Physical Design And Materialization in Scalable Data Management. University of Toronto (Canada).
- [30] Barrefors, B., 2015. Dynamic data management in a data grid environment.