

A pilot review of the Google cluster workload trace 2019, methodology and its alternatives: analysis of workload in large scale data centres

Akash Patel¹, Amit Nayak², Khushi Patel³, Anand Patel⁴

¹Department of Information Technology, Devang Patel Institute of Advance Technology and Research (DEPSTAR), Faculty of Technology and Engineering (FTE), Charotar University of Science and Technology (CHARUSAT), Anand, India

²Department of Computer Science and Engineering, Devang Patel Institute of Advance Technology and Research (DEPSTAR), Faculty of Technology and Engineering (FTE), Charotar University of Science and Technology (CHARUSAT), Anand, India

³Department of Computer Engineering, Devang Patel Institute of Advance Technology and Research (DEPSTAR), Faculty of Technology and Engineering (FTE), Charotar University of Science and Technology (CHARUSAT), Anand, India

⁴Information Technology Department, Dharmsinh Desai Institute of Technology (DDIT), Faculty of Technology, Dharmsinh Desai University (DDU), Nadiad, India

Article Info

Article history:

Received Jun 5, 2025

Revised Apr 14, 2026

Accepted May 17, 2026

Keywords:

Cloud computing

Data centers

Deep learning

Google cluster trace

Machine learning

Workload

ABSTRACT

Cloud data centres require shared services that are highly available, elastic capacity, managed operations, and robust recovery capabilities to facilitate the next generation of efficient, reliable, and diverse connected computing environments. Nevertheless, large-scale cloud infrastructures continue to fail regularly despite their availability, scalability, and cost efficiency, primarily due to low resource utilisation and inadequate early-stage failure management. The key to effective resource management and minimizing failures in such settings is understanding the nature of the workload and its failure modes. The current review considers the Google cluster workload trace 2019 to investigate workload and failure patterns and to generalise the results of 24 articles. The analysis is also compared with other major datasets, such as Microsoft Azure Trace, Tencent Trace, and Alibaba Trace. The paper establishes the relevance of Google cluster traces, describes the key contents of the 2019 dataset, and contrasts prior literature with respect to research objectives, trace datasets, significant results, and limitations. Moreover, it briefly describes methods for analyzing and modeling cluster traces and identifies gaps in the research that should be addressed to advance the study of cluster traces.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Amit Nayak

Department of Computer Science and Engineering

Devang Patel Institute of Advance Technology and Research (DEPSTAR)

Faculty of Technology and Engineering (FTE)

Charotar University of Science and Technology (CHARUSAT)

Changa 388421, Anand, India

Email: Amitnayak.it@charusat.ac.in

1. INTRODUCTION

In commercial settings, large-scale data centers have become the standard solution for a vast range of applications. This is mainly attributable to the fact that these centers offer significant potential for availability, scalability, and cost-effective operational expenses. Cloud computing systems fail due to their complexity, heterogeneity, openness, and dispersion. These failures can degrade cloud system dependability and cost users and providers a lot of money. Cloud providers may incur substantial losses from these errors.

Thus, data analysis on real-world sources is essential to solving the problems facing large-scale cloud data centers.

This paper presents an analysis and survey of the 2019 Google Cluster traces dataset [1], which Google provides. It includes 2.4 terabytes of workload traces collected from eight clusters throughout May 2019. This dataset has enabled a comparison of the development of scheduling decisions made by cluster managers, opening up novel opportunities. The new dataset covers alloc sets, best-effort batch scheduler, vertical scalability, and task dependencies. This information aids in job scheduling explanations. Unlike the 2011 trace dataset [2], the 2019 dataset allows us to extend our hypothesis to eight clusters across three continents and to examine nonconformist tendencies [3].

A wide range of research articles have been published to date on the Google cloud traces 2019 dataset. Additionally, some of these articles have compared this dataset with other publicly available cloud datasets. When it comes to research publications, however, there is a lack of high-quality articles.

In cloud computing research, several real-world data center trace datasets have been released, along with their respective performance standards. 2011 featured the release of Google cluster trace [2], which was considered to be one of the earliest and most significant datasets. Furthermore, the Alibaba Cluster Dataset [4] is another dataset recently made available. The purpose of this resource-consumption dataset is to enhance utilization and reduce infrastructure costs. It comes from the production cluster of Alibaba, which operates both online services and batch processes. On the other hand, the diverse requirements of workloads necessitate the development of novel scheduling algorithms to manage co-located heterogeneous application groups effectively [5].

This paper consists of the following sections. Section 1 is an Introduction. Section 2 describes the significance of Google cluster traces. Section 3 discusses the Google cluster trace 2019 dataset. Section 4 presents a comparative analysis using alternative datasets. Section 5 describes methods for analyzing and modeling cluster traces. Section 6 describes research challenges and open areas. The final section, 7, presents the conclusion.

Research objectives:

- Presents a systematic literature review of 29 articles that deploy the Google cluster workload trace 2019 dataset, which summarizes evidence on workload as well as failure behaviour in real-world cloud clusters.
- Gives a comparative study of the Google cluster traces with the MS Azure, Tencent and Alibaba traces including their similarities, differences and their complementary advantages to cloud research.
- Explains the meaning and contents of the Google cluster dataset on 2019 such as what kind of workloads, resources, and events are being recorded, and how these have been exploited in previous literature.
- Determines solutions to unresolved issues and underinvestigated findings in the current body of literature on cluster trace and outlines future research questions in cluster trace analysis specifically within the context of AI- and GPU-based workload constraint prediction.

2. SIGNIFICANCE OF GOOGLE CLUSTER TRACES

They conducted an in-depth analysis of the 2019 Google cluster trace dataset, comprising workload traces totaling 2.4 gigabytes from eight clusters located in different parts of the world. Google's production cloud has been tasked with examining the characteristics of aborted or failed jobs and correlating them with significant factors. These factors include resource utilization, work priority, scheduling class, job duration, and the number of times the task was resubmitted. According to our findings, several significant features of failed tasks contributed to the work's failure and could be used to develop an early failure prediction system [1].

Several researchers have used GARCH and ARIMA models to develop techniques for predicting the reliability and response times of online services. In addition, several prior studies found that virtual machines exhibited relatively consistent workload patterns. Using Hidden Markov Modeling, they developed a method to describe and predict virtual machine workload patterns. They developed several approaches, and this strategy was among them. The use of statistical techniques based on Google Cluster traces has led to comparisons of Google data centers with grid or high-performance computing (HPC) systems [6]. This comparison is made even though Google data centers are classified as cloud stages.

A thorough multi-view assessment strategy for dual-cloud workloads was described, and a case study on Google Cluster Tracing was used to illustrate the concept [4]. Several studies used the workload of the Google cluster to determine the distribution of cloud servers' maintenance time and time-to-failure [6]. These analyses were conducted to determine the likelihood of failure.

The Google Cluster trace Dataset comprises eight Borg cells, each located in a different time zone, such as New York, Helsinki, and Singapore. Machine tables, Collection tables, and Instance tables are the

three primary types of tables that are included in each Borg cell, which is comprised of five tables. To be more explicit, Machines are divided into two tables: the MachineEvents table describes events associated with machines, and the MachineAttributes table provides information on machine features. Collections are what Google refers to as jobs and alloc sets, instances are what they call tasks and alloc instances, and items are what they use to refer to collections or instances. The CollectionEvents and InstanceEvents tables, respectively, record the lifecycle of collections and instances [7].

The 2019 traces include only resource requests and utilization; they do not capture end users, data, or access patterns to storage systems and other services. Each record contains a microsecond timestamp relative to a reference point 600 seconds before the start of the trace period. The timestamps are 64-bit integers. For example, a 20-second event occurring 620 seconds after the trace begins would have a timestamp of 620 seconds. Two exceptional timestamp values are available for events beyond the trace window: 0 for pre-trace events and $263 - 1$ (MAXINT) for post-trace events. Limiting use measurement times to 300 seconds, offset rules are used to detect occurrences relative to the trace period. Measurements are accurate to the nearest second but presented in microseconds for clarity. These timestamps are generally accurate, although small variations may occur due to clock drift across cluster nodes, rather than scheduling or usage concerns [8].

When conducting research on the cloud, it is essential to ensure the availability of empirical datasets that demonstrate the extent to which the cloud meets customer requirements. For example, numerous public cloud platform traces for cluster trace analysis have been made available. These include the Facebook trace, the Taobao dataset, the Google cluster trace, and the Alibaba trace 2018. In November 2011, Google released a “cluster usage trace dataset”. There is a remarkable amount of detail in this trace, as well as a wide variety of tasks. The positions encompass a wider range of responsibilities, including providing online services such as search queries and other uses of MapReduce methods. The term “workload” refers to the number of jobs being received, the tasks associated with those jobs that users have submitted, and the load a particular system experiences when conducting tasks. Within the scope of this study, the trace dataset is characterized with respect to how tasks are processed and terminated within the cluster, how terminated jobs are distributed, and how scheduling class and priority influence terminated jobs [9], [10].

To capture long-range dependencies in cloud workloads, this analysis is conducted using workloads derived from a standard real-world dataset, Google Cluster trace. The research shows that the metrics under analysis are heavy-tailed; moreover, a mathematical form is derived to characterize long-range dependencies in aggregate workloads. The workloads in the Google cluster trace are analyzed in prior work, depending on the characteristics of the tasks and the resources used. They observed that occupations in the forest industry can be classified into three categories: short, medium, and long tasks. This classification is based on the resource-intensive nature of the jobs. Additionally, they worked on clustering workload patterns. On the other hand, their research does not focus on understanding the nature of cloud workloads that are dependent on an extended period of time. The varied nature of workloads in Google cluster trace is the subject of discussion in another piece of research. The report observes that the workload is highly dynamic, meaning it changes over time and is driven by numerous brief assignments that require rapid scheduling decisions. The research that they have done, on the other hand, does not analyze the long-range reliance that cloud workloads have [11].

A base predictor is a term that is widely used to refer to each of the numerous prediction models that are utilized in an ensemble method. This includes the use of multiple prediction models to forecast future outcomes. In Google Cluster trace 2011, the workload is analyzed by partitioning it into three domains: positive domain (POS), negative domain (NEG), and border domain. A three-way decision is the tool employed for this type of workload analysis. EMA prediction methodology is utilized for the given stable period. For the volatile period, the Holt-Winters prediction model is applied, and this is sensitive to jump patterns in the data. For the volatility period, a weighted prediction technique forms the core prediction model within the predictive framework, as the model responsible for determining the relevant outcomes. Outcome weighting is used to estimate the final projected workload based on weights obtained using reciprocal error techniques. Hence, models with higher accuracy are given greater weight, whereas those with low accuracy receive minimal weight [12].

3. OVERVIEW OF GOOGLE CLUSTER TRACES 2019 DATASET

An extensive dataset that is made accessible to the general public is known as Google cluster traces. More than 12,500 nodes were used to generate the data for this large dataset. During the period from February 8 to March 28, Google’s cluster traces comprise approximately 28 million tasks and 672,074 jobs. Each job has unique characteristics based on its scheduling category, priority, resource requirements, and resource utilization. The following summary covers Google Traces’ key statistics. The links between failed tasks, scheduled courses, and resource demands are broken to eliminate them. Basic information was derived

by analyzing failure patterns over a 29-day period. The Google Trace of 29 days reveals significant differences in patterns of failures of jobs and tasks. Incomplete activities in Google Trace went highly between days 2 and 10, which shows significant destruction of particular data centre resources [6].

Google has only recently released new trace data applicable at a large scale. Compared with the prior version, the May 2019 dataset contains eight distinct Google computer clusters, referred to as cells. In contrast to earlier research, this approach enables inter-cluster analysis, which was previously difficult. Jobs, machine events, and other utilization data are included in the Google Trace 2019 dataset. Jobs and machine events contain information on job and allocation scheduling, as well as on machine life-cycle operating processes. Use the provided data to compute the normalized CPU and memory usage for each activity, ranging from 0 to 1. The computation divides actual usage by the maximum number of Google compute units (GCUs) or memory capacity of all computers in a cell [13].

Google released its third and latest dataset. It was released in early 2020 and records cloud resource usage from eight clusters, one of which had more than 12,000 computers in May 2019. This dataset focuses on resource requests and consumption and does not include end users, unlike the 2011 dataset. The 2019 dataset adds task-reserved shared resources, CPU-use histograms for each five minutes, and job-parent relationships for master-worker relationships, such as MapReduce jobs [14]. Table 1 shows a comparison between the Google cluster trace 2011 and the Google cluster trace 2019 datasets.

The dataset of Google cluster 2019 can be successfully applied in the development of the predictive models to optimize cluster performance and resource utilization since it includes abundant metrics that can be used to analyze workload patterns and resource requirement within the cloud environment. These understandings facilitate more sophisticated methods like reinforcement learning (RL), in which frameworks like Q-learning and deep Q-networks may optimize the resource allocation dynamically by balancing the workloads and enhancing the efficiency of the systems. Also, predictive systems can use forecasting techniques like NARX and ARIMA to predict future workloads using past data to enable the proactive allocation of resources, minimize costs, and achieve high performance. Machine learning solutions are also instrumental in task assignment where artificial neural networks, as well as, ensemble algorithms can be used to determine the best node-task assignments resulting in an efficient use of resources. Nevertheless, regardless of its potential, the model accuracy and performance may be influenced by challenges, including data quality problems, and the dynamic and unpredictable nature of workloads, and to maximize the benefits of predictive analytics in cluster management, it is important to mitigate these limitations [15]-[17].

Table 1. Comparison between Google Cluster trace 2011 (v2) and 2019 (v3) datasets

	Version 2 (cell)	Version 3 (total)	Version 3 (cell avg)
Total events	37780	309833	38729
Initial machines	12477	91902	11487
New ADD	106	4653	582
REMOVE	8957	153779	19222
UPDATE	7380	324	40
REJOIN	8860	151077	18884
REMOVE (permanently)	98	2702	338
Repair probability	0.98917	0.98243	-

4. COMPARATIVE ANALYSIS WITH ALTERNATIVE DATASETS

Recent studies on cloud computing and workload trace analysis have focused on improving resource utilization, scheduling efficiency, workload prediction, and system reliability in large-scale distributed environments. These studies utilize real-world cluster trace datasets from major cloud providers such as Alibaba, Google, Tencent, and Microsoft Azure to evaluate different models, scheduling techniques, and workload generation frameworks. Table 2 in the Appendix [18]-[22]. Comparative analysis presented below highlights the objectives, datasets, major findings, and identified limitations of existing research works, thereby providing a clear understanding of current advancements and potential future research directions in cloud workload trace analysis.

5. METHODS FOR ANALYZING AND MODELING CLUSTER TRACES

A comprehensive collection of workload data that was collected from eight Google Borg compute clusters over the month of May 2019 is provided by the Google cluster workload traces 2019 dataset, which serves as a landmark to the development of cloud computing research. It is a cornerstone for

understanding the dynamics of large-scale cluster management and workload scheduling inside a cloud environment, and the dataset obtained from Kaggle is the source of the information. Given that the project is primarily concerned with the management of heavy traffic, the Google cluster workload traces 2019 dataset has emerged as an important resource that can be utilized for job submission, scheduling choices, and resource utilization across a variety of clusters. The purpose of this component is to conduct an analysis in order to evaluate the efficacy and outcome of our algorithm. A comparison of the method with more conventional load balancing techniques, such as least connections and least response time, is used to analyze the effectiveness of the algorithm. In order to do the comparison, critical variables such as resource usage, latency, throughput, and error rate were taken into consideration. The efficiency and dependability of load-balancing algorithms in real-world circumstances may be analyzed using these measures, which are acceptable for the purpose. This analysis is carried out by distributing 500 tasks to each of the algorithms, beginning with 100 tasks and increasing the number of tasks by 100 in each successive stage. In this, the metrics were watched and kept closely in a manner for analysis to determine which implementation had an effective outcome from all of them [7].

This research analyzed the Google Trace Dataset. This dataset is extremely valuable due to the fact that it is currently being worked on by a large number of academics, and the findings of this research have the potential to influence the next generation of cloud computing and data centers. Four perspectives were used to analyze the Google Trace to enhance cluster performance. Memory affects work performance more than the CPU, according to our initial data research. Despite memory being the barrier, Google recently updated servers with full CPUs to have full memory. Borg performance will improve by reducing RAM by 25% from 2218 full CPU and memory systems to 10,188 with 0.5 CPU and 0.2493 memory. It can be seen that more restrictive professions tend to delay less in the second experiment, although the third trial has given results showing that lower-priority jobs contribute a lot in terms of rescheduling and augment cloud overhead, which were not expected. It was proposed to minimize the time spent on each job and let the lower-priority jobs finish for optimal use of resources. Finally, it was observed that the cloud management system puts the maximum constraints to KILLED jobs, and this is the main observation of this study. It should be broken down into jobs with subtasks in order to minimize the number of jobs eliminated and improve performance [9].

Although Google cluster dataset of 2019 has been widely explored, there are still major operational insights to be made. In this paper, the authors suggest artificial intelligence models that are run on graphics processing units to forecast the workload limitations, which is one of the primary bottlenecks in cloud environments at the scale. The approach would help in better provisioning of its resources, lessening performance degradation, and extracting more value out of the existing cluster trace information by anticipating such constraints in the future. Research suggests combining machine learning algorithms like XGBoost with hyperparameter tuning approaches like swarm-based artificial bee colony optimization to save energy and time. Accurate constraint number forecasting saves time and energy. Our research endeavors to push the limits of the capability of this dataset, with the goal of discovering significant knowledge that has an influence in the actual world. Also, the Borg system should optimize its efficiency by analyzing an option that disregards VMs that have an unusually low number of resources: either CPU or RAM. Therefore, the future research area is providing sophisticated ideas to Borg about resource allocation techniques, with the aim to eliminate potential bottleneck effects such under-provisioned virtual machines create. It is further highly recommended that other major cloud service providers including Amazon Web Services (AWS), Microsoft Azure, and Alibaba be considered in subsequent research efforts. Comparative studies across these platforms can therefore provide the research community as a whole with valuable findings [9].

Within the realm of current cluster management systems, trace analysis has been performing a significant role in providing an understanding of the behaviors that they exhibit. This paper presents the findings of our inter-cluster analysis of the recently disclosed Google trace 2019 data set. Our analysis focused on the distribution of underlying physical machines, such as the number of machines and resources, and computational task request metrics, such as work duration, time used, and cycles per instruction. This further reveal that the cells possessing almost the same composition of machine resource could consume those resources in an entirely different manner by considering the pattern of request of tasks. A shape has been figured out through an unsupervised learning algorithm from various clusters; such shapes appear to be two types - concave and convex forms. In case of the convex shape, it would be the presence of coherent scale variation for all dimension, but concave depicts missing variance along with any dimension.

Provide a characterization analysis of a workload trace that was obtained over a period of two months from a production machine learning as a service cluster in Alibaba that had more than 6,000 GPUs. The explanation of cluster scheduling challenges is given. Such issues include low utilization of graphical processing units, long times to wait in the queues, the existence of certain hard-to-schedule applications requiring high-end GPUs but have specific scheduling requirements and may cause imbalanced loading across heterogeneous machines, causing bottlenecks on CPUs. These current solutions include a scheduler,

which allows for GPU sharing and uses a reserving-and-packing strategy in separating high-GPU jobs from low-GPU tasks to better manage the scheduling of PAI workloads. The observation is that most of the jobs are usually running and scheduled by groups, with less than one GPU per instance. Only very few business-critical jobs use the high-end GPUs in one system connected via NVLinks. The CPU seems to be a bottleneck for such low GPU resource jobs; in fact, it also hosts data pre-processing and simulation. Some problems remain, including possible CPU bottleneck and load imbalance for heterogeneous devices. The trace—the most detailed in terms of workloads and cluster scale—is published to promote further research on better GPU scheduling [17].

Within the scope of this research, A comprehensive strategy for the task failure detection model is proposed. Five traditional models of machine learning and three deep learning models were created using the GCT dataset to assess the performance of these different models in job and task failure prediction. The best model is the XGBoost classifier, which gives an accuracy score of 94.35% and an F-score of 0.9310 as obtained by the performance analysis. Two models were found at the task level that are at excellent levels. These models are based on decision tree and random forest. A score of 89.75% for accuracy and a score of 0.9154 for the F-score have been calculated for both models. Taking everything into consideration, the findings have demonstrated that the traditional ML models perform significantly better than the DL models when it comes to classifying the termination status of jobs and tasks. Furthermore, according to the findings of our research of the significance of features, the scheduling class and the CPU request are the most important characteristics for Traditional ML, whereas the disk space request and the memory request are the most important features for DL in terms of the prediction of task failure. On the other hand, when it comes to task-level prediction, the most essential element for traditional ML is resource demands, whereas priority is the most critical factor for DL models. In conclusion, our scalability analysis revealed that Traditional ML models are capable of making predictions in a reasonable amount of time, despite the fact that they are implemented on consumer-level hardware [14].

In the field of cloud computing, several studies have made use of clustering to analyze cloud trace data in order to improve resource management. Much of this research have been successful. During the clustering process, the dimensions and the approach that are utilized can have a significant impact on the quality of the cloud traces that are clustered. The paper presented two contributions: an in-depth analysis, considering the possibility of improving clustering quality using dimensions that are more efficient than traditional ones and procedures different from the usual ones used. The challenge of predicting which of the two elements will occur is something that has not yet been anticipated, so a detection method, EFection, was invented. A comparison of the performance of three internal validation measures—Davies-Bouldin, Silhouette Coefficient, and Calinski-Harabasz—was conducted to find the most suitable validation method for this technique while developing the EFection approach. This method allows one to find which dimension and method will produce the highest quality clustering of a particular cloud trace without an analysis of the entire dataset. In doing so, the Davies-Boulden index as well as the Pearson correlation coefficient is used. Performance metrics: Initially, the testability of EFection started with a verification of accuracy and precision to identify instances over various combinations of attributes with different methods. Then its utility is shown by making two applications in related analysis. A comparison with the most recent documented efforts in the literature revealed that EFection was potentially applicable in almost 83% of cases and had the highest accuracy preference. Additionally, it surpassed the most recent attempt by eleven percentage points and had higher stability than the previous approach [23], [24].

Using the Google cluster trace data set version 3, which was acquired from about 96,000 computers, Failure statistics and machine lifetimes over time were studied. In less than a minute from the previous failure of the machine, it was calculated that 13% that another machine on the same network switch will fail. A Markov chain model was designed in order to predict which machines are in which states at any time. With this model and computed probability estimates, very accurate predictions about machine status could be done for several days. Using the estimated machine states, the trend of the active machine was generated and compared to the trend provided in the Google cluster workload traces dataset. The error obtained was 1.76%. Within the compressed format, the data set version 3 has a size of 2.7 terabytes. The data collection was made available by Google in both the JSON format and as BigQuery tables. Regarding JSON, the size of the data collection is greater than 7 terabytes. With such a large data set comes a set of challenges that are entirely unique. The majority of the challenges that are related with the JSON format may be resolved by using the version of the data set that is hosted on Google BigQuery. BigQuery is quite expensive, but due to the JSON format, it can run on commodity hardware. For this research work, the data set was retrieved from the Google Storage bucket, and then it was downloaded compressed in JSON format. To increase memory efficiency, the compressed JSON was imported into the Spark Cluster and saved in Apache Parquet format because uncompressed JSON occupies too much of the system's memory. Only the needed data from the Apache Spark Parquet file was read, converted, and saved to a CSV file for further processing with

Python [22]. In order to increase the accuracy of the forecast, First, present a three-way ensemble data center workload prediction. It was stated that workloads were continuous, varying, and delays. Thus, using the annealing scheme, workload-division thresholding was discovered. Subsequently, by making a three-way decision principle to assign one of three types of forecasting models depending upon workload type and prior a priori error prediction, accuracy can be improved upon. Lastly, TWD-RCPM enhanced workload forecast accuracy by 69.0%, 68.6%, and 72.6% when tested on Google cluster trace CPU load monitoring logs compared to ARIMA, NN, and DMASVR-3WD [12].

The analysis of Google cluster traces needs to utilize sophisticated clustering techniques that can efficiently operate with high-dimensional and dynamically changing data. Incorporating current algorithms and models, researchers will be able to enhance the accuracy and interpretability of the results of clustering, which is necessary to analyze data meaningfully. Trimmed and sparse clustering are automated methods of clustering that emphasize important features and deemphasize noise so that the process is less sensitive to outliers and the parameters are more easily chosen by users. Frameworks such as BURST, too, are tailored to streaming time-series data, allowing real-time analysis through dynamically adapting to changes and estimating the best number of clusters. Besides accuracy, explainability is an important part of clustering and such tools as Cluster-Explorer can be used to improve interpretability by determining patterns and characteristics that are important in clusters and employ methods such as frequent-itemset mining. Nevertheless, some problems are still present, especially in the problem of guaranteeing easy and understandable outcomes with complex data, so more research is necessary to enhance the efficiency and versatility of clustering approaches [25]-[28].

6. RESEARCH CHALLENGES IN CLUSTER TRACE ANALYSIS

There are several challenges in the Google cluster workload traces dataset and its alternative datasets namely Alibaba Trace, MS Azure Trace, and Tencent Trace.

- Edge cloud AI inference on IoT data presents novel challenges. Edge clouds, like conventional cloud platforms, will operate numerous tenant apps on each server. These applications utilize edge server hardware, such as accelerators. Traditional resources like CPUs and server GPUs accept virtualization for application multiplexing, while edge accelerators do not. Using a DNN accelerator for several tenant applications might cause performance interference owing to a lack of isolation methods, such as virtualization. This can negatively impact response times for latency-sensitive IoT applications. Developing novel cluster resource management strategies is necessary to efficiently multiplex shared edge cloud resources for latency-sensitive applications [23].
- Several researchers found that analyzing cluster workload traces on a single platform might be challenging due to variability, limiting the applicability of an approach. Most microservices in the Alibaba cluster are more sensitive to CPU than memory, according to an analysis of Alibaba trace. Such exploration is encouraged by trace cross-platform comparison [25].
- Although the Google cluster dataset 2019 has been analyzed, insights remain unexplored. Our solution involves using AI and GPU power to estimate job limits, a significant challenge in cloud systems [9].
- The search for new clusters was attempted in the 2019 Google Cluster dataset, measured at 2.4 terabytes, huge in comparison with its predecessor that measured only at 44 gigabytes. Therefore, proper effort has to be put behind time and resources to maneuver the amount of data processed. To overcome these challenges, a Google cloud tool called BigQuery was used. It enables complex searches on the Google Cluster dataset without having to download the entire trace [9].
- Google Cluster trace is an image of their data center operations. As infrastructure develops, this trace may become outdated and limited. Without standards or procedures, decision-making might be challenging to understand. Privacy and confidentiality concerns make Google Trace challenging. It is challenging to manage and analyze huge amounts of data in large-scale systems, detect valuable patterns, and generate trustworthy results. Generalizability may be restricted by Google's rules, which may differ significantly from those of other businesses. Analyzing and categorizing the Google Cluster trace provides valuable insights into managing large-scale distributed systems, such as cloud computing and data centers [10].
- Researchers want to use unsupervised learning techniques to analyze work categories. A challenge arises when dealing with millions of multi-dimensional data items. Due to the curse of dimensionality, with larger data sets, each dimension greatly increases processing costs. Additionally, data with severe outliers and heavy bias might bias the clustering algorithm. These challenges are overcome by reducing the metrics to three based on the outcome of correlation analysis. Other strategies include scaling of outliers using Box-Cox transformation and centering the data by a standard scaler which also eliminates standard deviation [13].

- Our scheduler policy design has challenges, many of which are yet unaddressed. Next, some open challenges that may arise in GPU clusters with heterogeneous machines are going to be dealt with: mismatched specifications and request for instances, overcrowding low-performance GPU computers, and imbalanced loads in high-end machines [17].
- This challenge helps cloud service companies save expenses. The sustainable development goals are also supported. To train this model, consider using alternative datasets. Transfer learning techniques may be used to create high-quality prediction models for cloud resource management [14].
- Straggler tasks, which are sometimes related to the long tail problem, is still a burning issue in cluster computing since only a few slow running tasks can greatly postpone the execution of a whole job. These stragglers have adverse effects on the performance and resource efficiency of the system as a whole. Consequently, proper identification and control of such activities are critical, and superior frameworks and strategies are necessary to maximize resource allocation and enhance the efficiency of the system in general [29].

7. CONCLUSION

The survey has examined Google cluster workload trace datasets of 2011 and 2019, and the other data sets including the Alibaba, MS Azure and Tencent cluster traces. It compares the 2019 Google trace to be larger, more varied, and closer to the modern production workloads compared to the 2011 trace, with the alternative datasets offering supplementary insights of co-located workloads, cloud operations, and large-scale scheduling behavior. Nevertheless, there are a number of gaps in research in the literature. To begin with, the majority of the current studies are dedicated to one dataset, and there is no commonly adopted cross-trace benchmarking framework to assess the workload analysis techniques in various cloud settings. Second, preprocessing conventions, schema mapping, and trace filtering vary significantly between datasets restricting reproducibility and direct comparison. Third, the temporal variations in the nature of workloads, schedule patterns and resource consumption are not yet adequately investigated, particularly between the 2011 and 2019 Google traces. Lastly, future studies ought to focus on transferability, standardized measurements and standardized evaluation procedures to ensure that trace-based models can be tested efficiently using heterogeneous production datasets.

Overall, cluster trace datasets continue to be crucial in comprehending cloud workload behavior, yet the discipline still lacks standardized practices in analysis and validation across datasets. The fill-in of these gaps will ensure that future trace-based scheduling, prediction, and resource-management studies become more generalizable and useful in practice.

FUNDING INFORMATION

Authors state no funding involved.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Akash Patel	✓	✓	✓	✓	✓	✓		✓	✓	✓			✓	
Amit Nayak	✓	✓				✓				✓	✓	✓		
Khushi Patel		✓		✓	✓	✓				✓				
Anand Patel		✓		✓	✓	✓		✓		✓	✓			

C : Conceptualization

M : Methodology

So : Software

Va : Validation

Fo : Formal analysis

I : Investigation

R : Resources

D : Data Curation

O : Writing - Original Draft

E : Writing - Review & Editing

Vi : Visualization

Su : Supervision

P : Project administration

Fu : Funding acquisition

CONFLICT OF INTEREST STATEMENT

The authors declare no conflict of interest.

DATA AVAILABILITY

Data generated and analyzed during this study are available from the corresponding author upon reasonable request. The data, which contain information that could compromise the privacy of research participants, are not publicly available due to certain restrictions. Requests should be directed to the corresponding author's contact details.

REFERENCES

- [1] F. H. Bappy, T. Islam, T. S. Zaman, R. Hasan, and C. Caicedo, "A deep dive into the google cluster workload traces: analyzing the application failure characteristics and user behaviors," in *Proceedings - 2023 International Conference on Future Internet of Things and Cloud, FiCloud 2023*, IEEE, Aug. 2023, pp. 103–108. doi: 10.1109/FiCloud58648.2023.00023.
- [2] Google, "Google cluster workload traces 2019." Accessed: May 20, 2026. [Online]. Available: <https://research.google/resources/datasets/google-cluster-workload-traces-2019/>
- [3] "Google cluster data 2011," Google. Accessed: May 20, 2026. [Online]. Available: <https://research.google/tools/datasets/cluster-workload-traces/>
- [4] N. Dezhnabad, S. Ganti, and G. Shojia, "Cloud workload characterization and profiling for resource allocation," in *Proceeding of the 2019 IEEE 8th International Conference on Cloud Networking, CloudNet 2019*, IEEE, Nov. 2019, pp. 1–4, doi: 10.1109/CloudNet47604.2019.9064138.
- [5] Alibaba, "Cluster data collected from production clusters in Alibaba for cluster management research." Accessed: May 20, 2026. [Online]. Available: <https://github.com/alibaba/clusterdata>
- [6] H. Bommala, V. Uma Maheswari, R. Aluvalu, and S. Mudrakola, "Machine learning job failure analysis and prediction model for the cloud environment," *High-Confidence Computing*, vol. 3, no. 4, p. 100165, Dec. 2023, doi: 10.1016/j.hcc.2023.100165.
- [7] V. E. Jyothi, S. Alugoju, K. Indrāja, N. S. Chowdary, A. Madhuri, and S. Sindhura, "Adaptive cloud load balancing with reinforcement learning: leveraging google cluster data," in *5th International Conference on Electronics and Sustainable Communication Systems, ICESC 2024 - Proceedings*, IEEE, Aug. 2024, pp. 572–579, doi: 10.1109/ICESC60852.2024.10689973.
- [8] M. Yildiz and A. Baiocchi, "Data-driven workload generation based on google data center measurements," in *IEEE International Conference on High Performance Switching and Routing, HPSR*, IEEE, Jul. 2024, pp. 143–148, doi: 10.1109/HPSR62440.2024.10635925.
- [9] D. Shahmirzadi, N. Khaledian, and A. M. Rahmani, "Analyzing the impact of various parameters on job scheduling in the Google cluster dataset," *Cluster Computing*, vol. 27, no. 6, pp. 7673–7687, Sep. 2024, doi: 10.1007/s10586-024-04377-8.
- [10] S. Jomah and S. Aji, "Google cloud trace: characterization of terminated jobs," in *Proceedings - 2nd International Conference on Advancement in Computation and Computer Technologies, InCACCT 2024*, IEEE, May 2024, pp. 517–522, doi: 10.1109/InCACCT61598.2024.10551146.
- [11] S. Gupta and A. D. Dileep, "Long range dependence in cloud servers: a statistical analysis based on Google workload trace," *Computing*, vol. 102, no. 4, pp. 1031–1049, Apr. 2020, doi: 10.1007/s00607-019-00779-4.
- [12] Y. Lin, A. Barker, and S. Ceasay, "Exploring characteristics of inter-cluster machines and cloud applications on Google Clusters," in *Proceedings - 2020 IEEE International Conference on Big Data, Big Data 2020*, IEEE, Dec. 2020, pp. 2785–2794, doi: 10.1109/BigData50022.2020.9377802.
- [13] Q. Weng et al., "MLaaS in the wild: workload analysis and scheduling in large-scale heterogeneous GPU clusters," *Proceedings of the 19th USENIX Symposium on Networked Systems Design and Implementation, NSDI 2022*, pp. 945–960, 2022.
- [14] T. N. T. Asmawi, A. Ismail, and J. Shen, "Cloud failure prediction based on traditional machine learning and deep learning," *Journal of Cloud Computing*, vol. 11, no. 1, p. 47, Sep. 2022, doi: 10.1186/s13677-022-00327-0.
- [15] S. M. Ali and G. Kecskemeti, "EFection: effectiveness detection technique for clustering cloud workload traces," *International Journal of Computational Intelligence Systems*, vol. 17, no. 1, p. 202, Aug. 2024, doi: 10.1007/s44196-024-00618-1.
- [16] A. Umer, A. N. Mian, and O. Rana, "Predicting machine behavior from Google cluster workload traces," *Concurrency and Computation: Practice and Experience*, vol. 35, no. 5, Feb. 2023, doi: 10.1002/cpe.7559.
- [17] R. Shi and C. Jiang, "Three-way ensemble prediction for workload in the data center," *IEEE Access*, vol. 10, pp. 10021–10030, 2022, doi: 10.1109/ACCESS.2022.3145426.
- [18] Y. Liang, N. Ruan, L. Yi, and X. Su, "An approach to workload generation for modern data centers: a view from Alibaba trace," *Benchmark Transactions on Benchmarks, Standards and Evaluations*, vol. 4, no. 1, p. 100164, Mar. 2024, doi: 10.1016/j.tbench.2024.100164.
- [19] C. Katsakioris, C. Alverti, K. Nikas, D. Siakavaras, S. Psomadakis, and N. Koziris, "FaaSRAil: employing real workloads to generate representative load for serverless research," in *HPDC 2024 - Proceedings of the 33rd International Symposium on High-Performance Parallel and Distributed Computing*, New York, NY, USA: ACM, Jun. 2024, pp. 214–226, doi: 10.1145/3625549.3658684.
- [20] J. Li, Q. Wang, P. P. C. Lee, and C. Shi, "An in-depth comparative analysis of cloud block storage workloads: findings and implications," *ACM Transactions on Storage*, vol. 19, no. 2, pp. 1–32, May 2023, doi: 10.1145/3572779.
- [21] C. Jiang, G. Han, J. Lin, G. Jia, W. Shi, and J. Wan, "Characteristics of Co-allocated online services and batch jobs in internet data centers: a case study from Alibaba cloud," *IEEE Access*, vol. 7, pp. 22495–22508, 2019, doi: 10.1109/ACCESS.2019.2897898.
- [22] Q. Wang, Y. Tian, X. Yu, L. Ding, and X. Zhang, "Where is the Traffic Going? A comparative study of clouds following different designs," *IEEE Transactions on Services Computing*, vol. 16, no. 2, pp. 1473–1484, Mar. 2023, doi: 10.1109/TSC.2022.3182047.
- [23] Q. Liang, W. A. Hanafy, A. Ali-Eldin, and P. Shenoy, "Model-driven cluster resource management for AI workloads in edge clouds," *ACM Transactions on Autonomous and Adaptive Systems*, vol. 18, no. 1, pp. 1–26, Mar. 2023, doi: 10.1145/3582080.
- [24] L. Ruan, X. Xu, L. Xiao, L. Ren, N. Min-Allah, and Y. Xue, "Evaluating performance variations cross cloud data centres using multiview comparative workload traces analysis," *Connection Science*, vol. 34, no. 1, pp. 1582–1608, Dec. 2022, doi: 10.1080/09540091.2021.2015289.
- [25] L. Sliwko, "Cluster workload allocation: a predictive approach leveraging machine learning efficiency," *IEEE Access*, vol. 12, pp. 194091–194107, 2024, doi: 10.1109/ACCESS.2024.3520422.

- [26] S. Ofek and A. Somech, "Explaining black-box clustering pipelines with cluster-explorer," *Proceedings of the VLDB Endowment*, vol. 18, no. 5, pp. 1495–1508, Jan. 2025, doi: 10.14778/3718057.3718075.
- [27] A. Giannoulidis, A. Gounaris, and J. Paparrizos, "BURST: rendering clustering techniques suitable for evolving streams," *Proceedings of the VLDB Endowment*, vol. 18, no. 11, pp. 4054–4063, Jul. 2025, doi: 10.14778/3749646.3749675.
- [28] J. A. Bernabé-Díaz, M. Franco, J. M. Vivo, and J. T. Fernández-Breis, "Optimizing clustering-based analytical methods with trimmed and sparse clustering," *Computers in Biology and Medicine*, vol. 194, p. 110436, Aug. 2025, doi: 10.1016/j.compbiomed.2025.110436.
- [29] T. Le Duc, C. Nguyen, and P.-O. Östberg, "Energy-aware straggler task management framework for sustainable cloud environment," *Authorea Preprints*, 2024, doi: 10.36227/techrxiv.173398069.98852332/v.

APPENDIX

Table 2. Comparative analysis

Reference and publication with year	Objectives of the research	Cluster trace dataset used	Findings of the research	Limitations/future work
[18] Elsevier, 2024	To provide STWGEN, a unique approach that uses statistical trace data to generate workload, to solve this challenge.	Alibaba Trace 2018	STWGEN accurately simulates batch task resource utilization by integrating task-level data with machine-level challenges. Extensive experiments confirm the superiority of the STWGEN method over existing baselines.	The scope of our research will be expanded in the future to account for work dependencies in workforce development.
[19] ACM, 2024	This research aims to fill the gap by creating a method for fitting open-source workloads from FaaS benchmarking suites to operational FaaS workload traces while preserving their essential statistical aspects.	MS Azure Trace	It analyses the execution modes of Faas Rail and compares load representativity with Azure Functions, Huawei's production-scale traces, and existing serverless literature.	FaaSRAIL currently does not modify function inputs, resulting in consistent expected execution times. Experimenting with different inputs is a potential next step. Not all experiments benefit from scaling the input trace over time. Experiments analyzing sandbox caching strategies or predictive preallocation optimizations may benefit from FaaSRAIL's Minute Range mode, which does not scale request over time.
[20] ACM, 2023	The research below compares operational cloud block storage workloads using block-level I/O traces from Alibaba Cloud and Tencent Cloud Block Storage, analyzing billions of requests.	Alibaba Cloud & Tencent Cloud Block Storage	Cloud block storage workloads and prominent public block-level I/O workloads from corporate data centers at Microsoft Research Cambridge are compared to determine their similarities and differences. These objectives are achieved through the overall provision of sixteen detailed data points and six high-level data points, talking about load intensity, geographical patterns, and temporal patterns. The analysis focuses on how the findings influence load balancing in cloud block storage, the efficiency of caching, and the administration of clusters.	Trace analysis has various limitations. Because the traces collected by AliCloud and Tencent Cloud capture the response time differently than the traces captured by MSRC, it's impossible to do latency analysis for I/O requests during a real deployment. Both of them also do not expose which programs are running on a particular disk, thus leaving out the association of individual workloads of an application to their corresponding I/O pattern.
[21] IEEE, 2023	Harmony is a deep learning-driven ML cluster scheduler that optimizes performance by scheduling heterogeneous training tasks using parameter server or all-reduce architecture to minimize interference and increase performance (i.e., training completion time reduction).	Real ML Workloads in Kubernetes Cluster	The model is built for online job arrival using a parameter server or all-reduce architecture. A two-step learning mechanism is applied to learn the placement policy without an interference model of workload. Initially, the scarce historical data is used for training a reward neural network, and then DRL is trained on placement decision with reward samples obtained from the reward model. According to the evaluation of a Kubernetes cluster, Harmony reduces the completion time for jobs by 16%-42% compared with representative scheduling methods.	-

Table 2. Comparative analysis (*Continued*)

Reference and publication with year	Objectives of the research	Cluster trace dataset used	Findings of the research	Limitations/future work
[17] IEEE, 2022	This research introduces a three-way ensemble forecast for data center workload to enhance accuracy.	Google Cluster trace	Simulated annealing was utilised to obtain the workload split cutoff value for stable, volatile and jitter periods. The three-way decision idea was used to assign prediction models based on workload characteristics and a priori error prediction to improve accuracy. In workload prediction accuracy trials utilizing CPU load monitoring logs from Google cluster trace, TWD-RCPM surpasses ARIMA, NN, and DMASVR-3WD by 69.0%, 68.6%, and 72.6%.	Our approach is broad, scalable, and adaptable to many settings. Future work involves extending it. Further improvements of the system will include the analysis of memory, I/O, and disk usage, along with the introduction of new workload prediction models. Future research will also be based on energy consumption, with the prediction model applied to analyze energy savings in data center resource allocation.
[16] ORCA, Cardiff University	Our research analyzes workload traces to understand machine operations across platforms in Google data centers.	Google Cluster trace Version3	On average, 98.24% of machines recover/join in 47 minutes, with a standard variation of 4.5 hours. While the number of active machines fluctuates over the collecting period, the overall trend remains consistent if new machines are not included. Failures occur often, yet less than 2% of computers join the cluster.	Using machine state predictions, an active machine timeline can be generated. The model predicts the machine state with an accuracy of MAPE at 1.76% compared to the actual machine timeline.
[21] IEEE, 2019	Learn about the features of co-allocated online services and batch processes in an Alibaba Cloud production cluster with 1.3k servers.	Alibaba Cloud	This research presents insights that can help communities and data center operators in understanding workload characteristics, resource use, and failure recovery architecture.	Predict work failures and transfer instances to other computers to prevent problems.
[22] IEEE, 2022	According to our research, both cloud architectures have advantages and drawbacks, and their route inflations have distinct causes. The architecture of the ISP cloud can outperform the DC-cloud due to its advantages in data centers and large-scale network infrastructure.	Alibaba and Tencent Cloud Trace	The findings minimize latency by 11.0% to 30.5% for Alibaba, 9.8% to 18.6% for Tencent, and 54.1% to external destinations for CTYun. The results indicate that both cloud systems have opportunity for development, although the ISP-cloud may have better performance due to its inherited benefits from the ISP infrastructure.	The tools of measurement will be deployed on the platforms like OpenNetLab for launching queries from the network edge, such as a university, home, and cellular networks. The purpose is to do long-term measurement research tracking the evolution of cloud routing regulations, with the flattening trend of the Internet in consideration.

BIOGRAPHIES OF AUTHORS



Akash Patel     is an academic and researcher affiliated with the Devang Patel Institute of Advance Technology and Research (DEPSTAR), CHARUSAT University. His research interests span cloud computing, computer networking, internet of things (IoT), and image processing. He has contributed to multiple research publications in reputed journals and international conferences, focusing on emerging technologies such as edge computing, deep learning applications, and intelligent systems. He can be contacted at email: akashpatel.dit@charusat.ac.in.



Dr. Amit Nayak     is an associate professor at Charotar University of Science and Technology (CHARUSAT), with extensive experience in teaching and research in computer networks, cloud computing, IoT, data center networks, and wireless communication. He has made significant contributions to the research community through numerous publications in reputed journals and international conferences. He can be contacted at email: amitnayak.dit@charusat.ac.in.



Dr. Khushi Patel     is an assistant professor at Devang Patel Institute of Advanced Technology and Research, CHARUSAT. She has an extensive publication record with significant citations, reflecting her contributions to areas such as edge computing, blockchain-based systems, IoT security, cloud computing, and artificial intelligence. Her research also spans advanced topics like multipath TCP security, fine-grained access control, and AI-assisted secure IoT systems. She has actively contributed to both foundational and applied research, addressing real-world challenges in distributed and secure computing environments. She can be contacted at email: khushipatel.ce@charusat.ac.in.



Anand Patel     is affiliated with Dharmsinh Desai University, with research interests in Security, Machine Learning, and Natural Language Processing (NLP). He has contributed to multiple research publications, including studies on green cloud computing, RAG-based chatbot systems, blockchain applications in industrial automation, and IoT security. His recent work also includes advanced neural network applications for medical image analysis. His contributions reflect a strong inclination toward interdisciplinary research in intelligent and secure computing systems. He can be contacted at email: anandpatel.it@ddu.ac.in.