

Prediction of Workloads in Cloud using ARIMA-ANN

Suriya S.¹, Surya Arvinth M.²

¹Associate Professor, Department of Computer Science and Engineering, PSG College of Technology, Coimbatore, Tamil Nadu, India.

²PG Scholar, Department of Computer Science and Engineering, PSG College of Technology, Coimbatore, Tamil Nadu, India.

E-mail: ¹ss.cse@psgtech.ac.in, suriya84@gmail.com, ²23mz04@psgtech.ac.in

Abstract

This study introduces an innovative hybrid ARIMA-ANN model personalized for cloud workload prediction. Unlike existing models that focus solely on linear or nonlinear patterns, the approach combines the strengths of ARIMA for time-series linear trends and ANN for nonlinear data complexities. This integration ensures higher accuracy, as validated using the MIT Supercloud dataset. The methodology leverages data pre-processing, sensitivity analysis, and advanced validation techniques, demonstrating improved accuracy in scenarios of high workload variability. This model supports cloud providers in resource optimization and dynamic load management.

Keywords: ARIMA, ANN, Load Balancing, Time Series data, MIT Supercloud

1. Introduction

Cloud computing, a foundation of distributed computing, relies on dynamic resource management to deliver optimal Quality of Service (QoS) and meet Service Level Agreements (SLAs). Infrastructure as a Service (IaaS) providers face challenges managing volatile CPU and memory demands. Accurate workload prediction ensures resource optimization, reducing over-provisioning and cost.

Existing methods, such as ARIMA and ANN, handle specific workload patterns but lack the versatility required for diverse and unpredictable workloads. This study introduces a

hybrid ARIMA-ANN model, blending linear forecasting with nonlinear learning capabilities. The novelty lies in addressing complex workload scenarios, validated using a real-world dataset. This approach achieves superior prediction accuracy and robustness, essential for dynamic resource allocation in cloud systems.

2. Related Work

The study focuses on the prediction models for efficient cloud resource provisioning, emphasizing the role of auto-scaling algorithms and comparing QoS parameters between conventional and efficient provisioning methods. It highlights designing prediction mechanisms and utilizing methods for workload determination [1]. In their research study [2], the authors discuss the shift to cloud computing, the rise of edge computing to complement cloud deployments, and challenges in orchestrating edge-cloud applications. It underscores the significance of machine learning in workload characterization, component placement, and application elasticity, classifying algorithms into supervised and unsupervised categories. This study [3] uses a hybrid ARIMA-ANN model to forecast future CPU and memory utilization in cloud resource provisioning. The model detects linear and nonlinear components in cloud traces, while the artificial neural network (ANN) uses residuals. The Savitzky-Golay filter finds a range of forecast values, reducing forecasting error by introducing a range of values. The accuracy of the prediction is tested using statistical analysis using Google's 29-day trail and BitBrain. Cloud computing is gaining popularity, requiring accurate prediction of computing resource usage for efficient management. However, excessive costs can be a concern. The study [4] presents a novel approach that uses data-driven prediction algorithms to generate short- and long-term cloud resource usage predictions. The solution readjusts to different load characteristics and usage changes. Preliminary tests showed 36% better prediction quality, and real-life historical data showed 9.28% to 80.68% better prediction quality. This research study [5] is a systematic survey and comparative study of machine learning-driven cloud workload prediction models. It discusses the importance of predictive resource management, operational design, motivation, and challenges. The study classifies different prediction approaches into five categories, focusing on theoretical concepts and mathematical functioning. The proposed method in [6] presents COSCO2, a workload prediction framework optimized with a sheep flock optimization algorithm. It accurately

predicts workloads, reducing energy consumption and outperforming existing methods for datasets like NASA and Saskatchewan HTTP traces. The study [7] proposes an autonomic prediction suite to improve the accuracy of cloud computing's auto-scaling system. It suggests that selecting the right time-series prediction algorithm based on the incoming workload pattern can increase the prediction accuracy. The research conducts theoretical investigations and empirical validations, designing a self-adaptive prediction suite that automatically chooses the most suitable algorithm. The research study [8] discusses designing cloud client prediction models using machine learning, identifying Support Vector Regression (SVR) as the best model for non-linear workload patterns. The study [9] addresses challenges in volatile resource demands and proposes a meta-algorithm for algorithm selection based on past performance, with insights into dynamic regret and optimal solutions. The method proposed in [10] introduces a DCRNN, a deep learning model for accurate workload prediction, to improve forecasting accuracy and minimize the error between the predicted and the actual workloads. The method put forth in [11] describes a framework for provisioning virtual machines using Kalman filter-based data preprocessing, enhancing service quality by reducing provisioning latency. Experimental results show that the framework reduces latency and improves cloud service quality, using Alicloud as an experimental infrastructure. The study presents RPPS, a Cloud Resource Prediction and Provisioning scheme, which predicts future demand and performs proactive resource provisioning for cloud applications. It uses the ARIMA model, combines coarse-grained and fine-grained resource scaling, and adopts a VM-complementary migration strategy. The prototype has high prediction accuracy and good resource scaling, making it a valuable solution for enterprises facing demand fluctuations in cloud data centers [12]. The research [13] presents a machine learning-based solution for cloud resource optimization, using anomaly detection to reduce costs while maintaining QoS. It emphasizes the importance of initial training and future research in model reuse without system duplication. The study [14] focuses on QoS-based resource provisioning to optimize allocation, reduce execution time and costs, and improve resource scheduling. It emphasizes the role of workload analysis in understanding QoS requirements. The method put forth in [15] presents the RRAU approach for resource management, improving metrics such as Job Completion Time (JCT) and monetary cost compared to conventional methods. It emphasizes optimizing VM runtimes and asset utilization. The researchers in [16] propose a capacity planning approach utilizing hybrid spot instances for improved utilization and reliability. It enhances throughput during out-of-bid situations and suggests future integration of reactive modules for better QoS impact.

analysis. The study in [17] introduces a container-based video surveillance cloud platform with predictive resource allocation, optimizing service density and cost prediction. It demonstrates the effectiveness of Docker technology over VM-based platforms in supporting microservices like video on demand. This article [18] presents a method for elastic resource provisioning using data clustering in cloud service platforms. The method consists of tasks clustering, task prediction, dynamic resource provisioning, and scheduling. It partitions tasks based on similarity and forecasts future tasks. The method achieves high guarantees, resource utilization, and total energy consumption in the Google cloud traces dataset. The study [19] addresses monitoring challenges in hybrid clouds, proposing a solution that automates management across hybrid environments. The study [20] introduces PCA-BPN for estimating simulation workloads in cloud manufacturing, achieving superior accuracy in execution time prediction, especially in factory simulation contexts.

3. Comparison of Various Computing Types

3.1 Centralized Computing

Centralized computing relies on a central server for all computations and data storage. This model is simple to implement and administer, with centralized management and security, but it struggles with scalability and is prone to single points of failure. Examples include traditional mainframe systems.

3.2 Distributed Computing

Distributed computing involves interconnected nodes working collaboratively. It offers scalability and reliability with fault tolerance and optimized resource usage. However, performance can be limited by network bandwidth. Examples include peer-to-peer networks and Hadoop.

3.3 Grid Computing

Grid computing coordinates geographically dispersed resources to function as a unified system. It allows scalability and redundancy but may face latency and bandwidth challenges. Control is typically centralized, with shared resource utilization. Examples include SETI@home and the Globus Toolkit.

3.4 Cluster Computing

Cluster computing connects multiple nodes to work together, providing fault tolerance and scalability by adding nodes. Management can be centralized or distributed, with shared resources tailored for specific tasks. Examples include HPC and Beowulf clusters.

3.5 Cloud Computing

Cloud computing provides elastic, on-demand resource sharing through virtualized environments. Resources are distributed but centrally managed, with data stored across multiple regions. Prominent platforms include AWS, Google Cloud Platform, and Microsoft Azure.

4. Types of Workload

4.1 Static Workload

Static workloads involve a fixed number of requests, with constant memory, processor capacity, and bandwidth. These workloads follow consistent usage patterns and do not scale dynamically. Examples include web servers and email services. Figure 1 represents the pattern of static workload.

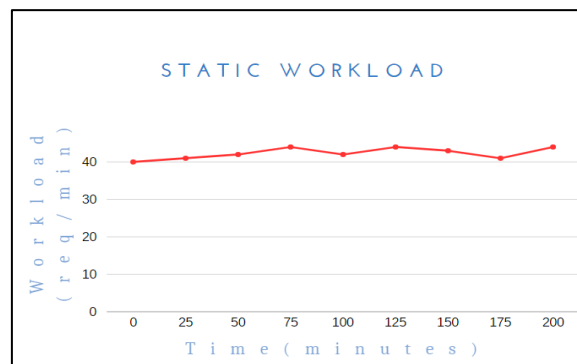


Figure 1. Static Workload Pattern

4.2 Periodic Workload

Periodic workloads involve recurring cycles, such as daily, weekly, or seasonal peaks. They require scalability to handle peak loads efficiently while avoiding resource

underutilization during off-peak periods. These workloads demand flexible resource allocation. Figure 2 represents the periodic workload pattern.

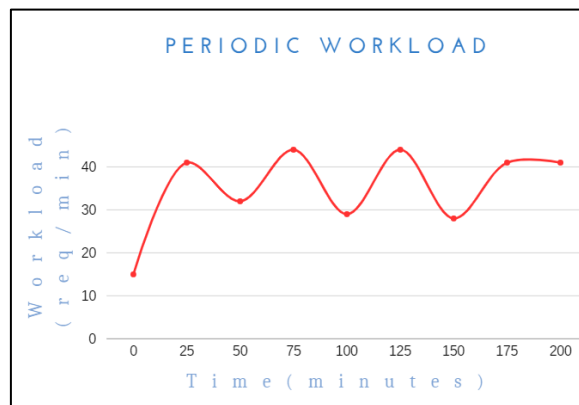


Figure 2. Periodic Workload Pattern

4.3 Unpredictable Workload

Unpredictable workloads are dynamic and lack regular patterns, making resource planning challenging. These workloads experience random surges due to factors like unexpected traffic or environmental conditions. Cloud providers must address variability through dynamic resource allocation. Figure 3 depicts the unpredictable workload pattern



Figure 3. Unpredictable Workload Pattern

4.4 Continuously Changing Workload

Continuously changing workloads exhibit gradual increases or decreases in resource demand over time. Resource allocation adjusts dynamically to align with the evolving scope and scale of tasks. Figure 4 represents continuously changing workload pattern.

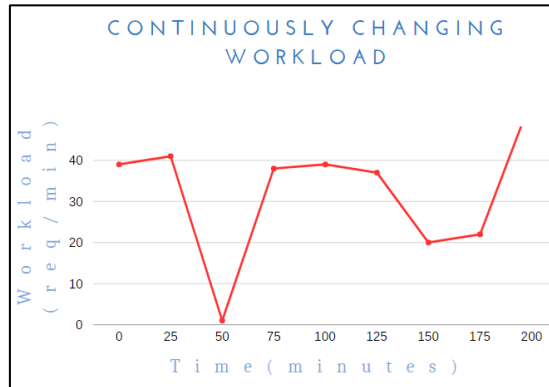


Figure 4. Continuously Changing Workload Pattern

4.5 Once in a Lifetime Workload

This workload represents peak utilization events unlikely to recur, requiring manual provisioning or decommissioning of resources. Such workloads often involve one-time high resource consumption. Figure 5 represents the Once in a Lifetime Workload pattern

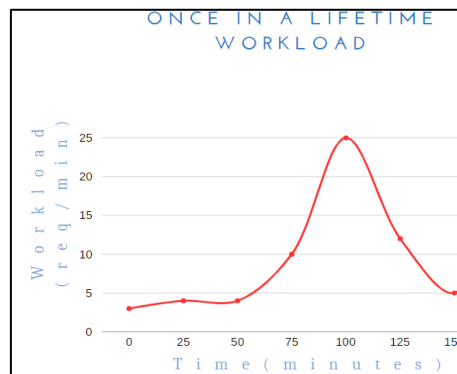


Figure 5. Once in a Lifetime Workload Pattern

5. Proposed Methodology

5.1 Time series forecasting Models

Time series can be defined with data for $y(t)$ it's a series of dispersed data values taken at successive time intervals and t is time and time is a variable that increases continuously from zero to one and one to two in equal intervals. It is a manner that they arrange the observations right from the observation that is observed most to the observation observed least. As usual,

the single variable has a time series while more than one variable is encoded in the multivariate time series. However, time series could be of two types: and these are: Continuous time series and Discrete time series. When the time series is continuous, this is because the observations the events occur at a continuous time point, while for the data recorded on discrete time series, the observations are taken at equal time steps. It has four components: These are seasonal working shifts, regular working shifts, at random working shift and periodic working shifts. On the basis of classification of numerous time series, economists divided them into four sub-series. They are fashionable, oscillatory, periodic, and non-periodic and are feasible from the observed facts. With reference to the impacts of the four elements, as well as the multiplicative and the additive models, use the time series computations [3].

Multiplicative Model

$$Y(t) = T(t) * S(t) * C(t) * I(t)$$

Additive Model

$$Y(t) = T(t) + S(t) + C(t) + I(t)$$

T(t)-Trend Component

S(t) Seasonal Component

C(t) Cyclic Component

I(t)-Irregular component

5.2 ARIMA – Autoregressive Integrated Moving Average

Autoregression represents a concept suggesting that a variable that varies is regressed on preceding or lag values. Integrated (I) refers to the level of the simple transformation of an immature form of data in the time series in the smoothing process that aims at replacing each of the value of time series with the variation of the data value from the preceding data value. MA employs the stochastic component of an observation in a moving average model fitted to lag observations.

The parameters of this model are:

P: the number of lag observations in the model more precisely identified as the lag order.

D: the number of times the raw observations are differenced; also the degree of differencing.

Q: the second parameter that could be changed is the size of the moving average window or simply the order of the moving average.

$$y(t) = \theta_0 + \phi_1 y_{t-1} + \phi_2 y_{t-2} \dots + \phi_p y_{t-p} + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} \dots + \theta_q \varepsilon_{t-q}$$

y_t and ε_t are the actual value and the random error at time period t respectively, $\phi_i = 1, 2, \dots, p$ and $\theta_i = 1, 2, \dots, q$ are the model parameters [3].

5.3 ANN

The ANN structure is very similar to the neuron structure of the human brain. ANNs approximate the various nonlinearities in the data. To model the time series data, the relationship between the output (y_t) and the inputs ($y_{t-1}, y_{t-2}, \dots, y_{t-p}$) is expressed as

$$y_t = W_0 + \sum_{j=1}^q w_j \cdot g(w_{0j} + \sum_{i=1}^p w_{ij} \cdot y_{t-i}) + \varepsilon_t$$

w_{ij} ($i = 0, 1, 2, \dots, p; j = 1, 2, \dots, q$) represents the model parameters also called the connection weights; p is the number of input nodes and q is the number of hidden nodes and g is the transfer function of the hidden layer. Indeed, the ANN model can be considered almost a nonlinear autoregressive model. This model's coefficient is then computed from the input of the ANN where the latter is fed with the observed data sequence and the latter trained with this sample sequence. The performance is checked on a sample after the training process is over [3].

5.4 Data Collection and Preprocessing

Dataset: MIT Supercloud Dataset

Source and Samples: Kaggle - 96,893 entries with 23 columns.

Features: Memory usage, and GPU memory utilization.

Preprocessing:

- Empty values are replaced by NaN and then NaN value elimination.
- Time-series transformation for sequential analysis.
- Scaling and normalization for feature consistency.

5.5 Model Training and Optimization

- All data is scaled to a range between 0 and 1 using MinMaxScaler. This is essential for neural networks to ensure uniform weight updates during training.
- ANN uses lagging to transform the time-series data into a supervised learning problem:
 - X (Features): The data at time t.
 - y (Targets): The data at time t+1.
- Loss function: Mean Squared Error (MSE).
- Optimizer: Adam optimizer with learning rate 0.05. Adam (Adaptive Moment Estimation) combines the benefits of momentum and RMSprop optimizers. It dynamically adjusts learning rates for each parameter, making it efficient for non-stationary objectives.
- Epochs: 100.
- Batch size: 32.

5.6 Evaluation Metrics and Tools

The code evaluates the models using metrics and tools specific to the methods applied. For the ANN model, Mean Squared Error (MSE) is used to measure the average squared difference between predicted and actual values, computed using `sklearn.metrics.mean_squared_error`, with lower MSE indicating better performance. For the ARIMA model, evaluation relies on statistical summaries provided by `statsmodels.tsa.arima.model`. ARIMA, which assess model fit and complexity, includes metrics like Akaike Information Criterion (AIC), Bayesian Information Criterion (BIC), and Log-Likelihood. Tools like `sklearn.neural_network.MLPRegressor` handles ANN training,

while `sklearn.preprocessing.MinMaxScaler` scales features to improve model convergence. Visualization is performed using `Matplotlib` to plot actual vs. predicted values for ANN. Additional metrics such as Root Mean Squared Error (RMSE), Mean Absolute Error (MAE), or R^2 Score could enhance the analysis further, providing more comprehensive insights into model performance.

6. Architecture

The ARIMA-ANN model, to predict the area of confidence values for CPU and memory utilization is presented in Figure. 6. The measurement for the analysis was done using the Google cluster data collection. Note that pre-processing of time series data requires null elimination. Finally, the time series is passed to the ARIMA model, where the implementation of predictions of CPU and memory usage of each request for the ensuing N requests is executed using the data. ARIMA cannot model the residues which are described as the nonlinear components of the model. The residuals, which were collected, are fed along with the upcoming data into the network model. It constructs the model of the prognosis of CPU and memory demands.

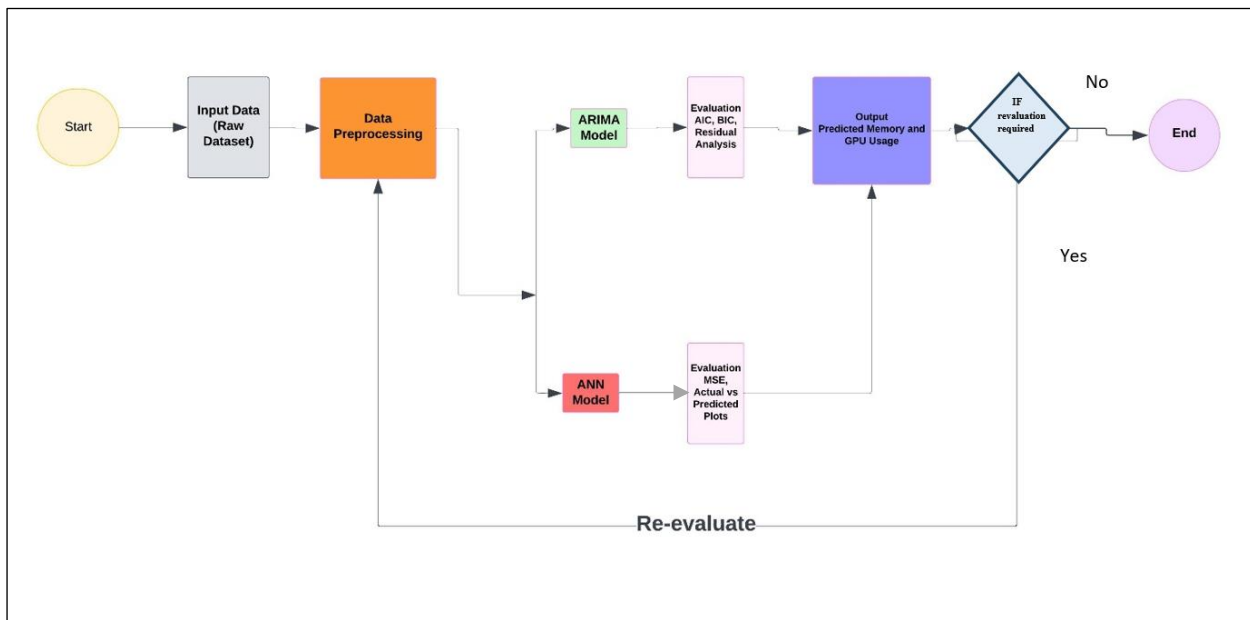


Figure 6. Architecture of Hybrid ARIMA – ANN Model [3]

7. Experimental Setup

ISSN: 2582-1369

Figure 7 Sample Data

8. Results and Evaluation

8.1 Prediction of Memory Utilization

The summary of the model used in this implementation is given in Figure 8. The Memory Utilization actual vs predicted plot is given in Figure 9.

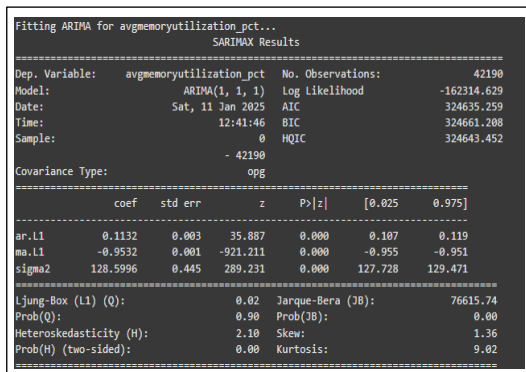


Figure 8 Model Summary

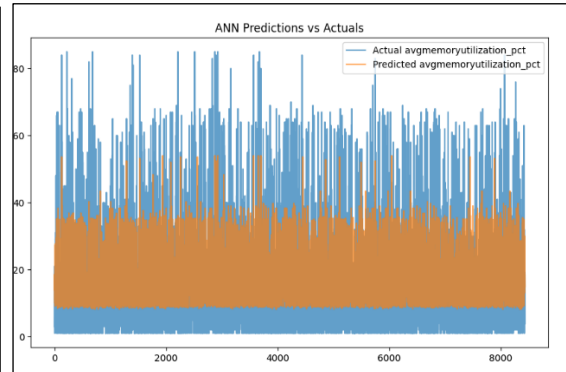


Figure 9 Memory Actual vs Predicted

The ANN mean square error between actual and predicted values are calculated and it was approximately 0.04122451258775828

8.2 Prediction of Memory

The summary of the model used in this implementation is given in Figure 10. The GPU actual vs predicted plot is given in Figure 11.

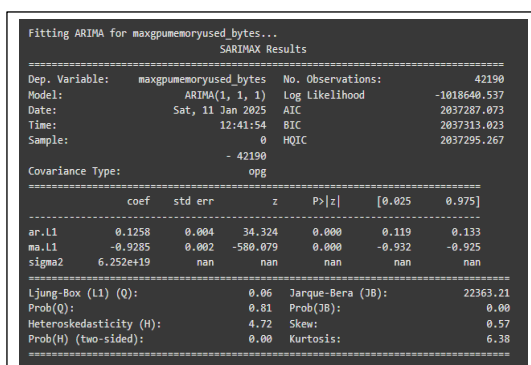


Figure 10. Model Summary

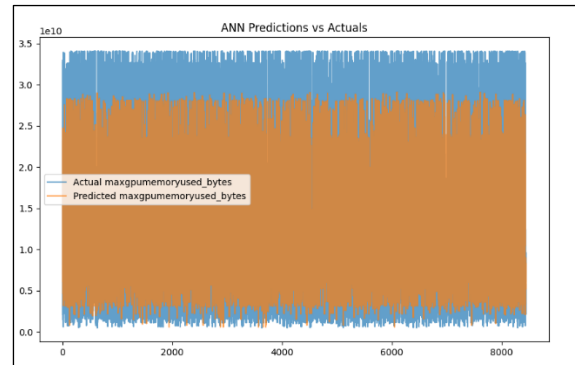


Figure 11. GPU Actual vs Predicted

9. Conclusion

However, to provide high QoS with optimum resource usage in clouds, it is imperative to predict workloads correctly. The workload characteristics of the hybrid model, which include linear and non-linear, afford cloud providers the tools necessary to manage inventories a priori and increase preparedness for varying workloads and demand. However, in case of any change in the workload pattern, it would be necessary to bring the model into the parameter estimation process to keep an accurate forecast. This open and flexible strategy also reduces the dangers of under-providing services, including service downtime, high energy expenses, and client loss. By using dynamic sliding windows and giving the latest data a higher weight, predictions are enhanced, supplying providers of cloud services with a powerful instrument to regulate resources for dynamic workloads.

References

- [1] Gadhavi, L. J., and M. D. Bhavsar. "Prediction-Based Efficient Resource Provisioning and Its Impact on QoS Parameters in the Cloud Environment." *International Journal of Electrical and Computer Engineering (IJECE)* 8, no. 6 (2018): 5359–5370.
- [2] Duc, T. L., R. G. Leiva, P. Casari, and P. O. Östberg. "Machine Learning Methods for Reliable Resource Provisioning in Edge-Cloud Computing: A Survey." *ACM Computing Surveys (CSUR)* 52, no. 5 (2019): 1–39.
- [3] Devi, K. L., and S. Valli. "Time Series-Based Workload Prediction Using the Statistical Hybrid Model for the Cloud Environment." *Computing* 105, no. 2 (2023): 353–374.
- [4] Nawrocki, P., P. Osypanka, and B. Posluszny. "Data-Driven Adaptive Prediction of Cloud Resource Usage." *Journal of Grid Computing* 21, no. 1 (2023): 6.
- [5] Saxena, D., J. Kumar, A. K. Singh, and S. Schmid. "Performance Analysis of Machine Learning Centered Workload Prediction Models for Cloud." *IEEE Transactions on Parallel and Distributed Systems* 34, no. 4 (2023): 1313–1330.
- [6] Karthikeyan, R., V. Balamurugan, R. Cyriac, and B. Sundaravadivazhagan. "COSCO2: AI-Augmented Evolutionary Algorithm Based Workload Prediction Framework for

Sustainable Cloud Data Centers." *Transactions on Emerging Telecommunications Technologies* 34, no. 1 (2023): e4652.

- [7] Nikraves, A. Y., S. A. Ajila, and C. H. Lung. "An Autonomic Prediction Suite for Cloud Resource Provisioning." *Journal of Cloud Computing* 6 (2017): 1–20.
- [8] Bankole, A. A., and S. A. Ajila. "Cloud Client Prediction Models for Cloud Resource Provisioning in a Multitier Web Application Environment." In *2013 IEEE Seventh International Symposium on Service-Oriented System Engineering, USA* 156–161. IEEE, 2013.
- [9] Comden, J., S. Yao, N. Chen, H. Xing, and Z. Liu. "Online Optimization in Cloud Resource Provisioning: Predictions, Regrets, and Algorithms." *Proceedings of the ACM on Measurement and Analysis of Computing Systems* 3, no. 1 (2019): 1–30.
- [10] Al-Asaly, M. S., M. A. Bencherif, A. Alsanad, and M. M. Hassan. "A Deep Learning-Based Resource Usage Prediction Model for Resource Provisioning in an Autonomic Cloud Computing Environment." *Neural Computing and Applications* 34, no. 13 (2022): 10211–10228.
- [11] Hu, Y., B. Deng, and F. Peng. "Autoscaling Prediction Models for Cloud Resource Provisioning." In *2016 2nd IEEE International Conference on Computer and Communications (ICCC), United States* 1364–1369. IEEE, 2016.
- [12] Fang, W., Z. Lu, J. Wu, and Z. Cao. "RPPS: A Novel Resource Prediction and Provisioning Scheme in Cloud Data Center." In *2012 IEEE Ninth International Conference on Services Computing, Honolulu, HI, USA* 609–616. IEEE, 2012.
- [13] Nawrocki, P., and P. Osypanka. "Cloud Resource Demand Prediction Using Machine Learning in the Context of QoS Parameters." *Journal of Grid Computing* 19, no. 2 (2021): 20.
- [14] Singh, S., and I. Chana. "Q-Aware: Quality of Service Based Cloud Resource Provisioning." *Computers & Electrical Engineering* 47 (2015): 138–160.
- [15] Baldoss, P., and G. Thangavel. "Optimal Resource Allocation and Quality of Service Prediction in Cloud." *Computers, Materials & Continua* 67, no. 1 (2021).

- [16] Sadashiv, N., S. D. Kumar, and R. S. Goudar. "Cloud Capacity Planning and HSI Based Optimal Resource Provisioning." In 2017 Second International Conference on Electrical, Computer and Communication Technologies (ICECCT), Coimbatore, India 1–6. IEEE, 2017.
- [17] Zhang, H., H. Ma, G. Fu, X. Yang, Z. Jiang, and Y. Gao. "Container Based Video Surveillance Cloud Service with Fine-Grained Resource Provisioning." In 2016 IEEE 9th International Conference on Cloud Computing (CLOUD), San Francisco, CA, USA 758–765. IEEE, 2016.
- [18] Fei, B., X. Zhu, D. Liu, J. Chen, W. Bao, and L. Liu. "Elastic Resource Provisioning Using Data Clustering in Cloud Service Platform." *IEEE Transactions on Services Computing* 15, no. 3 (2020): 1578–1591.
- [19] Naik, V. K., K. Beaty, N. Vogl, and J. Sanchez. "Workload Monitoring in Hybrid Clouds." In 2013 IEEE Sixth International Conference on Cloud Computing, 816–822. Santa Clara, CA, USA IEEE, 2013.
- [20] Chen, T. "A PCA-BPN Approach for Estimating Simulation Workload in Cloud Manufacturing." In 2015 Seventh International Conference on Ubiquitous and Future Networks, Sapporo, Japan 826–830. IEEE, 2015.