



GRAPH-NEURAL-NETWORK-DRIVEN PROACTIVE RESOURCE ALLOCATION FOR CLOUD WORKLOADS: A PREDICTIVE GRAPH-MODELING FRAMEWORK

K V S R Pavan Kumar^{1*} Dr. Bechoo Lal² Dr. Bysani Venkata Srinivasulu³

¹Research Scholar, Department of Computer Science, Shri JYT University, Rajasthan, India

²Research Guide, Shri JYT University, Rajasthan, India

³Research co Guide, Department of Computer Science and Engineering, Vignan institute of Technology and Science, Hyderabad, India

*Corresponding Author: K V S R Pavan Kumar, Email: kvsrraghava@gmail.com

ABSTRACT

Cloud infrastructures face high variability in workload demands, making static or reactive resource allocation inefficient: either resources are over-provisioned (wasting cost) or under-provisioned (violating SLAs). This paper proposes a novel framework that models the cloud infrastructure and workload interactions as a graph, and employs a Graph Neural Network (GNN) to forecast future resource demands for nodes and edges in the graph. By leveraging this prediction, the framework drives a proactive resource allocator that adjusts provisioning and workload placements ahead of time. Experimental results based on realistic workload traces demonstrate that the GNN-based predictor achieves significantly higher accuracy compared to conventional time-series methods, enabling the allocator to reduce cost by X% and SLA violation rate by Y%. The contributions include (1) a graph abstraction of the cloud system and workloads; (2) a spatio-temporal GNN architecture for demand prediction; and (3) a resource allocation engine integrating prediction with provisioning decisions.

Keywords: Graph Neural Network, Predictive Resource Allocation, Cloud Workload Forecasting, Proactive Autoscaling, Resource Graph Modeling.

I. INTRODUCTION

Cloud computing has revolutionized the delivery of computational services by offering scalable, on-demand access to shared resources such as storage, computing, and networking [1]. As organizations increasingly migrate mission-critical applications to cloud and edge infrastructures, efficient resource allocation becomes essential to ensure optimal utilization while meeting Service Level Agreements (SLAs) [2]. However, the highly dynamic and unpredictable nature of workloads makes static or rule-based resource management approaches inefficient and reactive [3]. Over-provisioning leads to unnecessary operational costs, while under-provisioning degrades performance and increases SLA violations [4].

Traditional workload forecasting techniques—including time-series models such as ARIMA and LSTM—capture temporal patterns of

individual resources but fail to consider spatial dependencies between interconnected cloud entities [5]. In real-world cloud infrastructures, services and applications interact as interdependent components connected via communication and data flow networks. Ignoring these relationships limits prediction accuracy and hinders proactive decision-making [6]. Consequently, a paradigm shift is required to model cloud infrastructures as graph-structured systems, where nodes represent compute resources and edges represent communication dependencies or workload interactions.

Graph Neural Networks (GNNs) have recently emerged as a powerful deep learning architecture for processing graph-structured data, offering an ability to capture relational dependencies and complex topological information through message passing [7]. Unlike



traditional deep learning models, GNNs learn representations by aggregating neighborhood information, making them suitable for modeling interdependent systems such as cloud services, microservices, and distributed resource environments [8]. By leveraging this property, GNNs can predict workload patterns across connected nodes more accurately, enabling proactive and context-aware resource allocation [9].

Recent works have applied GNNs to a range of cloud computing problems, including anomaly detection, load balancing, and service dependency analysis, demonstrating notable improvements in predictive performance [10]. Li et al. proposed a graph-evolution-based workload prediction model that dynamically captures temporal variations across cloud infrastructures, outperforming conventional LSTM predictors [11]. Similarly, Nguyen et al. highlighted the potential of GNNs for microservice-based cloud applications, showcasing their capability to model dependencies between distributed components [12]. In reinforcement learning contexts, GNNs have been used to improve task scheduling and resource scaling decisions in multi-agent systems, further validating their effectiveness in dynamic resource management scenarios [13].

Despite these advances, existing studies predominantly address either prediction or allocation independently, with limited integration between the two. Moreover, most predictive models do not exploit the graph nature of cloud environments, leading to suboptimal decisions when scaling or redistributing workloads [14]. To overcome these limitations, this paper proposes a Graph Neural Network (GNN)-based workload prediction and resource allocation framework that models cloud infrastructure as a graph to jointly forecast future workloads and guide resource provisioning decisions.

The proposed framework operates in two phases:

Workload Prediction Phase – A spatio-temporal GNN model captures both structural (graph) and temporal dependencies to predict future resource demands with high accuracy.

Resource Allocation Phase – A proactive allocation engine uses the predicted demands to perform cost-efficient scaling and workload distribution while ensuring SLA compliance.

By integrating predictive analytics with graph-based learning, this research aims to reduce resource wastage, minimize SLA violations, and enable intelligent, proactive resource management for next-generation cloud and edge systems. The experimental results demonstrate that the proposed model achieves superior prediction accuracy and cost efficiency compared to traditional methods, marking a significant advancement in AI-driven cloud orchestration [15].

II. LITERATURE REVIEW

Efficient workload prediction and resource allocation in cloud environments have been extensively explored, yet many traditional techniques still rely on time-series and statistical models that cannot effectively capture the complex spatial and temporal dependencies present in distributed infrastructures [16]. Early studies employed models such as ARIMA and Holt-Winters to predict resource usage trends, providing a foundation for predictive autoscaling [17]. However, such methods assume stationarity and independence between data points, making them unsuitable for highly dynamic and interconnected workloads.

The evolution of deep learning (DL) introduced new forecasting capabilities through models like Recurrent Neural Networks (RNNs), Long Short-Term Memory (LSTM), and Temporal Convolutional Networks (TCNs) [18]. These architectures demonstrated superior accuracy in learning nonlinear temporal relationships, enabling better short-term workload prediction.



For instance, Al-Dulaimy et al. achieved improved CPU utilization forecasts by applying deep LSTM-based models on Google cluster traces [19]. Nevertheless, these models still fail to incorporate topological relationships between nodes, which are crucial for multi-tier and microservice-based applications.

With the emergence of Graph Neural Networks (GNNs), researchers began representing cloud infrastructures as graphs, where nodes denote compute units (VMs, containers) and edges represent dependencies such as communication links or task flow [20]. Wu et al. provided a comprehensive survey on GNNs and emphasized their capability to learn from both graph structure and node features, making them suitable for cloud systems [21]. Subsequent works leveraged Graph Convolutional Networks (GCNs) and Graph Attention Networks (GATs) for analyzing service dependencies, anomaly detection, and network traffic prediction [22].

In 2023, Nguyen et al. developed a GNN-based framework for microservice dependency analysis, achieving enhanced accuracy in identifying resource bottlenecks and service-level degradation [23]. Similarly, Li et al. proposed a Graph Evolution Learning (GEL) model that captured temporal evolution of workload graphs, leading to improved resource utilization and lower energy consumption [24]. These studies demonstrated that GNNs can outperform traditional sequence models in predicting correlated workloads across distributed nodes.

Parallel to this, several works have explored reinforcement learning (RL) and multi-agent learning to improve decision-making in dynamic resource environments. The integration of RL with predictive analytics has enabled systems to dynamically adjust provisioning policies in real time [25]. Kaur et al. combined deep RL and clustering for multi-cloud resource optimization, showing improved scalability and SLA adherence [26]. However, RL systems typically

rely on external prediction modules and lack explicit modeling of inter-node dependencies, creating a gap between prediction accuracy and decision effectiveness.

In 2024, Zhou et al. presented a comprehensive review of Deep Reinforcement Learning (DRL) methods for resource scheduling, concluding that hybrid systems combining GNN-based prediction and RL-based control could provide significant operational benefits [27]. Meanwhile, Kayalvili et al. proposed a prediction-enabled RL allocator, where GNN-driven predictions were used to guide proactive scaling decisions, reducing cost by 25% compared to heuristic approaches [28]. This fusion of graph-based learning and control signals the next phase of intelligent cloud resource management.

Another research direction involves spatio-temporal modeling, where both spatial (graph connectivity) and temporal (time progression) dynamics are captured simultaneously. Models such as ST-GCN, T-GAT, and Spatio-Temporal Graph Attention Networks (STGAT) have shown effectiveness in traffic and energy forecasting domains and are now being adapted for cloud systems [29]. The proposed framework in this study builds upon these advances by employing a Spatio-Temporal GNN architecture that simultaneously learns node-level workload patterns and inter-node influence propagation.

Overall, existing literature reveals three key gaps that this study aims to address:

1. Limited graph representation: Most existing models neglect the graph-structured relationships inherent in cloud infrastructures.
2. Weak coupling between prediction and allocation: Prediction models and resource allocation engines are often developed independently.
3. Reactive decision-making: Few frameworks employ proactive, learning-based approaches that anticipate demand and adjust resources preemptively.



To bridge these gaps, the proposed GNN-based workload prediction and resource allocation framework integrates a graph representation of cloud topology with a spatio-temporal neural model for accurate forecasting, feeding its outputs into an adaptive allocation engine for proactive decision-making. This unified approach aims to optimize cost, utilization, and SLA adherence simultaneously, advancing the state-of-the-art in intelligent cloud management systems [30].

B.V. Srinivasulu's works on intelligent multimedia processing and mental-health analytics collectively demonstrate how adaptive computational techniques improve both data efficiency and behavioral prediction. In "Enhancing Lossless Image Compression through Smart Partitioning, Selective Encoding, and Wavelet Analysis," the author proposes a hybrid compression framework that partitions images based on structural regions, applies selective encoding to reduce redundancy, and leverages wavelet transforms for high-fidelity representation, ultimately achieving improved compression ratios without compromising visual quality [31]. Extending similar intelligence-driven optimization to the psychological domain, his study "Context-Aware Machine Learning Approaches for Enhanced Depression Prediction from Social Media Across Dynamic Conditions" integrates contextual signals—temporal patterns, linguistic sentiment shifts, and user-interaction cues—into machine-learning models, enabling more accurate and adaptable identification of depressive indicators across changing online behaviors [32]. Together, these studies illustrate how context awareness and intelligent feature processing can advance both technical domains of image compression and human-centric applications like mental-health prediction. [31][32].

The two conference studies demonstrate how AI and machine learning are reshaping both domain-specific problem solving and socio-

economic decision-making. The work "Redefining Agricultural Disease Management Using AI and ML" showcases how advanced image analytics, pattern recognition, and predictive modeling can detect crop diseases at early stages, enabling proactive interventions and improving agricultural productivity through data-driven decision support [33]. Complementing this technological application in agriculture, the study "The Role of AI-Enhanced Financial Literacy and HR Marketing in Personal Investment Decisions: A Quantitative Analysis of Urban and Rural Households" analyzes how AI-powered advisory tools, personalized financial-literacy modules, and targeted HR marketing strategies influence investment awareness and decision quality across diverse demographics, revealing the potential of intelligent systems to bridge knowledge gaps and promote informed financial behavior [34]. Together, these contributions highlight the broad societal impact of AI in both sustainable agriculture and financially inclusive economic development. [33][34].

III. SYSTEM MODEL AND METHODOLOGY

3.1 Overview of the Proposed Framework

The proposed framework aims to develop an intelligent, proactive workload prediction and resource allocation system using Graph Neural Networks (GNNs). The central idea is to represent the cloud infrastructure as a dynamic graph, where computing entities (e.g., virtual machines, containers, or microservices) are modeled as nodes, and their communication, data exchange, or resource dependencies are modeled as edges.

The framework operates in two sequential phases:

1. **Workload Prediction Phase:** A **spatio-temporal GNN** predicts future resource demands across all nodes by analyzing both structural (topological) and temporal patterns.



2. Proactive Resource Allocation Phase:

The predicted workload values are used to allocate or migrate resources optimally before demand surges occur, minimizing cost and maintaining SLA compliance.

This dual-phase design ensures that the system not only learns workload behavior but also proactively adjusts resources to achieve optimal utilization and stability.

3.2 Graph Representation of Cloud Infrastructure

The cloud environment is modeled as a time-evolving directed graph

$$G_t = (V, E, X_t)$$

where:

- $V = \{v_1, v_2, \dots, v_n\}$ is the set of nodes, representing computing resources such as VMs or containers.
- $E = \{e_{ij}\}$ is the set of edges, where each edge e_{ij} denotes a dependency, communication, or data flow from node v_i to node v_j .
- $X_t \in \mathbb{R}^{n \times d}$ is the feature matrix at time t , where each node v_i has a feature vector x_i^t containing metrics such as CPU utilization, memory usage, network bandwidth, queue length, and SLA status.

Each graph snapshot G_t represents the state of the cloud system at time t . As workloads fluctuate, node features and even edge connections may change, producing a sequence of graphs $\{G_{t-k}, G_{t-k+1}, \dots, G_t\}$ over time.

The objective is to forecast future resource demands $Y_{t+\Delta}$ (e.g., expected CPU and memory requirements) for each node given historical graphs:

$$Y_{t+\Delta} = f(G_{t-k}, \dots, G_t)$$

where $f(\cdot)$ is a predictive function learned by the GNN model.

3.3 Spatio-Temporal Graph Neural Network Design

The proposed predictive model combines graph convolution layers for spatial dependency learning with temporal sequence modules for time-series prediction.

3.3.1 Spatial Learning (Graph Convolution)

Each graph snapshot G_t is passed through a Graph Convolutional Network (GCN) or Graph Attention Network (GAT) to compute node embeddings:

$$H_t = \sigma(A \tilde{X}_t W_g)$$

where:

- H_t is the node embedding matrix at time t ,
- A is the normalized adjacency matrix with self-loops,
- W_g is the learnable weight matrix,
- σ is a nonlinear activation function (ReLU or LeakyReLU).

This operation aggregates information from each node's neighbors, capturing inter-node influences such as shared workloads, communication overhead, and service dependencies.

3.3.2 Temporal Dependency Modeling

The embeddings from multiple time steps $\{H_{t-k}, H_{t-k+1}, \dots, H_t\}$ are fed into a temporal learning module such as a Gated Recurrent Unit (GRU) or Temporal Convolutional Network (TCN):

$$Z_t = \text{GRU}(H_{t-k}, \dots, H_t)$$

This module captures temporal correlations and trends in workload evolution across nodes, enabling accurate short- and medium-term predictions.

3.3.3 Forecasting Head

Finally, a fully connected regression layer produces the predicted workload for each node:

$$\hat{Y}_{t+\Delta} = \text{FC}(Z_t)$$

where $\hat{Y}_{t+\Delta}$ represents predicted CPU, memory, or network utilization at future time $t+\Delta$.



3.3.4 Loss Function

Model training minimizes the prediction error using a Mean Squared Error (MSE) loss:

$$L = \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2$$

Optionally, regularization terms can be added to penalize overfitting and enforce smoothness between adjacent nodes in the graph.

3.4 Proactive Resource Allocation Engine

The Resource Allocation Engine (RAE) consumes the predicted workload vector $\mathbf{Y}^{t+\Delta}$ and generates optimal scaling and placement actions.

3.4.1 Decision Parameters

- **Scaling Decision:** Determines whether to scale up, down, or maintain current resources.
- **Migration Decision:** Identifies overloaded nodes and redistributes tasks to underutilized nodes.
- **Scheduling Policy:** Determines task priority and placement to optimize cost and response time.

3.4.2 Decision Model

The allocation process is formulated as a multi-objective optimization problem:

$$\min C = \alpha C_{\text{resource}} + \beta C_{\text{SLA}} + \gamma C_{\text{migration}}$$

subject to:

$$U_i(t+\Delta) \leq C_i(t+\Delta), \forall i \in V$$

where:

- C_{resource} = total operational cost,
- C_{SLA} = cost due to SLA violations,
- $C_{\text{migration}}$ = overhead cost due to VM or container migrations,
- U_i = predicted resource usage,

- C_i = available capacity.
Weights α, β, γ are dynamically tuned to balance cost and performance.

3.4.3 Control Logic

The engine operates in three stages:

1. **Prediction Input:** Receives predicted workload from the GNN model.
2. **Decision Generation:** Computes optimal scaling and migration actions.
3. **Execution and Feedback:** Applies decisions to the infrastructure and collects performance feedback for continuous refinement.

3.5 Adaptive Learning and Feedback Loop

The framework incorporates a closed-loop adaptive learning mechanism that continuously refines prediction accuracy and decision quality. After each allocation cycle:

1. Actual resource utilization and SLA performance are collected.
2. Deviations between predicted and actual workloads are calculated.
3. Model parameters are updated using recent data to adapt to new workload patterns.

This feedback mechanism ensures that the system evolves over time, handling seasonal workloads, dynamic user behavior, and hardware scaling events.

IV. EXPERIMENTAL SETUP AND RESULTS

4.1 Experimental Environment

The proposed GNN-based framework was implemented and evaluated in a simulated cloud computing environment that emulates the behavior of a large-scale infrastructure-as-a-service (IaaS) platform.

The experimental setup consists of the following key components:

Parameter	Specification
Simulation Tool	Python-based custom simulator built using TensorFlow and PyTorch
Number of Nodes	200 virtual machines (VMs) or containers
Graph Structure	Directed graph representing communication dependencies
Time Interval	5-minute sampling window
Prediction Horizon	30 minutes ahead



GNN Architecture	2 Graph Attention Layers (64 units each) + 1 GRU layer (128 units)
Training Epochs	100
Optimizer	Adam (learning rate = 0.001)
Batch Size	32
Hardware	Intel Xeon 3.2 GHz, 64 GB RAM, Ubuntu 22.04 environment

The cloud infrastructure was modeled as a time-evolving graph, where nodes represent compute entities (VMs) and edges represent communication links or resource dependencies. Each node maintains real-time metrics including CPU usage, memory consumption, I/O rates, and network throughput.

4.2 Baseline Models for Comparison

To validate the performance of the proposed model, results were compared against several baseline techniques commonly used in cloud workload prediction and resource allocation:

Method	Description
ARIMA	Traditional time-series forecasting model used for univariate resource prediction.
LSTM	Deep learning sequence model capturing temporal dependencies without spatial graph context.
1D-CNN + LSTM	Hybrid deep model combining convolutional feature extraction with temporal learning.
GNN (Static)	Graph model without temporal learning; processes only instantaneous snapshots.
Proposed GNN-ST	The proposed Spatio-Temporal Graph Neural Network integrating graph and sequence learning for proactive allocation.

Each baseline was integrated into the same resource allocation engine to ensure fair comparison of overall performance in terms of both prediction accuracy and resource management efficiency.

4.3 Evaluation Metrics

The evaluation of the framework was based on two major performance categories — prediction accuracy and resource allocation efficiency. The following metrics were used:

A. Prediction Metrics

Mean Absolute Error (MAE):

$$MAE = \frac{1}{N} \sum_{i=1}^N |y_i - \hat{y}_i|$$

Root Mean Square Error (RMSE):

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2}$$

The dataset includes both synthetic workloads generated through dynamic patterns (sinusoidal, bursty, and diurnal variations) and real-world traces from public sources such as the Google Cluster Trace dataset. This combination ensures a diverse workload profile for training and testing.

Mean Absolute Percentage Error (MAPE):

Measures prediction accuracy as a percentage error relative to actual usage.

B. Allocation and Efficiency Metrics

- **Resource Utilization (RU):** Average utilization of CPU and memory resources after allocation.
- **SLA Violation Rate (SVR):** Percentage of tasks failing to meet SLA thresholds.
- **Operational Cost (OC):** Aggregate cost of active resources during evaluation.
- **Adaptation Latency (AL):** Time taken by the system to respond to workload fluctuations.

These metrics jointly quantify the accuracy, responsiveness, and cost-effectiveness of the proposed approach.



4.4 Performance Analysis

4.4.1 Prediction Performance

The **GNN-ST model** significantly outperformed all baseline methods in forecasting future workloads. The model captured both **spatial correlations** among dependent nodes and **temporal patterns** over time. Table 1 summarizes the comparative prediction performance across all models.

Table 1 – Prediction Accuracy Comparison

Model	MAE	RMSE	MAPE (%)
ARIMA	0.183	0.267	10.4
LSTM	0.141	0.201	8.7
CNN + LSTM	0.126	0.182	7.3
GNN (Static)	0.119	0.176	6.9
Proposed GNN-ST	0.092	0.138	5.1

The proposed Spatio-TemporalGNN achieved approximately 25% lower RMSE compared to LSTM and 40% lower MAPE compared to ARIMA.

This improvement results from the model's ability to encode graph connectivity and workload propagation effects, which traditional sequential models cannot represent.

4.4.2 Resource Allocation Efficiency

The GNN-based prediction outputs were fed into the proactive Resource Allocation Engine (RAE). The performance was evaluated on key operational indicators, summarized in Table 2.

Table 2 – Resource Allocation Performance Comparison

Model	Resource Utilization (%)	SLA Violation Rate (%)	Operational Cost (USD/hr)
ARIMA + Allocator	63.4	6.8	128.9
LSTM + Allocator	70.5	4.3	112.7

CNN-LSTM + Allocator	73.2	3.6	105.4
GNN (Static) + Allocator	76.8	3.2	101.9
Proposed GNN-ST + Allocator	88.7	1.4	79.3

The results demonstrate that the proposed framework achieved 88.7% average resource utilization, representing a 17–20% improvement over baseline methods.

Moreover, SLA violations dropped to 1.4%, and operational costs were reduced by 25–35%, validating the framework's ability to make precise and proactive allocation decisions.

4.4.3 Adaptation and Convergence

The system's adaptability was evaluated by subjecting the workload to sudden bursts and diurnal demand patterns. The GNN-ST model adapted within an average adaptation latency of 1.7 minutes, compared to 4–6 minutes for LSTM-based models.

Training convergence was achieved within 60 epochs, after which the model stabilized with minimal oscillations in prediction error and consistent reward improvement for allocation decisions.

4.5 Discussion of Results

The experimental findings confirm that modeling cloud infrastructure as a graph enables a richer understanding of workload dependencies and resource interactions.

The GNN-ST architecture effectively integrates spatial relationships (across services and nodes) and temporal dynamics (across time intervals), leading to superior prediction performance.

Key observations include:

1. **Proactive Scaling:** The framework scaled resources 30–60 seconds earlier



than reactive systems, preventing overloads.

2. **Cost Efficiency:** Intelligent scaling reduced idle resource activation and lowered operating costs.
3. **Reliability:** The reduction in SLA violations demonstrates improved reliability and user satisfaction.
4. **Scalability:** Experiments with up to 500 nodes showed linear scalability with negligible degradation in prediction latency.

Overall, the proposed GNN-based system demonstrates substantial improvements in prediction accuracy, resource efficiency, and operational reliability, making it a promising approach for next-generation cloud resource management systems.

V. CONCLUSION AND FUTURE SCOPE

CONCLUSION

This research introduced an intelligent and proactive Graph Neural Network (GNN)-based workload prediction and resource allocation framework designed for modern cloud computing environments. By modeling the infrastructure as a dynamic graph, where nodes represent computing resources and edges capture inter-node dependencies, the system effectively integrates spatial and temporal learning to forecast future resource demands with high precision.

The proposed Spatio-Temporal GNN (GNN-ST) model successfully captures workload correlations and dependency propagation across cloud components, outperforming conventional deep learning methods such as LSTM and CNN-LSTM in both accuracy and responsiveness. Experimental evaluations demonstrate that the proposed framework achieves up to 40% improvement in prediction accuracy, 25–35% reduction in operational cost, and significant improvement in SLA adherence compared to baseline systems.

Through its proactive resource allocation engine, the framework enables predictive autoscaling, intelligent workload migration, and cost-optimized provisioning. Moreover, the feedback-driven adaptive learning mechanism ensures continuous model evolution in dynamic environments, reducing latency and improving resource utilization.

The results validate that integrating graph-based modeling with AI-driven prediction leads to an autonomous, scalable, and efficient resource management system suitable for large-scale cloud and edge infrastructures.

FUTURE SCOPE

The proposed framework can be extended by integrating reinforcement learning for autonomous decision-making and energy-aware optimization for sustainable cloud management. Future work may also focus on dynamic graph learning and real-world deployment in platforms like Kubernetes or OpenStack to validate scalability and practical performance.

REFERENCES

- [1] R. Buyya, C. Vecchiola, and S. ThamaraiSelvi, *Mastering Cloud Computing*, 2nd ed., McGraw Hill, 2023.
- [2] M. Al-Dulaimy, A. Shawkat, and L. Garcia, “Workload Prediction for Cloud Computing Using Machine Learning Techniques,” *J. Cloud Comput.*, vol. 9, no. 1, pp. 1–14, 2023.
- [3] S. Kayalvili, R. Senthilkumar, S. Yasotha, and R. S. Kamalakannan, “An Optimized Resource Allocation in Cloud Using Prediction-Enabled Reinforcement Learning,” *Scientific Reports*, vol. 15, no. 1, 2025.
- [4] U. K. Lilhore, S. Alex, and M. Khan, “A Multi-Objective Approach to Load Balancing in Cloud Environments,” *Scientific Reports*, vol. 15, 2025.
- [5] T. Kaur, D. Anand, U. Kaur, and S. Verma, “Integrative Resource Management in Multi-Cloud Computing: A DRL-Based Approach,” *EAI/SIS Journal*, vol. 8, no. 2, 2024.



- [6] P. Sharma and T. S. Pawar, "Dynamic Resource Allocation Using Reinforcement Learning in Cloud Computing: A Review," *J. Cloud Comput.*, vol. 10, no. 1, 2022.
- [7] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang, and P. S. Yu, "A Comprehensive Survey on Graph Neural Networks," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 32, no. 1, pp. 4–24, 2021.
- [8] H. X. Nguyen, S. Zhu, and M. Liu, "Graph Neural Networks for Microservice-Based Cloud Applications," *Sensors*, vol. 22, no. 23, art. 9492, 2022.
- [9] D. Konidaris, "Intent-Based Resource Allocation in Edge and Cloud Environments," *Algorithms*, vol. 18, no. 10, 2025.
- [10] A. Agrawal, "Cluster and Reinforcement-Learning-Based Multi-Objective Resource Allocation in Cloud Systems," *J. Cloud Comput.*, vol. 12, 2025.
- [11] J. Li et al., "Forecasting Resource Usage Pattern Changes in Clouds via a Multi-View Contrast Graph-Evolution Learning Algorithm," *Future Gener. Comput.Syst.*, vol. 155, pp. 37–49, 2024.
- [12] M. A. Hady, S. Hu, M. Pratama, J. Cao, and R. Kowalczyk, "Multi-Agent Reinforcement Learning for Resource Allocation Optimization: A Survey," *Comput. Intell.*, vol. 41, no. 2, 2025.
- [13] S. Kumari and D. Mishra, "Adaptive and Fair Resource Allocation in Cloud Data Centers Leveraging A3C Deep Reinforcement Learning," *arXiv preprint arXiv:2506.00929*, 2025.
- [14] G. Zhou, W. Tian, R. Buyya, and L. Song, "Deep Reinforcement Learning-Based Methods for Resource Scheduling in Cloud Computing: A Review," *Artif. Intell.Rev.*, vol. 57, article 124, 2024.
- [15] J. Li and Y. Zhao, "Graph Neural Network-Based Load Forecasting for Dynamic Cloud Resource Management," *IEEE Trans. Cloud Comput.*, vol. 12, no. 3, pp. 1452–1466, 2025.
- [16] D. Roijers, P. Vamplew, S. Whiteson, and R. Dazeley, "A Survey of Multi-Objective Sequential Decision-Making," *J. Artif. Intell.Res.*, vol. 48, pp. 67–113, 2023.
- [17] G. Tesfahun, "Forecasting Cloud Workloads Using ARIMA Models," *Proc. IEEE ICCT*, pp. 420–427, 2023.
- [18] P. Sharma and T. S. Pawar, "Dynamic Resource Allocation Using Deep Learning in Cloud Computing," *J. Cloud Comput.*, vol. 10, no. 1, 2022.
- [19] M. Al-Dulaimy, A. Shawkat, and L. Garcia, "Workload Prediction for Cloud Computing Using Machine Learning Techniques," *J. Cloud Comput.*, vol. 9, no. 1, 2023.
- [20] R. Buyya, C. Vecchiola, and S. T. Selvi, *Mastering Cloud Computing*, 2nd ed., McGraw Hill, 2023.
- [21] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang, and P. Yu, "A Comprehensive Survey on Graph Neural Networks," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 32, no. 1, pp. 4–24, 2021.
- [22] H. X. Nguyen, S. Zhu, and M. Liu, "Graph Neural Networks for Microservice-Based Cloud Applications," *Sensors*, vol. 22, no. 23, art. 9492, 2022.
- [23] M. A. Hady, S. Hu, M. Pratama, J. Cao, and R. Kowalczyk, "Multi-Agent Reinforcement Learning for Resource Allocation Optimization: A Survey," *Comput. Intell.*, vol. 41, no. 2, 2025.
- [24] J. Li et al., "Forecasting Resource Usage Pattern Changes in Clouds via a Multi-View Contrast Graph-Evolution Learning Algorithm," *Future Gener. Comput.Syst.*, vol. 155, pp. 37–49, 2024.
- [25] S. Kumari and D. Mishra, "Adaptive, Efficient and Fair Resource Allocation in Cloud Data Centers Leveraging A3C Deep Reinforcement Learning," *arXiv preprint arXiv:2506.00929*, 2025.
- [26] T. Kaur, D. Anand, and S. Verma, "A DRL-Based Multi-Objective Framework for Resource Allocation in Cloud Systems," *EAI Endorsed Trans. Smart Cities*, vol. 9, no. 1, 2024.



- [27] G. Zhou, W. Tian, R. Buyya, R. Xue, and L. Song, "Deep Reinforcement Learning-Based Methods for Resource Scheduling in Cloud Computing: A Review," *Artif. Intell.Rev.*, vol. 57, article 124, 2024.
- [28] S. Kayalvili, R. Senthilkumar, S. Yasotha, and R. S. Kamalakannan, "An Optimized Resource Allocation in Cloud Using Prediction-Enabled Reinforcement Learning," *Scientific Reports*, vol. 15, no. 1, 2025.
- [29] Y. Zhang, J. Chen, and Q. Wu, "Spatio-Temporal Graph Attention Networks for Multi-Node Demand Forecasting," *IEEE Access*, vol. 12, pp. 21021–21035, 2024.
- [30] J. Li and Y. Zhao, "Graph Neural Network-Based Load Forecasting for Dynamic Cloud Resource Management," *IEEE Trans. Cloud Comput.*, vol. 12, no. 3, pp. 1452–1466, 2025.
- [31] B.V.Srinivasulu "Enhancing Lossless Image Compression through Smart Partitioning, Selective Encoding, and Wavelet Analysis" *SSRG International Journal of Electronics and Communication Engineering*, Scopus indexed(Q3), ISSN: 2348-8549/<https://doi.org/10.14445/23488549/IJECE-V11I5P120>, Volume 11 Issue 5, 207-219, May 2024,.
- [32] B.V.Srinivasulu "Context-Aware Machine Learning Approaches For Enhanced Depression Prediction From Social Media Across Dynamic Conditions", *Scopus(Q3) ISSN: 2073-607X, 2076-0930.*, Volume 13 Issue 03 Year 2021,IJCNIS. 2021.
- [33] Attended a International Conference on "Redefining Agricultural Disease Management Using AI and ML" ISBN: 979-8-3315-1208-8 by IEEE conference – ICICCS- 2025 on 19 th – 21 st March,2025.
- [34] Attended a International Conference on "The role of AI – Enhanced Financial literacy and HR marketing in personal investment decisions a quantitative analysis of urban rural household" by IEEE conference – ICTBIG- 2024 on 13 th – 14 th December,2024.