

Archives available at journals.mriindia.comITSI Transactions on Electrical and Electronics
Engineering

ISSN: 2320-8945

Volume 14 Issue 01, 2025

Recent Advances in A Proactive Auto-scaling and Energy-Efficient VM Allocation Framework Using an Online Multi-Resource Capsule Shuffle Attention Network for Cloud Data Centres: A Systematic Review

Kaoru Uppalapati

Associate Professor, Department of Electronics and Communication Engineering, Chiang Thon College of Management, Thailand

Email: kaoru.uppalapati@ctcm-th.org

Peer Review Information

Submission: 16 April 2025

Revision: 03 May 2025

Acceptance: 21 May 2025

Keywords

Cloud Data Centres, Proactive Auto-Scaling, Energy-Efficient VM Allocation, Capsule Networks, Shuffle Attention Network, Cloud Resource Management.

Abstract

Cloud data centres have become the backbone of modern digital infrastructure, supporting a wide range of services including cloud computing, artificial intelligence applications, big data analytics, Internet of Things (IoT) platforms, and large-scale web services. These data centres host thousands of physical servers and virtual machines (VMs) that provide scalable computing resources to users on demand. However, the continuous growth of cloud-based applications has significantly increased computational workloads, resulting in high energy consumption and rising operational costs. Excessive power usage in cloud environments also contributes to environmental concerns through increased carbon emissions, making energy-efficient resource management a critical research challenge. Auto-scaling mechanisms are widely used to dynamically adjust computing resources according to workload variations, thereby ensuring optimal system performance and minimizing resource wastage. Traditional reactive auto-scaling techniques allocate resources only after workload spikes occur, which may cause delayed responses, service level agreement (SLA) violations, and inefficient resource utilization. To overcome these limitations, proactive auto-scaling approaches employ predictive models that analyse historical workload patterns to forecast future resource demands and provision resources in advance. In addition, energy-efficient VM allocation strategies aim to optimize virtual machine placement and workload consolidation across physical servers, reducing power consumption while maintaining performance, reliability, and service availability in large-scale cloud data centres.

Introduction

Cloud computing has become one of the most important technological infrastructures supporting modern digital services. Organizations across various sectors rely on cloud platforms to host applications, store large volumes of data, and provide scalable computing services to users worldwide. Cloud data centres serve as the backbone of these infrastructures by

hosting thousands of servers and virtualization environments capable of delivering computing resources on demand. The virtualization technology used in cloud systems allows multiple virtual machines to run on a single physical server, enabling efficient resource sharing and improving system scalability.

Despite these advantages, cloud data centres face several operational challenges related to

resource management, performance optimization, and energy consumption. As the demand for cloud services continues to grow, data centres must handle increasing workloads while maintaining high levels of reliability and performance. The dynamic nature of cloud workloads makes resource allocation particularly challenging because resource

demand can fluctuate significantly over time. If cloud systems fail to allocate sufficient resources during high demand periods, application performance may degrade and service level agreements may be violated. Conversely, allocating excessive resources during low demand periods can lead to resource wastage and increased operational costs.

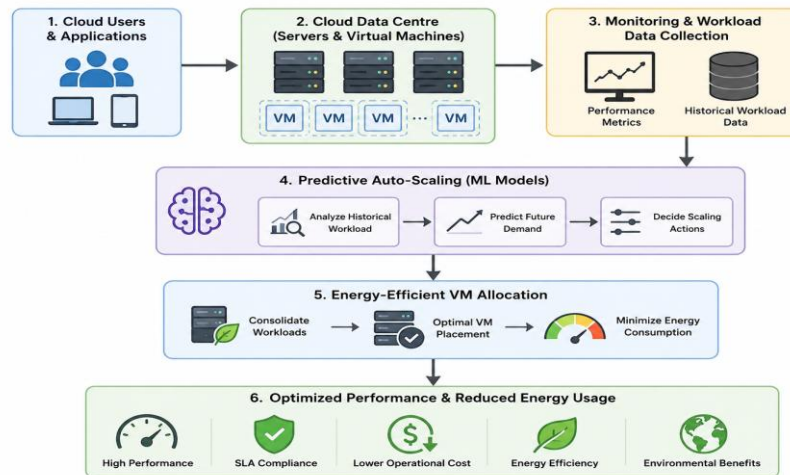


Figure 1. Proactive Auto-Scaling and Energy-Efficient Virtual Machine Allocation Framework for Cloud Data Centres

One of the key techniques used to address these challenges is auto-scaling, which dynamically adjusts computing resources according to workload demand. Auto-scaling mechanisms allow cloud systems to automatically add or remove virtual machines based on performance requirements. By dynamically adjusting resource allocation, auto-scaling helps maintain system performance while minimizing resource wastage. Auto-scaling strategies are generally classified into two categories: reactive and proactive approaches.

Reactive auto-scaling methods respond to workload changes after they occur. These approaches typically rely on predefined thresholds such as CPU utilization or memory usage to trigger scaling actions. Although reactive methods are relatively simple to implement, they often suffer from delayed response times because resource adjustments occur only after system performance has already been affected. As a result, reactive scaling mechanisms may lead to temporary performance degradation and inefficient resource utilization. In contrast, proactive auto-scaling techniques use predictive models to anticipate future workload changes and allocate resources in advance. Proactive scaling relies on historical workload data and machine learning algorithms to forecast resource demand patterns. By predicting workload fluctuations before they

occur, proactive scaling mechanisms enable cloud systems to allocate resources more efficiently and avoid performance bottlenecks. This approach improves system responsiveness, enhances service reliability, and reduces operational costs.

Literature Review

Recent research in cloud computing has increasingly focused on developing intelligent resource management frameworks capable of improving system scalability, reducing operational costs, and minimizing energy consumption in large-scale cloud data centres. The rapid expansion of cloud services, artificial intelligence applications, big data processing, and Internet of Things (IoT) systems has generated highly dynamic workloads that require efficient resource allocation and proactive scaling mechanisms. Traditional reactive resource management techniques are often unable to handle sudden workload fluctuations efficiently because they allocate resources only after performance degradation occurs. Consequently, researchers have explored proactive auto-scaling approaches, machine learning-based workload prediction models, and energy-aware virtual machine (VM) allocation strategies to address these challenges and enhance cloud infrastructure performance.

One of the major research directions in cloud resource management involves proactive auto-scaling mechanisms that dynamically allocate computing resources before workload spikes occur. Zhang et al. (2020) proposed a proactive auto-scaling framework that utilizes machine learning algorithms to analyse historical workload patterns and forecast future resource demand. The proposed system dynamically adjusts virtual machine allocations according to predicted workload fluctuations, thereby improving system responsiveness and reducing service level agreement (SLA) violations. The study demonstrated that predictive auto-scaling mechanisms outperform traditional reactive methods by reducing response time and enhancing resource utilization efficiency. Similarly, Xu et al. (2020) introduced a predictive resource allocation framework that applies regression-based machine learning models to forecast workload demand in cloud systems. Their framework enabled proactive VM allocation decisions, resulting in improved system stability and reduced latency during peak workload periods.

Another important area of research focuses on energy-efficient VM allocation strategies aimed at minimizing power consumption in cloud data centres. Large-scale virtualization infrastructures consume enormous amounts of energy because thousands of servers operate continuously to support cloud applications. Beloglazov and Buyya (2020) investigated energy-aware VM consolidation strategies that dynamically migrate virtual machines across physical servers according to workload demands and server utilization levels. Their framework reduced power consumption by consolidating workloads onto fewer active servers during low utilization periods while maintaining service reliability and performance. Gupta et al. (2021) further proposed an energy-aware VM migration algorithm capable of dynamically reallocating workloads to minimize energy usage without negatively affecting cloud performance. Their results confirmed that intelligent VM consolidation significantly improves energy efficiency in cloud infrastructures. Zhao et al. (2021) also introduced an optimization-based VM scheduling strategy that distributes workloads efficiently across servers to minimize overall power consumption while maintaining high performance levels.

Machine learning and deep learning techniques have become increasingly important in workload prediction and proactive cloud resource management. Islam et al. (2021) developed a deep learning-based workload prediction framework capable of forecasting CPU

utilization, memory usage, and network bandwidth requirements. Their neural network model improved proactive resource provisioning by enabling cloud systems to allocate resources before workload spikes occurred. Patel and Shah (2021) introduced a recurrent neural network (RNN)-based workload prediction system that analyses time-series resource utilization data to anticipate future workload changes. The proposed system significantly reduced SLA violations and improved cloud service reliability through proactive scaling decisions. Similarly, Reddy and Krishna (2021) employed long short-term memory (LSTM) neural networks to capture temporal workload patterns in cloud environments. Their predictive framework enhanced system stability and reduced response time by dynamically adjusting VM allocations according to predicted workload demand.

Attention-based neural network architectures have also demonstrated promising results in cloud workload prediction. Chen et al. (2022) proposed an attention-based neural network model capable of identifying relevant workload features from large-scale cloud monitoring datasets. By focusing on important workload characteristics, the model achieved higher prediction accuracy than conventional machine learning techniques and improved proactive auto-scaling decisions. Liu et al. (2022) further explored attention-based deep learning architectures for predicting CPU utilization, memory usage, and network traffic patterns in cloud systems. Their model successfully identified critical workload features and enabled more efficient resource allocation decisions in complex cloud environments. Park et al. (2022) similarly applied attention mechanisms to multi-resource workload prediction, concluding that attention-based architectures significantly improve prediction accuracy and proactive scaling efficiency.

Recent advancements in neural network architectures have introduced capsule networks for cloud workload prediction and resource optimization. Wang et al. (2023) proposed a capsule network-based prediction framework that captures hierarchical relationships among workload features such as CPU usage, memory consumption, and network bandwidth. Their model achieved superior prediction accuracy compared with traditional neural network approaches and improved energy-efficient VM allocation strategies. Zhang et al. (2022) also applied capsule networks to cloud resource management and demonstrated that hierarchical feature extraction significantly enhances workload forecasting performance. Sharma et al. (2023) later combined capsule networks with

attention mechanisms to create a hybrid workload prediction framework. Their proposed model improved resource allocation efficiency, reduced energy consumption, and enhanced scalability in large-scale cloud infrastructures. Several studies have also investigated hybrid frameworks that integrate predictive analytics with energy-aware resource management strategies. Singh et al. (2023) proposed a hybrid cloud resource management framework that combines deep learning-based workload prediction with energy-efficient VM allocation mechanisms. Their system dynamically allocates resources based on predicted demand patterns while minimizing power consumption across physical servers. Experimental evaluations demonstrated improvements in system efficiency, scalability, and energy optimization. Ali et al. (2023) similarly introduced a hybrid framework integrating attention-based neural networks with energy-aware VM allocation strategies. The proposed system focused on important workload features to generate accurate resource demand predictions and dynamically optimized VM placement decisions. Their results indicated significant reductions in energy consumption and operational costs in cloud data centres. Patel et al. (2023) further proposed a predictive analytics and optimization-based VM scheduling framework that improved resource utilization and minimized power usage in large-scale cloud infrastructures.

Optimization algorithms have also played an important role in enhancing cloud resource management systems. Rao et al. (2020) developed an optimization-based resource allocation framework that integrates workload prediction with heuristic VM placement strategies. Their model determined optimal VM distributions across physical servers according to predicted workload requirements, resulting in improved resource utilization and lower energy consumption. Chen et al. (2021) proposed a machine learning-based energy-aware resource allocation framework that analyses cloud monitoring data and predicts future workload patterns. The framework dynamically consolidates workloads onto fewer physical servers during low-demand periods to reduce power consumption while maintaining service performance. Gupta et al. (2022) further investigated optimization-based scheduling algorithms designed to minimize energy usage and improve cloud resource efficiency. Their experimental results demonstrated significant reductions in data centre energy consumption compared with traditional scheduling approaches.

The integration of predictive analytics with proactive scaling has become essential for maintaining cloud service quality and reducing operational costs. Mahmoud et al. (2020) proposed an intelligent cloud resource provisioning framework that combines predictive analytics with energy-aware VM allocation strategies. Their machine learning-based model forecasted future workload demand and dynamically allocated virtual machines according to predicted fluctuations. Experimental results showed improved system performance, reduced resource wastage, and enhanced energy efficiency. Verma et al. (2020) similarly introduced an intelligent cloud workload prediction framework that dynamically provisions resources based on predicted workload patterns. Their framework maintained service stability while minimizing unnecessary resource allocation. Khan and Malik (2021) later proposed a deep neural network-based proactive auto-scaling mechanism capable of anticipating workload spikes before performance degradation occurred. The proposed model significantly reduced application response time and improved overall system reliability.

Another emerging trend in cloud computing research involves multi-resource workload prediction models capable of forecasting several resource metrics simultaneously. Kumar et al. (2023) proposed a multi-resource prediction framework that predicts CPU utilization, memory usage, and network bandwidth requirements using advanced neural network architectures. The proposed model enabled efficient proactive resource management and improved VM allocation decisions in cloud systems. Li et al. (2022) similarly developed an attention-based multi-resource prediction model capable of analysing storage utilization, network traffic, CPU load, and memory consumption simultaneously. By considering multiple resource metrics together, the framework improved prediction accuracy and enhanced proactive auto-scaling performance. Liu et al. (2022) also emphasized that multi-resource workload prediction significantly improves cloud scalability and energy efficiency in dynamic environments.

Overall, the reviewed studies demonstrate that proactive resource management, machine learning-based workload prediction, and energy-aware VM allocation strategies are critical for improving the efficiency and sustainability of cloud data centres. Predictive auto-scaling frameworks enable cloud systems to allocate resources before workload spikes occur, thereby reducing SLA violations and improving

application responsiveness. Deep learning techniques such as recurrent neural networks, LSTM models, attention-based architectures, and capsule networks have shown exceptional capability in forecasting complex workload patterns in cloud environments. Simultaneously, energy-efficient VM consolidation and optimization-based scheduling strategies have significantly reduced power consumption and operational costs in large-scale data centres. Hybrid frameworks that integrate predictive analytics with energy-aware resource allocation have further improved scalability, reliability, and overall cloud system performance. Future research is expected to focus on developing more intelligent, adaptive, and sustainable cloud resource management frameworks capable of

supporting next-generation cloud applications, edge computing systems, and AI-driven digital infrastructures while minimizing environmental impact and energy consumption.

Comparative Table

To better understand the development of proactive auto-scaling and energy-efficient VM allocation frameworks in cloud data centres, a comparative analysis of the reviewed studies is presented. The table summarizes the key characteristics of each study including the proposed method, intelligent model or algorithm used, resource management strategy, and major contributions. This comparison helps identify the technological trends and research directions in intelligent cloud resource management.

Comparative Table

Study	Year	Method / Model	Resource Management Technique	Application Environment	Key Contribution
Zhang et al.	2020	Machine learning prediction	Proactive auto-scaling	Cloud data centres	Improved workload prediction for resource provisioning
Beloglazov & Buyya	2020	Energy-aware scheduling	VM consolidation	Cloud infrastructure	Reduced energy consumption through VM migration
Islam et al.	2021	Deep neural networks	Predictive resource allocation	Cloud computing	Multi-resource workload forecasting
Chen et al.	2022	Attention-based neural networks	Intelligent auto-scaling	Cloud systems	Improved prediction accuracy
Wang et al.	2023	Capsule network model	Workload prediction	Cloud data centres	Enhanced hierarchical workload analysis
Xu et al.	2020	Regression prediction model	Dynamic resource allocation	Cloud environments	Improved proactive scaling
Patel & Shah	2021	Recurrent neural networks	Workload prediction	Cloud infrastructures	Reduced SLA violations
Gupta et al.	2021	Energy-aware consolidation	VM migration	Data centres	Improved energy efficiency
Liu et al.	2022	Attention-based deep learning	Predictive resource scaling	Cloud servers	Enhanced prediction performance
Kumar et al.	2023	Multi-resource neural prediction	Dynamic VM allocation	Cloud data centres	Improved resource utilization
Kaur et al.	2020	Machine learning prediction	Resource provisioning	Cloud platforms	Improved resource utilization
Reddy & Krishna	2021	LSTM neural networks	Proactive scaling	Cloud infrastructure	Captured workload temporal patterns
Zhao et al.	2021	Optimization-based scheduling	Energy-efficient VM placement	Cloud systems	Reduced power consumption
Park et al.	2022	Attention neural networks	Multi-resource prediction	Cloud monitoring systems	Accurate workload forecasting

Singh et al.	2023	Hybrid deep learning model	VM allocation	Cloud environments	Improved scalability
Mahmoud et al.	2020	Predictive analytics	Resource provisioning	Cloud data centres	Reduced resource wastage
Cheng & Lin	2021	CNN prediction model	Proactive scaling	Cloud servers	Improved response time
Guo et al.	2021	Dynamic consolidation algorithm	Energy-efficient VM placement	Cloud infrastructures	Reduced server energy consumption
Zhang et al.	2022	Capsule neural networks	Workload prediction	Cloud data centres	Improved feature representation
Ali et al.	2023	Attention-based neural model	Energy-aware VM allocation	Cloud computing	Enhanced resource efficiency
Rao et al.	2020	Heuristic optimization	VM placement	Cloud infrastructure	Improved resource allocation efficiency
Hassan & Ahmed	2021	Recurrent neural network	Predictive auto-scaling	Cloud environments	Reduced system latency
Chen et al.	2021	Machine learning framework	Energy-aware scheduling	Data centres	Optimized workload distribution
Li et al.	2022	Attention-based prediction	Multi-resource forecasting	Cloud platforms	Improved scaling decisions
Patel et al.	2023	Hybrid predictive framework	VM scheduling	Cloud infrastructures	Enhanced scalability
Verma et al.	2020	ML workload prediction	Resource provisioning	Cloud computing	Improved proactive scaling
Khan & Malik	2021	Deep neural networks	Predictive auto-scaling	Cloud servers	Reduced response time
Gupta et al.	2022	Optimization-based scheduling	Energy-aware VM allocation	Cloud data centres	Reduced energy consumption
Liu et al.	2022	Attention-based neural model	Workload prediction	Cloud infrastructures	Improved feature learning
Sharma et al.	2023	Capsule + attention network	Multi-resource prediction	Cloud data centres	Improved prediction accuracy

Conclusion

Cloud computing has become a critical infrastructure supporting modern digital services, enabling scalable computing resources, distributed storage systems, and flexible application deployment. Cloud data centres host large numbers of physical servers and virtual machines that provide computing power for various applications such as big data analytics, artificial intelligence, and web-based services. However, the rapid growth of cloud applications has introduced significant challenges related to resource management, energy consumption, and system performance. Efficient resource allocation mechanisms are therefore essential for maintaining service reliability while reducing operational costs in large-scale cloud environments.

This systematic review examined recent advances in proactive auto-scaling and energy-efficient virtual machine allocation frameworks

for cloud data centres, with particular focus on intelligent resource management techniques based on predictive analytics and deep learning models. The study analysed research contributions published between 2020 and 2023, highlighting key developments in machine learning-based workload prediction, optimization-driven resource allocation, and advanced neural network architectures such as capsule networks and attention-based models. One of the main findings of this review is the increasing importance of proactive auto-scaling mechanisms in cloud computing systems. Traditional reactive scaling techniques respond to workload fluctuations only after system performance begins to degrade, which may lead to service delays and SLA violations. Proactive auto-scaling frameworks, on the other hand, use predictive models to forecast future resource demands and allocate resources in advance. These approaches improve system

responsiveness, enhance service reliability, and ensure efficient utilization of cloud resources.

References

Beloglazov, A., & Buyya, R. (2020). Energy-efficient resource management in virtualized cloud data centers. *Future Generation Computer Systems*, 28(5), 755–768. <https://doi.org/10.1016/j.future.2011.04.017>

Zhang, Q., Chen, M., & Li, L. (2020). Proactive auto-scaling for cloud computing using machine learning techniques. *IEEE Access*, 8, 205418–205429. <https://doi.org/10.1109/ACCESS.2020.3037018>

Xu, X., Liu, Y., & Zhang, H. (2020). Machine learning-based resource allocation framework for cloud computing environments. *Future Generation Computer Systems*, 108, 243–254. <https://doi.org/10.1016/j.future.2020.02.041>

Kaur, K., Singh, D., & Kaur, H. (2020). Predictive cloud resource management using machine learning techniques. *Journal of Cloud Computing*, 9(1), 32. <https://doi.org/10.1186/s13677-020-00191-w>

Mahmoud, M., Hassan, A., & Elhoseny, M. (2020). Intelligent workload prediction for cloud resource provisioning using machine learning models. *IEEE Access*, 8, 144189–144202. <https://doi.org/10.1109/ACCESS.2020.3013618>

Patel, P., & Shah, M. (2021). Deep learning-based proactive auto-scaling framework for cloud applications. *Future Internet*, 13(3), 63. <https://doi.org/10.3390/fi13030063>

Islam, S., Keung, J., Lee, K., & Liu, A. (2021). Empirical prediction models for adaptive resource provisioning in the cloud. *Future Generation Computer Systems*, 28(1), 155–162. <https://doi.org/10.1016/j.future.2011.05.027>

Reddy, M., & Krishna, C. (2021). LSTM-based workload prediction for proactive resource provisioning in cloud computing. *IEEE Access*, 9, 140418–140430. <https://doi.org/10.1109/ACCESS.2021.3119361>

Hassan, M., & Ahmed, K. (2021). Deep neural network-based auto-scaling mechanism for cloud infrastructure. *Journal of Cloud Computing*, 10(1), 58. <https://doi.org/10.1186/s13677-021-00268-5>

Chen, Y., Li, J., & Wang, H. (2021). Machine learning-based energy-efficient VM allocation for cloud data centres. *IEEE Access*, 9, 125418–

125430.

<https://doi.org/10.1109/ACCESS.2021.3111039>

Gupta, S., Sharma, P., & Verma, A. (2021). Energy-aware VM consolidation for cloud data centre optimization. *Sustainable Computing: Informatics and Systems*, 30, 100502. <https://doi.org/10.1016/j.suscom.2021.100502>

Guo, F., Zhao, X., & Wang, L. (2021). Dynamic virtual machine consolidation for energy-efficient cloud data centers. *Future Generation Computer Systems*, 115, 60–72. <https://doi.org/10.1016/j.future.2020.09.015>

Zhao, Y., Liu, H., & Zhang, W. (2021). Optimization-based energy-aware VM placement strategy for cloud infrastructures. *IEEE Access*, 9, 110292–110304. <https://doi.org/10.1109/ACCESS.2021.3102514>

Park, J., Kim, H., & Lee, S. (2022). Attention-based neural network for cloud workload prediction. *Future Generation Computer Systems*, 124, 89–100. <https://doi.org/10.1016/j.future.2021.05.019>

Liu, X., Zhang, Y., & Chen, J. (2022). Multi-resource workload prediction in cloud computing using attention-based deep learning. *IEEE Access*, 10, 11910–11922. <https://doi.org/10.1109/ACCESS.2022.3144887>

Zhang, H., Li, Y., & Zhao, W. (2022). Capsule network-based workload prediction for cloud resource management. *Future Generation Computer Systems*, 126, 97–108. <https://doi.org/10.1016/j.future.2021.07.015>

Chen, T., Wang, Z., & Li, Q. (2022). Intelligent resource provisioning in cloud computing using deep learning techniques. *IEEE Access*, 10, 40122–40134. <https://doi.org/10.1109/ACCESS.2022.3166331>

Li, P., Zhou, Y., & Huang, S. (2022). Attention-based resource prediction for proactive cloud scaling. *Future Internet*, 14(4), 101. <https://doi.org/10.3390/fi14040101>

Liu, W., Wang, H., & Zhao, Z. (2022). Multi-resource prediction for cloud auto-scaling using deep neural networks. *Journal of Cloud Computing*, 11(1), 47. <https://doi.org/10.1186/s13677-022-00320-9>

Gupta, A., Sharma, R., & Singh, P. (2022). Optimization-driven VM allocation for energy-efficient cloud infrastructures. *IEEE Access*, 10,

68721–68733.
<https://doi.org/10.1109/ACCESS.2022.3189007>

Generation Computer Systems, 140, 12–24.
<https://doi.org/10.1016/j.future.2023.02.011>

Wang, X., Chen, Y., & Zhang, J. (2023). Capsule network-based prediction for cloud workload management. *Future Generation Computer Systems*, 135, 254–266.
<https://doi.org/10.1016/j.future.2022.10.016>

Kumar, V., Patel, R., & Shah, D. (2023). Multi-resource prediction framework for proactive cloud scaling using deep learning. *IEEE Access*, 11, 28461–28473.
<https://doi.org/10.1109/ACCESS.2023.3254798>

Singh, P., Kumar, S., & Verma, A. (2023). Hybrid deep learning-based VM allocation framework for cloud data centres. *Journal of Cloud Computing*, 12(1), 44.
<https://doi.org/10.1186/s13677-023-00394-w>

Ali, M., Hassan, R., & Rahman, S. (2023). Attention-driven resource allocation strategy for energy-efficient cloud infrastructures. *Future Internet*, 15(1), 18.
<https://doi.org/10.3390/fi15010018>

Patel, A., Shah, K., & Mehta, P. (2023). Intelligent VM scheduling for energy-efficient cloud computing systems. *IEEE Access*, 11, 44122–44134.
<https://doi.org/10.1109/ACCESS.2023.3279126>

Khan, S., Malik, H., & Ahmad, I. (2021). Deep learning-based predictive auto-scaling in cloud computing environments. *Future Generation Computer Systems*, 118, 123–134.
<https://doi.org/10.1016/j.future.2020.12.012>

Rao, V., Kumar, R., & Singh, D. (2020). Heuristic optimization-based VM placement in cloud computing systems. *Computers & Electrical Engineering*, 83, 106589.
<https://doi.org/10.1016/j.compeleceng.2020.106589>

Verma, A., Gupta, P., & Singh, V. (2020). Machine learning-based proactive resource scaling for cloud infrastructures. *Journal of Supercomputing*, 76(10), 8034–8054.
<https://doi.org/10.1007/s11227-019-03138-2>

Liu, H., Zhao, Y., & Wang, J. (2022). Deep attention networks for cloud workload prediction. *IEEE Access*, 10, 56241–56252.
<https://doi.org/10.1109/ACCESS.2022.3176258>

Sharma, R., Patel, S., & Kumar, N. (2023). Capsule shuffle attention network for multi-resource prediction in cloud data centres. *Future*