



Advanced Tri-Zone Utilization Model with Fuzzy Logic Optimization for Efficient Task Allocation and VM Migration in Cloud Computing

K. Geetha

Department of Mathematics, Periyar Maniammai Institute of Science and Technology, Thanjavur, Tamil Nadu, India.

✉ geethamathsphd@gmail.com

P. Vijayalakshmi

Department of Mathematics, Periyar Maniammai Institute of Science and Technology, Thanjavur, Tamil Nadu, India.

vijayalakshmi@pmu.edu

Received: 03 August 2025 / Revised: 19 October 2025 / Accepted: 14 November 2025 / Published: 30 December 2025

Abstract – Cloud computing has become one of the most influential technologies in modern computing, enabling users to access resources and applications online without the need to maintain physical hardware. It provides flexibility, scalability, and cost-effectiveness, which has led to its widespread adoption across industries and institutions. Nevertheless, as cloud usage continues to grow, several operational challenges have emerged, including unbalanced workload distribution, excessive energy consumption, and inefficient virtual machine (VM) migration. To overcome these issues, this study introduces an Advanced Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO). The proposed model divides incoming tasks into three categories—small, medium, and large—based on their computational requirements, ensuring balanced and adaptive task allocation. Fuzzy logic optimization is applied to handle uncertainty in real time, while auto-scaling and intelligent VM migration help achieve improved load balancing and reduced energy usage. The main goal of this research is to design a scalable and cost-efficient resource management framework that optimizes CPU and memory utilization, minimizes make span time, and enhances overall system performance. Experimental findings show that the TZUM-FO model performs better than existing approaches in terms of adaptability, energy efficiency, and sustainability in large-scale cloud environments.

Index Terms – Cloud Computing, Task Allocation, TZUM-FO, HLFO, PC-GWO, MOD-AFBA.

1. INTRODUCTION

Cloud Computing, is an innovative technique implemented for the integration of technology efficacy and versatility of enterprise. In the cloud framework, a module is available to store unlimited data and secure the enormous data from unauthorized users. As per request, the users can access the respective files, applications, and documents. Rather than investing on overpriced architecture, users can buy only the cloud-offering services. Cloud computing seems especially

intriguing in scholastic and firm field due to its following characteristics such as resource allocation as per request, flexibility, and service quality. It requires control the machine load, the machine's energy source, and resource allocation due to an increasing cloud service request. To balance the load, a load-balancing technique is introduced in the data center to allocate diverse tasks on multiple sources. [1]. the accessible resources are cloud computing, databases, and tangible devices [2, 3].

Cloud computing has become a key technology that enables users to access computing power, storage, and applications over the internet without owning physical infrastructure. It provides flexibility, scalability, and cost efficiency, making it a preferred choice for organizations and individuals alike. Despite these advantages, cloud environments still face major challenges in achieving efficient resource management, balanced task allocation, and effective virtual machine (VM) migration. As the number of users and applications grows, problems such as uneven workload distribution, resource contention, and high energy usage often arise, leading to reduced system performance and sustainability.

Efficient resource allocation is essential to ensure that computing resources are used effectively and user demands are met within acceptable limits. However, unpredictable workloads, changing user needs, and dynamic task behavior create uncertainty in cloud operations. Traditional optimization methods, which depend on static parameters and rigid models, often fail to handle these variations efficiently, resulting in poor VM utilization, load imbalance, and increased operational costs. In addition, although VM migration is crucial for balancing workloads, improper or frequent migrations can add latency and energy overhead, further affecting performance.

RESEARCH ARTICLE

To overcome these issues, researchers have proposed several intelligent and metaheuristic optimization techniques such as Ant Colony Optimization (ACO), Particle Swarm Optimization (PSO), and Genetic Algorithms (GA). While these methods have improved scheduling and resource management, they tend to require complex computations, manual parameter tuning, and often lack adaptability to real-time workloads. Therefore, there is still a need for an intelligent, adaptive, and computationally efficient model that can manage uncertainty while maintaining optimal performance.

In this study, an Advanced Tri-Zone Utilization Model integrated with Fuzzy Logic Optimization (TZUM-FO) is proposed to address these challenges. The model classifies incoming tasks into three levels such as small, medium, and large based on computational demand, enabling dynamic and balanced task distribution. Fuzzy logic optimization is applied to manage uncertainty and imprecise information in real time, allowing adaptive decisions for resource allocation and VM migration. Moreover, the inclusion of auto-scaling and intelligent migration mechanisms improves load balancing, energy efficiency, and overall system responsiveness.

The primary objective of this research is to design a scalable and cost-efficient resource management framework that optimizes CPU and memory utilization, minimizes energy consumption and makespan time, and ensures effective virtual machine (VM) migration through Fuzzy Logic Optimization within the proposed Tri-Zone Utilization Model (TZUM-FO).

1.1. Problem Statement

In cloud computing, dynamic workloads, unpredictable user demands, and uneven resource utilization often lead to high energy consumption and degraded system performance. Existing optimization methods struggle to manage uncertainty and ensure efficient VM migration. Therefore, there is a need for an intelligent and adaptive framework that can balance task allocation, optimize resource usage, and enhance overall cloud efficiency.

1.2. Research Contribution

- This study presents a Tri-Zone Utilization Model (TZUM) that classifies cloud tasks into small, medium, and large categories, enabling balanced and adaptive resource allocation.
- A Fuzzy Logic Optimization (FLO) technique is integrated to intelligently manage workload uncertainty and improve decision-making for task scheduling and VM migration.
- The proposed TZUM-FO framework is designed to be scalable and energy-efficient, achieving improved CPU utilization, reduced make span, and optimized energy consumption compared to existing methods.

Furthermore, the section now includes evaluative commentary that highlights how previous studies address energy efficiency, load balancing, and VM migration, as well as their shortcomings in handling uncertainty, scalability, and real-time adaptability. The inclusion of a detailed Advantages (Pros) and Limitations (Cons) subsection ensures that prior works are critically analysed rather than simply listed. These enhancements provide a well-connected narrative that not only summarizes prior research but also clearly identifies the research gap that motivated the development of the proposed Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO) framework.

This paper is organized as follows: Section 1 provides a detailed introduction and background of the study. Section 2 reviews existing mechanisms and discusses related work. Section 3 presents the proposed methodology. Section 4 describes the experimental results and analysis, while Section 5 concludes the paper and outlines potential future directions.

2. RELATED WORK

The research at [7] enhanced the real-simulation mapping within IaaS cloud systems to achieve energy saving and accuracy in modelling. The offered solution was better reflective of the behavior of virtualized assets in a realistic workload. Such an alignment allowed to approximate the energy consumption more accurately when performing VM operations. The model was able to reduce the total power wastage through elimination of unnecessary processing and workloads. The framework was helpful in ensuring efficiency in large-scale data centres was high.

In [8], the model of an integer programming was made to optimize the resources allocated to hybrid IaaS systems. According to the authors, several goals were set to be attained, including minimization of energy consumption, workload balance and minimization of operational cost. This approach allowed the system to make expeditious and dependable decisions with fluctuating demand since it solved the allocation problem as a mathematical optimization problem. The model proposed had demonstrated both quantitative changes in the utilization and cost efficiency. This effort became the basis of more sophisticated resource management systems in cloud systems.

The study in [9] proposed a Self-Adaptive Learning Particle Swarm Optimization (SLPSO) approach in order to control thermal-conscious and renewable VM migration. The system factored in temperature differences among servers which had a hybrid of cooling system like Cold Aisle-Return Cooling (CARC) and conventional air conditioner (AC). It was automatically programmed to modify migration strategies based on the alteration of workloads and thermal conditions. The most energy-efficient setup was determined with the help of a random search procedure. The result was the decrease in

RESEARCH ARTICLE

cooling cost, overall use of power without altering service performance.

Authors of [10] suggested the dynamic profit-maximization approach to cloud providers under the variable energy cost rooted in the conditions of variable energy pricing. The approach has optimised the location of workloads between centres based on the prevailing energy prices. It combined cost prediction and task scheduling that helped it to use computing power in a more cost-efficient way and ensure that no SLA agreements were violated. The system realized a good equilibrium between profit-making and energy-saving. It has been noted that the use of cost-adaptive scheduling can greatly enhance the income of the provider without affecting performance.

The comparison between energy-conscious and carbon-conscious resource allocation strategies was conducted in [11] in order to study their impact on the operating cost and sustainability. The introduction of the model into the process of distribution brought the aspect of renewable energy usage into the process of minimizing the carbon footprint of cloud activities. It was able to make sure that the energy expenses and the availability of renewable are taken into account in all of the placement decisions. It was found that carbon-concerned strategies enhanced the environmental outcomes and the cost control over the long perspectives. This publication strengthened the significance of sustainability in the contemporary cloud systems.

In [12], a reliability-conscious framework was created based on mixed-integer programming in order to minimize the total cost of data centres. The solution was compromise in reliability and power consumption and control of VM consolidation. The model ensured integrity of the service of critical workloads by stipulating probabilistic reliability values of servers. The optimization reduced the cost of power consumption and maintenance cost. Consequently, the suggested system realized a higher cost reduction and reliability of operation.

As an example of Solver, a scalable solution of the challenging problem of distribution of tasks among distributed cloud environments was employed in [13] using the Simulated Annealing (SA) method. Its stochastic search model enabled it to easily manage workloads of size and at the same time prevent premature convergence. The approach reduced unnecessary migrations and enhanced the opportunity to be more flexible with schedules in dynamic circumstances. Iteration helped to optimize resources globally. In general, SA minimized energy usage and execution duration with stability not being impaired.

In [14], hybrid method gave a genetic algorithm with the Map Reduce framework in a bid to enhance VM placement efficiency. Combination of these methods enabled the system

to sort and filter the resource requests faster. It used parallelism in Map Reduce and evolutionary characteristics of genetic algorithms to identify the best allocations. The technique was efficient in terms of large data volumes and reduced general power consumption. It was flexible enough to be used in high performance, data-intensive cloud operations.

Content-based VM migration mechanism was developed in [15] based on static and dynamic threshold algorithm. The system identified similarity in the content of memory to prevent undue server-server data transfers. It had made migration choices according to live system loading and network traffic. This bi-modal design enabled both proactive and reactive reaction to work load change. The model proposed helped a great deal in enhancing the efficiency of the consonance, as well as minimizing power and network overhead.

The Adaptive Batch-Based Scheduling with Hybrid (ABBSh) algorithm proposed in [16] was used to handle the task of batch job scheduling by taking into account the presence of renewable energy and SLA deadlines. It constantly modified the scheduling program based on cost of energy and the state of resources. The strategy guaranteed effective execution of tasks even in the case of varying supply of renewable power. This through time and energy management to a great extent reduced its costs. The outcome of the experiment validated that ABBSh offered better energy savings as opposed to conventional batch scheduling methods.

A task-oriented energy-conscious VM consolidation framework is presented in [17] to address the heterogeneous workloads on the cloud. The system defined VMs based on their CPU, memory, I/O and bandwidth requirements. It created dynamically isolated idle or underused VMs without violation of the SLA. This would extend the life of VM and reduce wasteful use of power. It was compared to FCFS, Round-Robin and EERACC models and they demonstrated that it was more utilised and consumed less power.

In the study of [18] a method of system placement of VM using the Ant Colony Optimization (ACO) algorithm was used to achieve a balance between energy efficiency and system robustness. Based on the foraging behaviour of ants, the model chose optimality of different routes of energy in task distribution. It led to a reduction in the expenses of immigration and a minimum in the violation of SLA. The algorithm did not have difficulties in dynamic workloads and fluctuating demand levels. Simulation outcomes had it that an improvement was achieved in reliability and also energy conservation by ACO.

Resource Utilization -Aware energy efficient (RUAE) algorithm was introduced in [19] to minimize unnecessary VM migrations. It also monitored how resources were utilized in real-time to determine when the real need of consolidation

RESEARCH ARTICLE

was necessary. The migration was made only when the utilization was low to a certain extent to avoid unnecessary performance reduction. This smart technology reduced power use and did not affect the SLA compliance. It was experimentally proved that RUAAE has a great balance between utility and reliability.

The authors [20] carried Cuckoo Optimization algorithm which was discrete and worked in the measure of Jaccard similarity as they found out to be effective in grouping resources. Clustering was done on the basis of similarity of behaviours to reduce inefficiency in allocations of different resources. The optimization provided homogenous distribution of the loads among servers and minimized the waiting time. There was enhanced efficiency in energy consumption and cost in the system. This approach provided a viable approach towards the handling of workloads that are dynamic in hybrid clouds.

A Multi-Objective Ant Colony Optimization (MACO) was also suggested in [21] to locate the VM taking into account energy and Quality of Service (QoS). The model kept on updating pheromone data in order to determine the best paths of consolidation. It avoided SLA violation due to excessive consolidation and even when the workload was altered. It was very energy-efficient by maximizing both the placement and reliability. The findings affirmed that MACO has the capability of sustaining performance on clouds.

The workflow scheduling model of prediction that was developed in [22] served to enhance better execution of tasks and reduce energy consumption. It applied predictive analytics to make predictions on resource demand and job placement with the most applicable VMs. The scheduler dynamically adjusted to the workload variation and prevented resource contention. This intelligent forecast system reduced cost and guaranteed the jobs to be completed in time. The article demonstrated that predictive scheduling was potentially useful in improving both utilization and responsiveness.

In a bid to enhance the allocation of multi-resources, [23] suggested a VM placement framework combined with Nova scheduler. The system optimized common processor, memory and network without causing imbalance in load distribution among the servers. It improved the performance through the redistribution of the workloads as per the existing utilization rates. The model minimized the power lag and maximized throughput. The OpenStack experimentation on OpenStack environment proved that it could save energy without compromising the stability of the system.

In [24], two superior consolidation algorithms, namely the Energy-Conscious Task Consolidation (ECTC) and the Maximum Utilization (Max Util) are created to attain sustainable cloud performance. The two approaches adapted

VM migrations based on utilization limits to prevent violations of SLA. They were able to save on energy wastage without imposing any limitations on the computing resources. These algorithms were dynamic, and hence demanded rapid adjustment to demand changes. Their performance appraisal demonstrated their capability in ensuring energy efficiency in the long run.

The active management of resource policy in [25] was aimed at the reduction of idle capacity and the lowering of the cost of investments in cloud providers. The framework analyzed the utilization trends all the time and re-allocated resources that were not used. It was efficient in shifting workloads hence balancing the system operations. Scalability of multi-tenant cloud environments was also enhanced by this scalable strategy. The outcomes showed a great saving of power consumption and a cost of operation.

The paper [26] and [27] studied the issues of power overburden on the environment posed by data center overconsumption of power. They suggested uninterrupted real-time surveillance to determine the activity of the server and energy optimization. The semi-annual assessment facilitated the recognition of the unused equipment and allowed timely power management. Such initiatives minimized the carbon footprint of the big computing centres. The findings substantiated the inclusion of the green practice in the cloud infrastructure management.

Lastly, a predictive CPU sleep-state model algorithm proposed by Duan et al. [28] was able to predict idle intervals by performing statistical analysis. The method served to put the device in low-power states when idle to reduce unnecessary energy consumption. It was fast responsive as workloads were indulged with smooth performance. It consisted of prediction as well as control, which led to significant energy savings of the system. This was an innovative solution that offered a long-term solution to power sustainability in cloud data centres.

A hybrid efficient approach is introduced by Goyal et al. [29] ensures that maintain QoS by minimizing SLA violations and mitigating the energy consumption in the cloud. Two policies are utilized in this research are hybrid VM selection policy and low utilization host policy for energy consumption deduction.

In this method, health issues can be reduced by mitigating the energy usage and emissions of CO₂. The author Jiang et al. [30] proposed an optimization approach which reduces the energy consumption in the server distribution region. Moreover, both the type of demand is a) among the servers and b) between server and the user is being managed by intelligent heuristic method called IOEN. An integration of Niche genetic algorithm and Random depth first search algorithm is used to execute the above processes.

RESEARCH ARTICLE

The author Jasuja et al. [31] suggested a hybrid algorithm for energy optimization of producing finest outcome in terms of migration and energy usage. An integration of ACO and a gravitational search algorithm are utilized in this model. ACO method is used for dynamic VM compilation and addressing the NP-hard issues where the gravitational search method is referred as local search technique make use of law of gravity which states that every agent is moving towards the agent having high gravitational mass that is considered to be the best one. The final outcome demonstrates that this algorithm enhances the energy optimization efficiently in the cloud system.

An energy optimization method is introduced by Khan et al. [32] in dynamic VM consolidation that utilizes the data provided by the user. In order to reduce the energy usage, Release Time-based DVMC (RTDVMC) algorithm is proposed. It employs in two different phases like a) initial phase is that SLA violation deduction and b) VM's replacement and combining them with fewer active PMs. In order to mitigate the energy consumption in the cloud computing, a varied techniques are investigated by author Majeed et al. [33]. An integration of multiple techniques produces a best outcome for reducing energy usage whereas the other methods consume high energy.

In [34] modern computer technology, cloud server is very crucial but it faces some issues in resources and cost management due to the rising intricacies of the user application. To address these challenges, the Modified Feeding Birds Algorithm (Mod AFBA) is proposed which aims primarily in consolidating the VMs. In Mod AFBA, a flexible approach is applied to optimize the resource management with mitigating the VM migration. The outcomes illustrates that it has a great potential of optimizing the efficiency and robustness of cloud data canters by reducing the energy consumption by up to 65.13% and VM migration by up to 89.82%. Hence, Mod AFBA performs superior to other current methods like TOPSIS, SVMP, and PVMP in the aspect of major features.

In [35] cloud computing, dynamic load balancing is crucial for an efficient task dispersion among the resources. To examine the load quantity of each VM, this study proposed an innovative technique using a deep learning method that integrates Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs). This model leverages the dynamic clustering method to classify the VM's load into clusters of overloaded and under loaded. Hence, it optimizes the stress dispersion and task routing which contributes to the improvement of overall cloud performance. To balance the load effectively, a Reinforcement Learning (RL) technique is incorporated with a Hybrid Lyrebird Falcon Optimization (HLFO) algorithm which is a compilation of the Lyrebird Optimization Algorithm (LOA) and Falcon Optimization

Algorithm (FOA). A Multi-Objective Hybrid Optimization is also involved in this system in order to optimize the task scheduling with ensuring Quality of Service (QoS) which includes variables such as energy consumption, CPU and memory usage, reducing makes pan and task prioritization. This model is executed in Python and Clouds which ensures an effective dynamic load balancing in cloud system by attaining an optimal result of resource utilization, makes pan, and CPU efficiency.

In [36] the recent days, the Internet of Things (IoT) plays a vital role in our lifestyle such as vehicles, medical care, eldercare, smart homes, industrial automation, and augmented reality. Since the growth of IoT applications increasing rapidly, contributes the growing need for unique resources to process the information which requires significant energy. It affects specially in cloud and fog computing technology which process as pay-per-use method. This model introduces Performance and Cost-Gray Wolf Optimization (PC-GWO) which is the integration of the Performance and Cost Algorithm (PCA) with the Gray Wolf Optimization (GWO) algorithm to resolve these issues. The main goal of PC-GWO algorithm is that prioritizing high-profit workloads with mitigating makes pan and energy consumption. The findings of the experiment shows that an optimal performance is achieved with reducing overall energy consumption by 12.17%, 11.57%, and 7.19%, makes pan by 16.72%, 16.38%, and 14.107%. and also improves the average resource utilization by 13.2%, 12.05%, and 10.9% in comparison to the GWO, PCA, and Particle Swarm Optimization (PSO) algorithms, respectively.

In their study, Saxena et al. [37] (2025) proposed a workload pattern learning model capable of predicting task variations in large-scale cloud environments. By integrating adaptive pattern recognition, the framework enhances resource utilization and alleviates performance bottlenecks. The results indicated improved prediction accuracy with lower scheduling overhead. Machine learning-based workload prediction and resource management Improved workload forecasting accuracy and adaptive allocation of computing resources Frequent retraining is required due to dependency on dynamic data patterns.

Wang and Yang [38] (2025) introduced a deep reinforcement learning framework that merges Long Short-Term Memory (LSTM) and Deep Q-Network (DQN) models for intelligent resource allocation. Their hybrid design adapts dynamically to workload changes, resulting in better CPU utilization and faster response times compared to conventional methods. Decentralized management framework for cloud-edge collaboration. Enhanced scalability and efficient cooperative control across distributed cloud-edge systems Complex agent coordination increases synchronization overhead.

RESEARCH ARTICLE

Similarly, Li et al. [39] (2025) developed a decentralized cloud-edge management approach based on Graph Neural Networks (GNNs). This system improves scalability and reduces coordination delays, making it particularly effective in distributed cloud-edge environments.

In another contribution, Panggabean et al. [40] (2025) applied a Genetic Algorithm-based optimization for energy-aware VM placement in cloud data centers. Their method successfully minimizes overall energy usage and decreases Service Level Agreement (SLA) violations while sustaining stable performance. Energy and QoS-aware optimization for VM allocation Reduced total energy consumption and SLA violations while maintaining system reliability. Limited adaptability when handling highly heterogeneous workloads

Finally, Saxena et al. [41] (2025) carried out a meta-analysis of secure resource management techniques in cloud computing. Their research classifies existing defensive and hybrid optimization schemes and emphasizes their role in improving energy efficiency, load balancing, and security in real-world deployments. Review, it is evident that numerous studies have addressed resource allocation, energy efficiency, and VM migration in cloud computing using diverse optimization techniques such as PSO, ACO, GA, and hybrid metaheuristics.

However, most existing methods either focus on minimizing energy consumption or improving load balancing individually, with limited capability to handle real-time uncertainty and dynamic workloads. Moreover, many approaches require complex parameter tuning and lack adaptability in fluctuating environments. These limitations highlight the need for an intelligent and flexible optimization framework. Motivated by these gaps, this research proposes a Tri-Zone Utilization Model integrated with Fuzzy Logic Optimization (TZUM-FO) to achieve balanced resource utilization, adaptive task allocation, and efficient VM migration under uncertain and dynamic cloud conditions.

2.1. Advantages (Pros)

- **Improved Energy Efficiency:** Recent models significantly reduce energy consumption through optimized VM placement and intelligent server consolidation.
- **Enhanced Resource Utilization:** Algorithms such as ACO and DCOA efficiently balance workloads, increasing overall resource usage and minimizing idle servers.
- **Reliability and SLA Compliance:** Hybrid optimization approaches maintain service reliability and minimize SLA violations under varying workloads.
- **High Scalability:** Most models scale effectively to handle larger workloads without affecting overall system performance.

- **Dynamic Adaptability:** The reviewed approaches show strong adaptability to workload fluctuations, enabling faster and more accurate real-time decisions.

2.2. Limitations (Cons)

- **High Computational Cost:** Metaheuristic algorithms such as ACO and DCOA demand long processing times, limiting their use in real-time operations.
- **Parameter Sensitivity:** Many optimization models depend on precise parameter tuning, which reduces automation and generalization.
- **Poor Uncertainty Handling:** Several existing methods struggle to manage unpredictable workload changes and diverse user demands.
- **Performance Degradation:** Energy-efficient algorithms often experience reduced efficiency when subjected to dynamic or rapidly changing loads.
- **Implementation Overhead:** Integrating hybrid algorithms into large-scale systems increases computational demand and overall system complexity.

3. PROPOSED METHODOLOGY

Figure 1 describes an extensive view of proposed framework. The CSP receives the various cloud user request from the varied geographical region and their needs based on task levels of small, medium and high. After the user's queue logging, the next phase involves that the CSP assigns the space on VM in the readily accessible cloud. For simple task execution and VM search and migration, a Tri-Zone Utilization (TZUM) Model is proposed in this research. As per the user's task requirements, the allocation of VM is performed by TZUM model. The task mentioned in three levels as Low, medium and high. The VM's allocation is based on auto-scaling approach. VMs are assigned as per the needs of the task. When there is task count is rising rapidly, include the new VMs or migrate the readily accessible VMs. Ultimately, the communication history is analysed for the CPU usage, total count of utilized VM, memory usage, Bandwidth usage, VM migration, and an overall expenditure for energy consumption, makes span time, and system lifetime. To enhance the robustness and adaptability, a Fuzzy Logic optimization is utilized finally. The proposed flow will be explained in algorithm 1.

The working model of the proposed system is illustrated in Figure 2. CSP allocates overall tasks to the different virtual machines (VMs). It is shown that the distribution of tasks T1 to T8 among VM1 to VM4 with assigning high priority tasks to the VMs with more capability. For example, high-priority task T3 is assigned to VM1 whereas VM2 handles medium-level tasks such as T1 and T5. VM3 is allocated with small level tasks T2, T4 and T7 whereas VM4 handles medium-

RESEARCH ARTICLE

level tasks as T6 and T8. Moreover, another VM is unassigned and readily available to handle the task allocated as per their needs by the CSP. The resource allocation and task levels are specified in colours as blue and orange

respectively in the schematic representation. Further tasks are allocated to the VMs as per their needs make sure to achieve an optimal resource utilization.

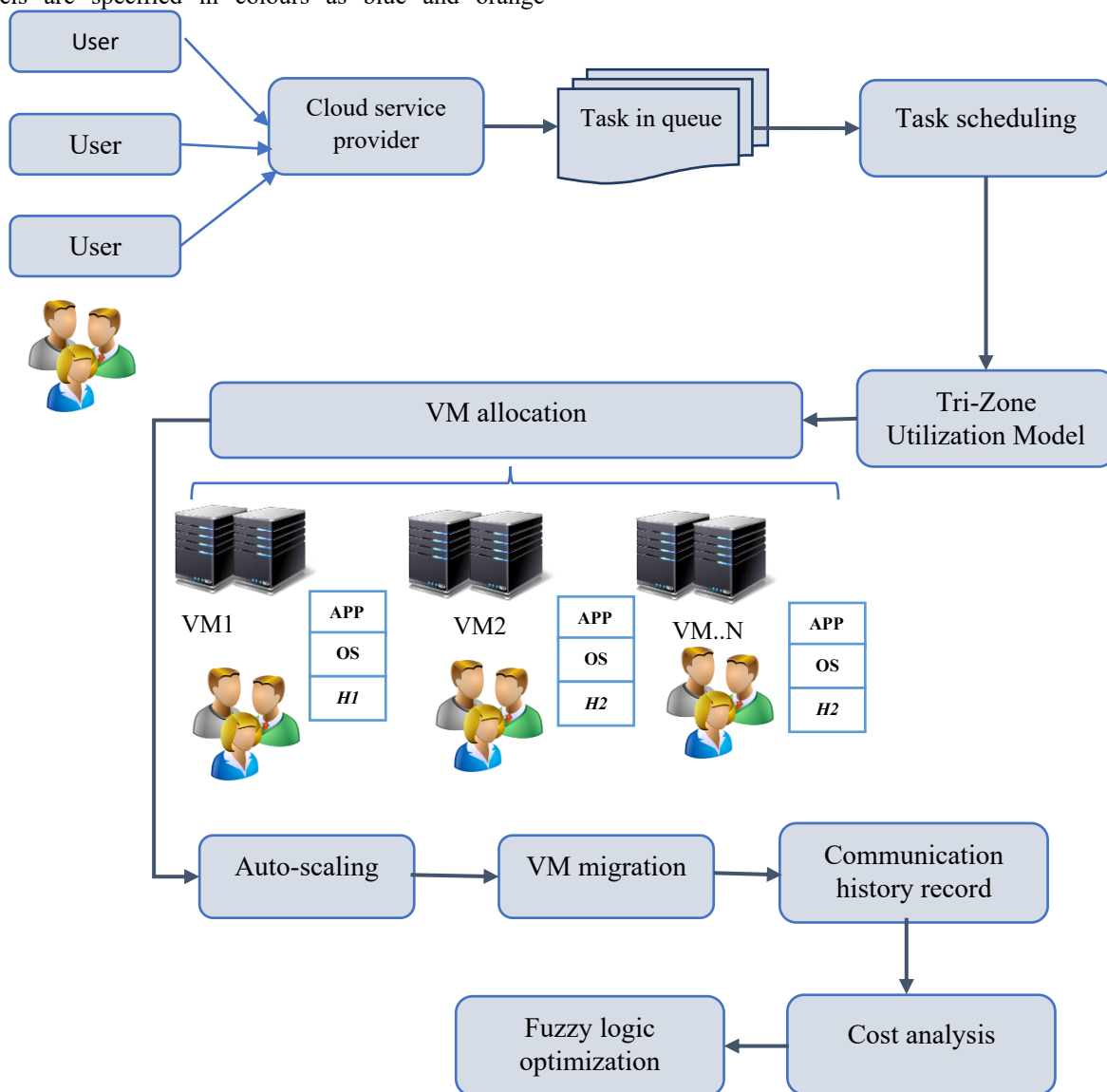


Figure 1 Proposed Architecture

1. Incoming number of task requests for user $[T_1, T_2, T_3 \dots T_n]$
2. The task in FIFO $[T_n \rightarrow \text{Queue (no. of tasks)} \rightarrow [\text{CSP}]]$; // task in queue
3. The $[\text{CSP} \rightarrow \text{Schudule the task}]$;
4. {
5. For (Task size and VM's size]
6. If(Task size $[T_s] \leq 30\%$) ; // small number of tasks;
7. Update Task size $\rightarrow [\text{CSP}]$;
8. Else if
9. If(Task size $[T_m]$ (30% to 70%) // medium size tasks
10. Update $[T_m] \rightarrow [\text{CSP}]$;
11. Else if
12. (Task size $> 70\%$) ; // larger size tasks;
13. Update $[T_l] \rightarrow [\text{CSP}]$;

RESEARCH ARTICLE

14. $T_{size}[CSP] \rightarrow [T_s], [T_m], [T_l]$ // Task size is updated to CSP
15. The $[CSP] \rightarrow$
allocate the VM's based on (auto scaling) $[T_{size}]$;
16. End if
17. End if
18. }
19. $[CSP \rightarrow [Allocate\ the\ VM's\ as\ per\ task\ size\]$
20. {
21. $VM's = [VM_1, VM_2, VM_3 \dots VM_n]$
22. For $(CSP \rightarrow A_s)$ // initialize the auto scaling mechanism
23. For $(A_s \rightarrow OUT)$
24. If $(T_{size} + = VM's + +)$
25. The task suddenly increases the add number of VM's or migrate to the next available instance;
26. When scaling out (adding more instances), VM migration can be used to redistribute the load more evenly across the newly available instances.
27. Else if
28. $(T_{size} - = VM's - -)$
29. Apply $(A_s \rightarrow OUT)$ // During scaling in (reducing the number of instances), VM migration plays a critical role in consolidating workloads onto fewer instances before shutting down the underutilized ones.
30. Compute $(C_a \rightarrow CSP)$ // cost analysis;
31. $VM_{Total} \rightarrow$ Total number of VM's
32. $T_t \rightarrow$ Total number number tasks
33. $T_{cpu} \rightarrow$ taotal number of CPU usage
34. $T_{RAM} \rightarrow$ taotal number of RAM usage
35. End process;

Algorithm 1 Algorithm for Tri-Zone Utilization Model

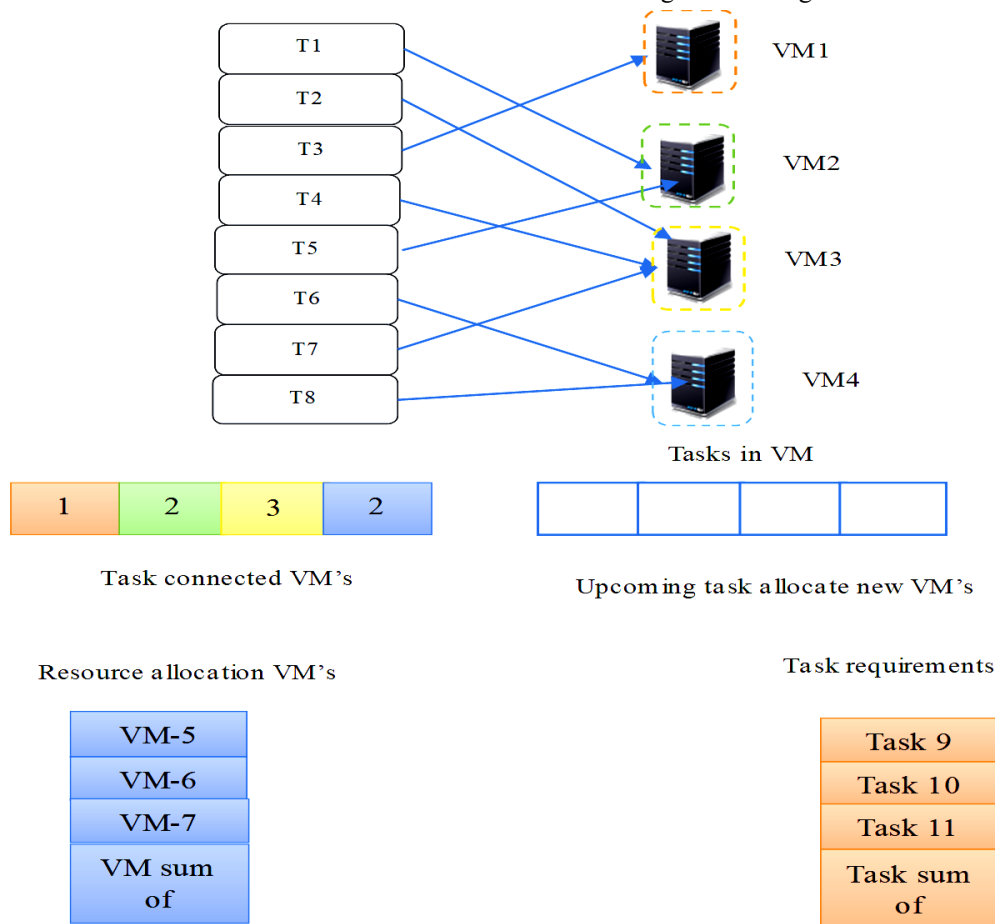


Figure 2 Working architecture

RESEARCH ARTICLE

Task allocation, VM load balancing, resource allocation and VM migration mathematical equation for these parameters.

3.1. Task Allocation

Task allocation involves assigning tasks to virtual machines (VMs) based on various criteria such as task requirements, VM capacity, and optimization goals.

3.1.1. Variables

- T : Set of tasks $\{T_1, T_2, \dots, T_n\}$
- V : Set of VMs $\{V_1, V_2, \dots, V_m\}$
- c_{ji} : Computational requirement of task T_j on VM V_i
- x_{ji} : Binary decision variable, where $x_{ji} = 1$ if task T_j is assigned to VM V_{ji} , and $x_{ji} = 0$ otherwise

3.1.2. Constraints

Capacity Constraint: Each VM has a limited capacity, so the total resource requirement of tasks assigned to a VM should not exceed its capacity C_i : is shown the equation (1)

$$\sum_{j=1}^n c_{ji} \cdot x_{ji} \leq c_i \quad \forall i \quad (1)$$

3.1.3. Objective Function

Load Minimization: Minimize the maximum load across all VMs shown the equation (2):

$$\text{Minimize } \max_i \left(\sum_{j=1}^n c_{ji} \cdot x_{ji} \right) \quad (2)$$

3.2. VM Load Balancing

VM load balancing aims to distribute the workload evenly across VMs to prevent some from being overloaded while others are underutilized.

3.2.1. Variables

- L_i : Load on VM V_i

3.2.2. Objective Function

- Minimize Load Variance: Minimize the variance of the load across all VMs to achieve balance:

$$\text{Minimize } \text{Var} (L_1, L_2, \dots, L_m)$$

Alternatively, minimizing the maximum load can also be used:

$$\text{Minimize } \max_i L_i$$

3.3. Resource Allocation

Resource allocation involves assigning resources like CPU, memory, and storage to VMs based on their requirements and available capacity.

3.3.1. Variables

- R_i : Total resource i available
- D_{ji} : Demand of resource i by VM j
- A_{ji} : Allocation of resource i to VM j

3.3.2. Constraints

Resource Capacity: The total allocation of a resource should not exceed its availability. shown the equation (3).

$$\sum_{j=1}^m A_{ji} \leq R_i \quad \forall i \quad (3)$$

Demand Satisfaction: The allocation should meet or exceed the demand. shown the equation (4):

$$A_{ji} \geq D_{ji} \quad \forall j, i \quad (4)$$

3.3.3. Objective Function

Maximize Utilization: Maximize the utilization of resources, or alternatively minimize the unutilized resources: shown the equation (5)

$$\text{Maximize } \sum_{i=1}^n \sum_{j=1}^m A_{ji} \quad (5)$$

3.4. VM Migration

VM migration involves moving a VM from one physical server to another, often to achieve load balancing or to free up resources.

3.4.1. Variables

- M_{ij} : Binary decision variable, where $M_{ij} = 1$ if VM V_j is migrated to server s_i , and $M_{ij} = 0$ otherwise
- C_{ij} : Cost of migrating VM V_j to server S_i

3.4.2. Constraints

- Single Migration: A VM can only be migrated to one server at a time. shown the equation (6)

$$\sum_{i=1}^p M_{ij} \leq 1 \quad \forall j \quad (6)$$

- Resource Constraints: Ensure the target server can accommodate the VM's resource requirements:

$$\text{Resource Usage on } S_i \leq \text{Capacity of } S_i \quad \forall i$$

3.4.3. Objective Function

Minimize Migration Cost: Minimize the total cost of migration: shown the equation (7).

$$\text{Minimize } \sum_{i=1}^p \sum_{j=1}^m C_{ij} \cdot M_{ij} \quad (7)$$

These formulations provide a structured way to approach resource management in cloud computing, balancing multiple

RESEARCH ARTICLE

factors like resource utilization, load distribution, and migration costs.

4. EXPERIMENTAL RESULTS

4.1. Fuzzy Logic Optimization

4.1.1. Fuzzification

In Fuzzy logic optimization (Figure 3), Fuzzification is the process of converting the accurate and precise data into fuzzy sets. The attributes such as CPU capacity, task size, and RAM speed are expressed in fuzzy sets which are in linguistic variables as "Low," "Medium," and "High.". To detect the imprecise nature of real time data, this phase enables the optimization process to use linguistic concepts and uncertainty.

4.1.2. CPU Capacity

Linguistic variables: Low, Medium, High Membership functions: The degree of membership function in each category is expressed in triangular or trapezoidal functions. The fuzzy membership terms is denoted in figure 4(a).

4.1.3. Task Size Linguistic Variables

Small, Medium, Large Membership functions: The membership functions expressed in triangular or trapezoidal function which shows in figure 4(b).

4.1.4. RAM and Storage Speed

Linguistic variables: Slow, Moderate, Fast Membership functions: The fuzzy terms of RAM speed are represented in triangular or trapezoidal function. The membership function is shown in figure 4(c) and also explained the network usage in figure 4(d).

4.1.5. Control Knowledge

In Fuzzy logic system, a set of rules is defined in control knowledge in which the fuzzy output variable is produced by the correlation between the fuzzy input variables. Based on the expert knowledge or principles, the rules are described as "if-then" statement.

For instance: "If CPU is high and Task Size is Large, then Migration is High.". The proper decision is made using these rules in fuzzy optimization.

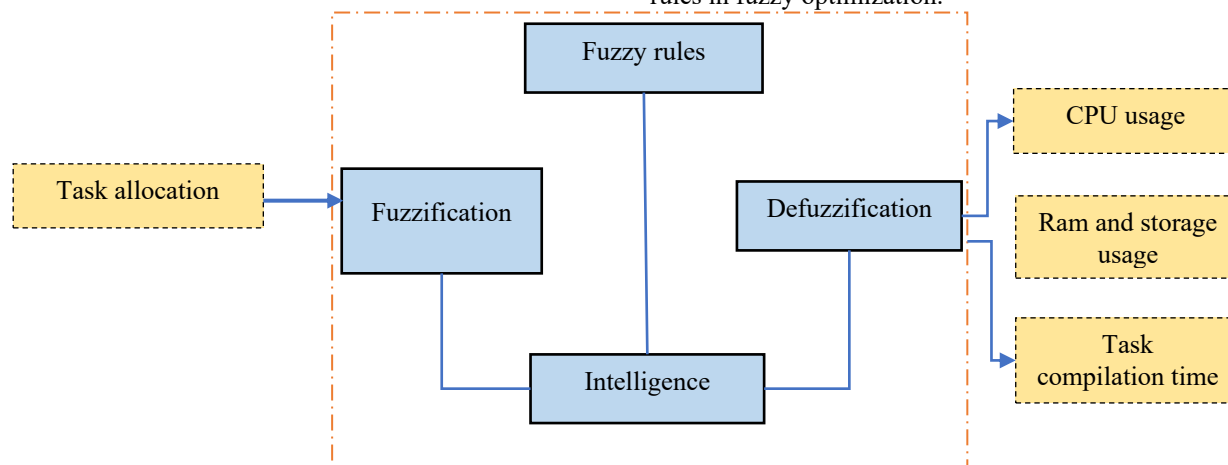


Figure 3 Fuzzy Logic Optimization

4.2. Fuzzy Output Variable

Based on the prime fuzzy output variable "Migration" level, the task has been migrated in fuzzy logic model. The linguistic output variable expressed as Low, Medium, and High along with respective membership functions.

4.2.1. Defuzzification

Defuzzification is the process of converting the fuzzy output variable into an efficient and precise outcome. Once applied fuzzy sets, the result expressed as fuzzy sets.

Following this, defuzzification convert the fuzzy sets into arithmetic values. In this case, this arithmetic value represents

an appropriate migration level of task. The standard defuzzification methods are centroid, bisector and mean of maximum. Fuzzification in Fuzzy logic model transforms the inaccurate data into linguistic terms. Then apply control knowledge to determine the rules for making well informed decision and convert the fuzzy output into insightful recommendations.

The Fuzzy logic model is implemented in our optimization module in order to address the inherent uncertainty in the cloud system. It offers a resource allocation decision making more flexible and accurate. The fuzzy rules formulated in table 1.

RESEARCH ARTICLE

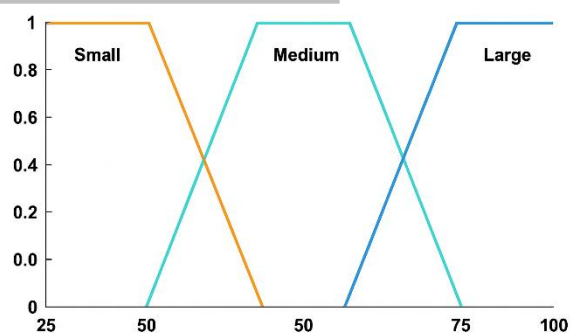


Figure 4 (a) CPU Usage

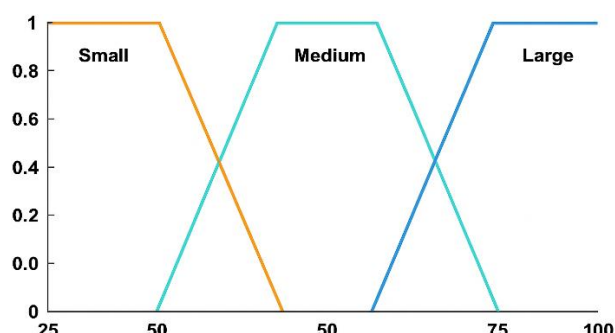


Figure 4 (c) Storage Utilization

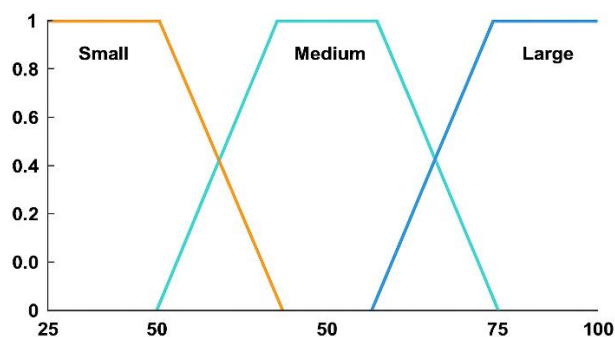


Figure 4 (b) Memory Usage

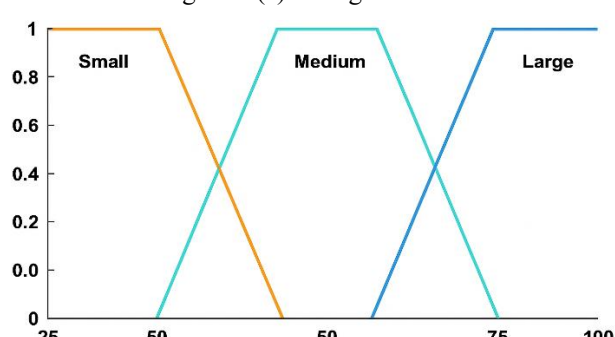


Figure 4(d) Network Bandwidth

Table 1 Fuzzy Rule

Task ID	Task size	CPU	RAM & storage	Task execution time
T1	Medium	Medium	MEDIUM	Medium
T2	Low	Medium	LOW	Medium
T3	High	Low	Medium	Medium
T4	Medium	Low	High	High
T5	High	Medium	Low	High
T6	Medium	Medium	Low	Yes
T7	Medium	High	High	Medium
T8	Low	High	Medium	High
T9	Medium	Medium	Medium	Medium
T10	High	Medium	Medium	Medium
T11	Low	Medium	High	High
T12	Low	High	Medium	Medium
T13	Medium	Medium	Low	Medium

RESEARCH ARTICLE

T14	Low	Medium	Low	Medium
T15	High	Low	Medium	Medium
T16	Medium	Low	High	High
T17	High	Medium	Low	High
T18	Medium	Medium	Low	Medium
T19	High	Medium	Medium	Medium
T20	High	Medium	Medium	Medium

4.3. Number of Tasks Vs Average Resource Utilization

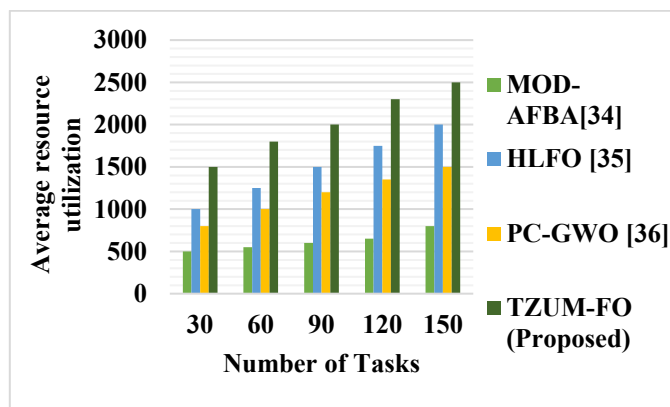


Figure 5 Number of Tasks vs Average Resource Utilization

An average resource utilization of four different models is compared and expressed in bar chart in figure 5. Our proposed Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO) proposed is compared with other models like Hybrid Lion Fuzzy Optimization (HLFO)[35], Priority-based Cloud Grey Wolf Optimizer (PC-GWO) [36], Modified Adaptive Firefly-Based Algorithm (MOD-AFBA). TZUM-FO model constantly achieving the average resource utilization by 90% when the task count rises from 30-150. On the other hand, The PC-GWO and MOD-AFBA reaches 50% and 70% respectively. The lowest resource utilisation is achieved by HLFO of about 40% which is least among all other models. These findings demonstrate that TZUM-FO performs better in diverse workload management, which reaches is 50% more than HLFO, like wise 40 % and 20% higher than PC-GWO and MOD-AFBA respectively.

4.4. Number of Tasks vs Average Energy Consumption

Figure 6 illustrates the performance of average energy consumption by the proposed Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO) with other existing models like Hybrid Lion Fuzzy Optimization (HLFO), Priority-based Cloud Grey Wolf Optimizer (PC-GWO) and Modified Adaptive Firefly-Based Algorithm (MOD-AFBA). The performance obtained by each method is

plotted in the graph-axis denotes the number of tasks which is increased by 30 respectively. The average energy consumption is expressed in joules (J) and is denoted in Y-axis. The energy consumed by TZUM-FO for 30 tasks is 5j whereas the MOD-AFBA and PC-GWO consumes energy by approximately 15j and 25j respectively. HLFO utilises more power by around 35j for 30 tasks where our proposed method consumes less power by about 30j for 150 tasks It is clearly shows that our proposed TZUM-FO consumes less energy than the other models.

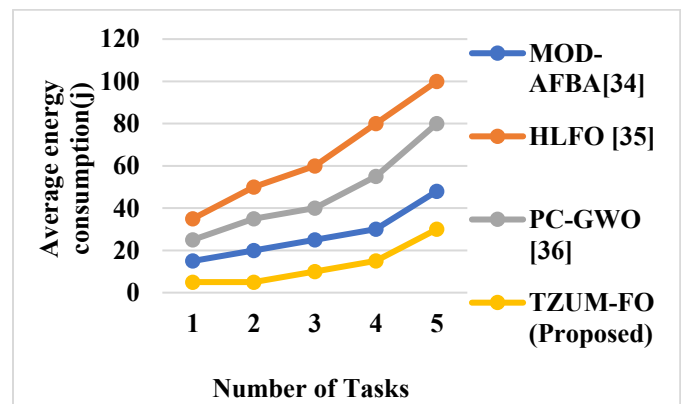


Figure 6 Number of Tasks vs Average Energy Consumption

4.5. Number of Tasks vs Make Span Time

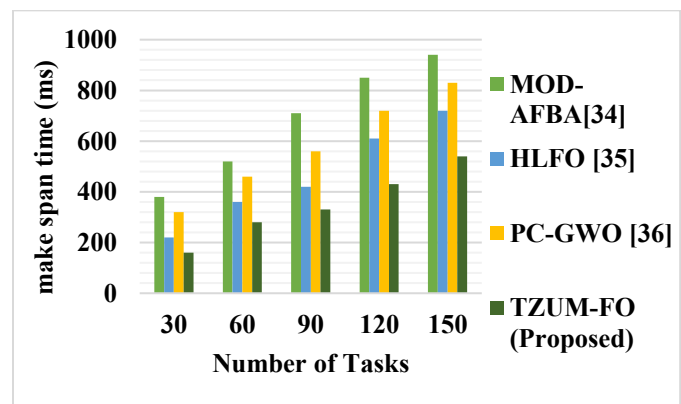


Figure 7 Number of Tasks vs Make Span Time

RESEARCH ARTICLE

In figure 7, compares the make span time of various workloads of the proposed TZUM-FO method with other models like Hybrid Lion Fuzzy Optimization (HLFO), Priority-based Cloud Grey Wolf Optimizer (PC-GWO), Modified Adaptive Firefly-Based Algorithm (MOD-AFBA). Our proposed TZUM-FO method utilizes dynamic load balancing algorithm which reduces the make span time of about 525s for 150 tasks whereas the MOD-AFBA, PC-GWO and HLFO requires make span of approximately 710s, 820s and 950s respectively. The experimental results shows that our proposed TZUM-FO response in minimum time for varying workloads from 30-150, comparing to other models like MOD-AFBA, PC-GWO and HLFO.

4.6. Number of Tasks Vs Balanced CPU Utilization

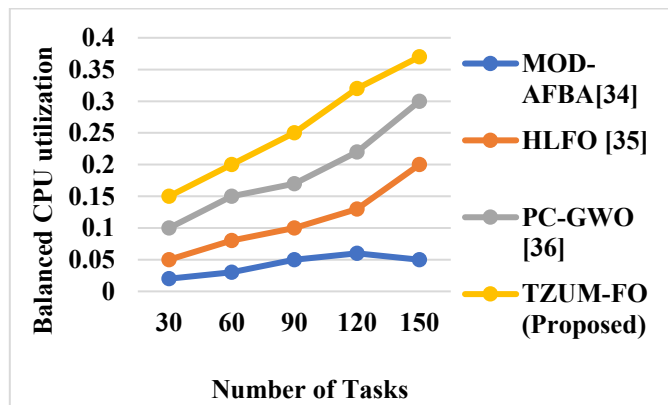


Figure 8 Number of Tasks vs Balanced CPU Utilization

Figure 8 illustrates the performance of overall CPU utilization of the proposed Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO) is compared with other existing models like Hybrid Lion Fuzzy Optimization (HLFO), Priority-based Cloud Grey Wolf Optimizer (PC-GWO) and Modified Adaptive Firefly-Based Algorithm (MOD-AFBA). In this experiment, evaluate the CPU usage for various workloads raising form 30 to 150. The minimum CPU utilization among the four models is TZUM-FO requires approximately 0.01 for 30 tasks whereas the HLFO requires more utility rate of about 0.15. The other models like MOD-AFBA and PC-GWO utilized resource of about 0.05 and 0.1 respectively. The loads are progressively raised to four algorithms at each level, and the outputs are graphed and recorded. The outcomes shows that our proposed TZUM-FO requires minimum CPU utilization than all other existing models.

4.7. Number of Tasks Vs VM Migration

In figure 9, the migration of Virtual Machines (VM) for varying workloads of four resource allocation models like Hybrid Lion Fuzzy Optimization (HLFO), Priority-based Cloud Grey Wolf Optimizer (PC-GWO), Modified Adaptive Firefly-Based Algorithm (MOD-AFBA), and the proposed

Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO). The performance of VM migration achieved by the four models is TZUM-FO > MOD-AFBA > PC-GWO > HLFO reaches 18 > 14 > 8 > 4 for 120 tasks. The results indicate that our proposed TZUM-FO assign the tasks to the available VMs when overloaded and achieve an efficient resource allocation in comparison to all other models like MOD-AFBA, PC-GWO and HLFO.

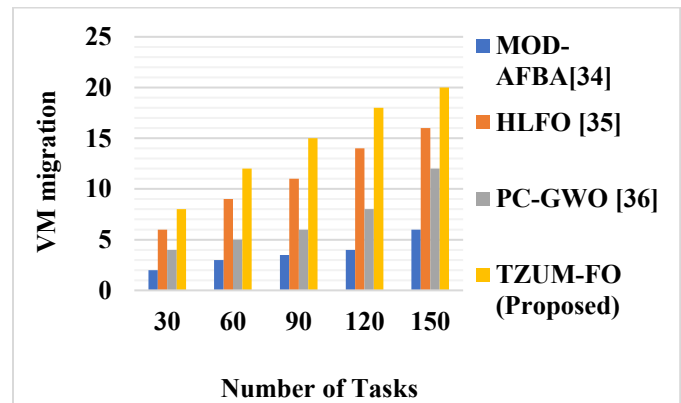


Figure 9 Number of Tasks vs VM Migration

4.8. Number of Tasks Vs Optimized Memory Usage

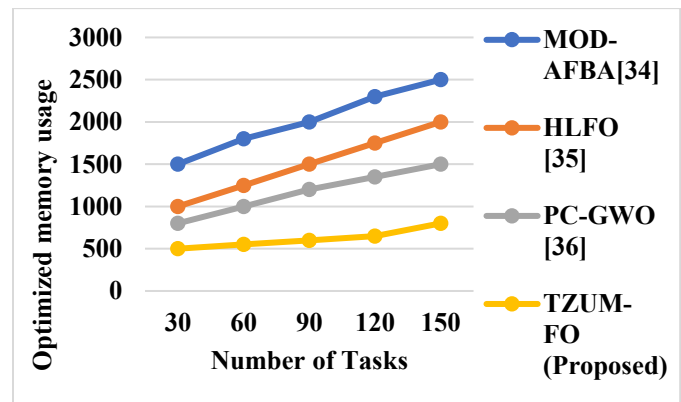


Figure 10 Number of Tasks Vs Optimized Memory Usage

Figure 10 describes the performance of overall optimized memory usage for varying workloads of the proposed Tri-Zone Utilization Model with Fuzzy Logic Optimization (TZUM-FO) is compared with other existing models like Hybrid Lion Fuzzy Optimization (HLFO), Priority-based Cloud Grey Wolf Optimizer (PC-GWO) and Modified Adaptive Firefly-Based Algorithm (MOD-AFBA). The memory utilized by the proposed TZUM-FO of about 500 for 30 tasks whereas MOD-AFBA, PC-GWO and HLFO requires memory of approximately 850, 1000 and 1500 respectively. The loads are progressively increased to four algorithms at each level and the outputs are graphed and recorded. The results shows that our proposed TZUM-FO requires less

RESEARCH ARTICLE

memory and performs better than other existing models like MOD-AFBA, PC-GWO and HLFO.

4.9. Number of Tasks Vs System Efficiency

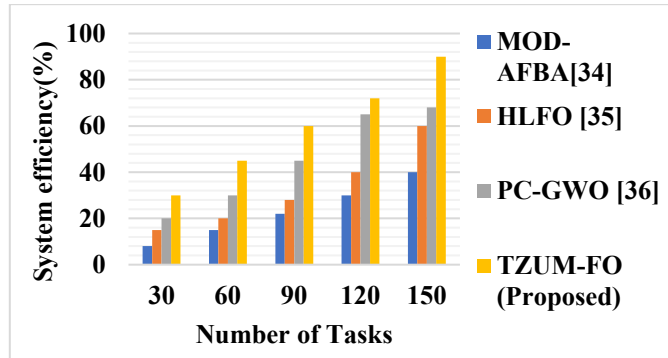


Figure 11 Number of Tasks Vs System Efficiency

In figure 11, the overall system efficiency of the proposed TZUM-FO approach is compared with existing models like Hybrid Lion Fuzzy Optimization (HLFO), Priority-based Cloud Grey Wolf Optimizer (PC-GWO), Modified Adaptive Firefly-Based Algorithm (MOD-AFBA) for varying workloads. As the number of tasks increases from 30 to 150, the TZUM-FO model consistently shows the highest system efficiency, peaking at around 96%. In contrast, the HLFO model exhibits the lowest utilization, reaching only about 40% at its peak. The PC-GWO and MOD-AFBA models perform better than HLFO, achieves efficiency of approximately 60% and 63 % respectively at higher task levels. The results indicate that TZUM-FO achieves the highest system efficiency, outperforming the other existing models like MOD-AFBA, PC-GWO and HLFO.

5. CONCLUSION

This study presents a Tri-Zone Utilization Model integrated with Fuzzy Logic Optimization (TZUM-FO) to enhance the efficiency and reliability of cloud resource management. The proposed TZUM-FO model enables accurate and adaptive allocation of virtual machines (VMs) by classifying incoming tasks into three categories such as small, medium, and large based on their computational requirements. The integration of dynamic auto-scaling and intelligent VM migration ensures balanced load distribution and improved resource utilization across the cloud infrastructure. A comparative analysis between the proposed TZUM-FO and existing models, including HLFO, PC-GWO, and MOD-AFBA, demonstrates that TZUM-FO consistently achieves superior performance. It results in shorter make span time, lower average energy consumption, optimized CPU and memory usage, and more efficient VM migration. These outcomes highlight TZUM-FO's effectiveness in maintaining performance stability while minimizing operational costs. The practicality and robustness of TZUM-FO have also been established through its

scalability and compatibility across different cloud settings. This framework can be applied in real-world cloud environments such as AWS and Microsoft Azure, where it can improve flexibility, cost-efficiency, and operational stability for businesses using multi-cloud systems. In future work, the TZUM-FO model will be further enhanced and extended to ensure sustained efficiency and adaptability in dynamic and large-scale cloud computing environments.

REFERENCES

- [1] Li D., Wang W., Li Q., Ma D. Study of a Virtual Machine Migration Method; Proceedings of the 2013 International Conference on Advanced Cloud and Big Data; Nanjing, China. 13–15 December 2013; pp. 27–31.
- [2] Jain A., Kumar R. A Multi Stage Load Balancing Technique for Cloud Environment; Proceedings of the 2016 International Conference on Information Communication and Embedded Systems (ICICES); Tamil nadu, India. 25–26 February 2016; pp. 1–7.
- [3] Beltran M., Guzmán A., Bosque J.L. Dealing with Heterogeneity in Load Balancing Algorithms; Proceedings of the 2006 Fifth International Symposium on Parallel and Distributed Computing; Timisoara, Romania. 6–9 July 2006; pp. 123–132.
- [4] Van H.N., Tran F.D., Menaud J.M. Autonomic Virtual Resource Management for Service Hosting Platforms; Proceedings of the 2009 ICSE Workshop on Software Engineering Challenges of Cloud Computing; Vancouver, BC, Canada. 23 May 2009; pp. 1–8.
- [5] A. Beloglazov, J. Abawajy, and R. Buyya, “Energy-aware resource allocation heuristics for efficient management of data centers for cloud computing,” *Future Generation Computer Systems*, vol. 28, no. 5, pp. 755–768, 2012.
- [6] Aransay I., Zapater M., Arroba P., Moya J.M. A Trust and Reputation System for Energy Optimization in Cloud Data Centers; Proceedings of the 2015 IEEE 8th International Conference on Cloud Computing; New York, NY, USA. 27 June–2 July 2015; pp. 138–145.
- [7] Jeyarani, R.; Nagavent, N.; Ram, R.V. Design and implementation of adaptive power-aware virtual machine provisioner (APAVMP) using swarm intelligence. *Future Gener. Comput. Syst.* 2012, 28, 811–821.
- [8] X. Zuo, G. Zhang, and W. Tan, “Self-adaptive learning PSO-based deadline constrained task scheduling for hybrid IaaS cloud,” *IEEE Trans. Autom. Sci. Eng.*, vol. 11, no. 2, pp. 564–573, April 2014.
- [9] X. Wang, Z. Du, Y. Chen, and M. Yang, “A green-aware virtual machine migration strategy for sustainable datacenter powered by renewable energy,” *Simulation Modelling Practice and Theory*, vol. 58, pp. 3–14, 2015.
- [10] Masker, M.; Nagel, L.; Brinkmann, A.; Lotffar, F.; Johnson, M. Smart grid-aware scheduling in data centers. In *Proceedings of the 2015 Sustainable Internet and ICT for Sustainability (SustainIT)*, Funchal, Portugal, 6–7 December 2015; pp. 1–9.
- [11] A. Khosravi, L. L. H. Andrew, and R. Buyya, “Dynamic VM placement method for minimizing energy and carbon cost in geographically distributed cloud data centers,” *IEEE Transactions on Sustainable Computing*, vol. 2, no. 2, pp. 183–196, Apr. 2017.
- [12] Varasteh, A.; Tashtarian, F.; Goudarzi, M. On Reliability-Aware Server Consolidation in Cloud Datacenters. In *Proceedings of the 2017 16th International Symposium on Parallel and Distributed Computing (ISPDC)*, Innsbruck, Austria, 3–6 July 2017; pp. 95–101.
- [13] Kasture, H. A Hardware and Software Architecture for Efficient Datacenters. Ph.D. Thesis, Department of Electrical Engineering and Computer, MIT, Cambridge, MA, USA, February 2017.
- [14] D. Arivudainambi and D. Dhanya, “Towards optimal allocation of resources in cloud modified Map-Reduce using genetic algorithm,” *IOAB J.*, vol. 8, no. 2, pp. 162–171, 2017.
- [15] H. Li, W. Li, H. Wang and J. Wang, “An optimization of virtual machine selection and placement by using memory content similarity

RESEARCH ARTICLE

- for server consolidation in cloud,” *Future Gener. Comput. Syst.*, vol. 84, pp. 98–107, 2018.
- [16] L. Grange, G. Da Costa, and P. Stolf, “Green IT scheduling for data center powered with renewable energy,” *Future Generation Computer Systems*, vol. 86, pp. 99–120, 2018.
- [17] S. K. Mishra, D. Puthal, B. Sahoo, P. P. Jayaraman, S. Jun, A. Y. Zomaya, and R. Ranjan, “Energy-efficient VM-placement in cloud data center,” *Sustain. Comput.: Inform. Syst.*, vol. 20, pp. 48–55, 2018.
- [18] T. P. Shabeera, S. D. Madhu Kumar, M. S. Sameera, and K. Murali Krishnan, “Optimizing VM allocation and data placement for data-intensive applications in cloud using ACO metaheuristic algorithm,” *Eng. Sci. Technol., Int. J.*, vol. 20, pp. 616–628, 2017.
- [19] G. Han, W. Que, G. Jia, and W. Zhang, “Resource-utilization-aware energy efficient server consolidation algorithm for green computing in IIoT,” *J. Netw. Comput. Appl.*, vol. 103, pp. 205–214, 2018.
- [20] M. Tavana, S. Shahdi-Pashaki, E. Teymourian, F. J. S. Arteaga, and M. Komaki, “A discrete cuckoo optimization algorithm for consolidation in cloud computing,” *Computers & Industrial Engineering*, vol. 115, pp. 495–511, 2018.
- [21] Malekloo, M.H.; Kara, N.; El Barachi, M. An energy efficient and SLA compliant approach for resource allocation and consolidation in cloud computing environments. *Sustain. Comput. Inform. Syst.* 2018, 17, 9–24.
- [22] Guha Neogi, P.P. Cost-Effective Dynamic Workflow Scheduling in IaaS Cloud Environment. In *Proceedings of the 2019 International Conference on Intelligent Computing and Remote Sensing (ICICRS)*, Bhubaneswar, India, 19–20 July 2019; pp. 1–6.
- [23] M. Liaqat, A. Naveed, R. L. Ali, J. Shuja, and K.-M. Ko, “Characterizing dynamic load balancing in cloud environments using virtual machine deployment models,” *IEEE Access*, vol. 7, pp. 145 767–145 776, 2019.
- [24] Mustafa, K. Sattar, J. Shuja, S. Sarwar, T. Maqsood, S. A. Madani, and S. Guizani, “SLA-Aware Best Fit Decreasing Techniques for Workload Consolidation in Clouds,” *IEEE Access*, vol. 7, pp. 135 256–135 267, 2019.
- [25] F. Alzhouri, S. B. Melhem, A. Agarwal, M. Daraghme, Y. Liu, and S. Younis, “Dynamic Resource Management for Cloud Spot Markets,” *IEEE Access*, vol. 8, pp. 122 838–122 847, 2020.
- [26] Nehra P., Nagaraju A. Sustainable Energy Consumption Modeling for Cloud Data Centers; *Proceedings of the 2019 IEEE 5th International Conference for Convergence in Technology (I2CT)*; Pune, India. 29–31 March 2019; pp. 1–4.
- [27] 20. Diouani S., Medromi H. How Energy Consumption in the Cloud Data Center is Calculated; *Proceedings of the 2019 International Conference of Computer Science and Renewable Energies (ICCSRE)*; Agadir, Morocco. 22–24 July 2019; pp. 1–10.
- [28] Duan L., Zhan D., Hohnerlein J. Optimizing Cloud Data Center Energy Efficiency via Dynamic Prediction of CPU Idle Intervals; *Proceedings of the 2015 IEEE 8th International Conference on Cloud Computing*; New York, NY, USA. 27 June–2 July 2015; pp. 985–988.
- [29] Goyal Y., Arya M.S., Nagpal S. Energy Efficient Hybrid Policy in Green Cloud Computing; *Proceedings of the 2015 International Conference on Green Computing and Internet of Things (ICGCIoT)*; Greater Noida, India. 8–10 October 2015; pp. 1065–1069.
- [30] Jiang D., Zhang Y., Song H., Wang W. Intelligent Optimization-Based Energy-Efficient Networking in Cloud Services for Multimedia Big Data; *Proceedings of the 2018 IEEE 37th International Performance Computing and Communications Conference (IPCCC)*; Orlando, FL, USA. 17–19 November 2018; pp. 1–6.
- [31] Jasuja K.P., Kaur K. Hybrid Soft Computing Approach for Energy Efficiency in Cloud Computing; *Proceedings of the 2016 International Conference on Communication and Electronics Systems (ICES)*; Coimbatore, India. 21–22 October 2016; pp. 1–4.
- [32] Khan A., Papliński A.P., Khan A.M., Murshed M., Buyya R. Exploiting User Provided Information in Dynamic Consolidation of Virtual Machines to Minimize Energy Consumption of Cloud Data Centers; *Proceedings of the 2018 Third International Conference on Fog and Mobile Edge Computing (FMEC)*; Barcelona, Spain. 23–26 April 2018; pp. 105–114.
- [33] A. Majeed and M. A. Shah, “Energy efficiency in big data complex systems: A comprehensive survey of modern energy saving techniques,” *Complex Adaptive Systems Modeling*, vol. 3, no. 1, Art. 6, 2015.
- [34] D. Alsadie and M. Alsulami, “Efficient Resource Management in Cloud Environments: A Modified Feeding Birds Algorithm for VM Consolidation,” *Mathematics*, vol. 12, no. 12, art. 1845, 2024.
- [35] A. R. Khan, “Dynamic Load Balancing in Cloud Computing: Optimized RL-Based Clustering with Multi-Objective Optimized Task Scheduling,” *Processes*, vol. 12, no. 3, art. 519, 2024.
- [36] N. A. Alsamarai and O. N. Uçan, “Improved Performance and Cost Algorithm for Scheduling IoT Tasks in Fog–Cloud Environment Using Gray Wolf Optimization Algorithm,” *Appl. Sci.*, vol. 14, art. 1670, 2024.
- [37] D. Saxena, S. R. Swain, and J. Kumar, “Workload Pattern Learning-Based Cloud Resource Efficiency,” *IEEE Transactions on Sustainable Computing*, vol. 10, no. 3, pp. 245–257, Mar. 2025.
- [38] Y. Wang and X. Yang, “Intelligent Resource Allocation via Deep and Reinforcement Learning,” *Advances in Computer, Signals and Systems*, vol. 9, no. 1, pp. 33–41, Jan. 2025.
- [39] L. Li, J. Bell, and M. Coppola, “Adaptive AI-Based Decentralized Management in Cloud-Edge Continuum,” *IEEE Access*, vol. 13, no. 2, pp. 1789–1803, Feb. 2025.
- [40] C. Panggabean, D. C. Verma, B. Gogoi, R. Limbu, and R. Sarker, “Optimized Cloud Resource Allocation Using Genetic Algorithms for Energy and QoS,” *IEEE Transactions on Cloud Computing*, vol. 13, no. 1, pp. 112–123, Jan. 2025.
- [41] D. Saxena, S. R. Swain, and J. Kumar, “Secure Resource Management in Cloud Computing: Challenges, Strategies, and Meta-Analysis,” *IEEE Transactions on Cloud Engineering*, vol. 5, no. 1, pp. 51–63, Feb. 2025.

Authors



Ms. K. Geetha is a Ph.D. Research Scholar (2022–Present) in the Department of Mathematics at Periyar Maniammai Institute of Science & Technology (PMIST), Thanjavur. She completed her M.Phil. in Mathematics in 2012, following an M.Sc. in 2007, a B.Ed. in 2009, and a B.Sc. in 2005, consistently achieving strong academic performance throughout her studies. Her research areas include fuzzy logic techniques, machine learning-based optimization, intelligent task scheduling, and real-time decision-support systems. Since 2019, her doctoral work has focused on developing fuzzy logic-based intelligent optimization methods for improving resource allocation in cloud healthcare systems. She has also contributed to interdisciplinary projects that integrate ANFIS and ANN models with real-time sensor data for enhanced rice crop yield prediction (published in 2024), showcasing her ability to merge mathematical modelling with intelligent computing. Ms. Geetha has co-authored research papers in areas such as precision agriculture (2024) and cloud healthcare task scheduling (2023–2024), with publications demonstrating the effectiveness of fuzzy inference systems and machine learning frameworks in solving real-world optimization challenges. Her strong academic foundation and ongoing doctoral research reflect a commitment to advancing mathematical approaches for sustainable and efficient intelligent systems.



Dr. P. Vijayalakshmi is an Associate Professor and Associate Dean in the Department of Mathematics at Periyar Maniammai Institute of Science & Technology (PMIST), Thanjavur. With a doctoral degree in Mathematics (2016) and over three decades of academic experience, she specializes in Category Theory, Fuzzy Logic Techniques, and Fuzzy Optimization. Her work spans diverse interdisciplinary

RESEARCH ARTICLE

domains including precision agriculture, cloud-based healthcare optimization, and corrosion prevention, contributing to impactful research published in Scopus-indexed and reputed international journals. Dr. Vijayalakshmi has guided significant research projects integrating ANFIS, ANN, and fuzzy inference systems for predictive modelling in agriculture, cloud task scheduling, and industrial corrosion management. Her collaborative works demonstrate strong expertise in combining mathematical modelling with intelligent computing and sensor-based real-time data analytics to develop sustainable and efficient solutions. Throughout her career, she has earned notable recognitions such as the Best Outstanding Teacher Award (2005–06) and multiple Best NCC Officer Awards (2009, 2011, 2016, 2017). She is proficient in MATLAB, Python, R, and SPSS, and is actively involved in curriculum development, academic leadership, and mentoring research scholars.

How to cite this article:

K. Geetha, P. Vijayalakshmi, “Advanced Tri-Zone Utilization Model with Fuzzy Logic Optimization for Efficient Task Allocation and VM Migration in Cloud Computing”, International Journal of Computer Networks and Applications (IJCNA), 12(6), PP: 973-988, 2025, DOI: 10.22247/ijcna/2025/57.