

Multi-Objective Binary Whale Optimization-Based Virtual Machine Allocation in Cloud Environments

Ankita Srivastava, Babasaheb Bhimrao Ambedkar University, India

Narander Kumar, Babasaheb Bhimrao Ambedkar University, India*

 <https://orcid.org/0000-0003-4633-1879>

ABSTRACT

With the rising demands for the services provided by cloud computing, virtual machine allocation (VMA) has become a tedious task due to the dynamic nature of the cloud. Millions of virtual machines (VMs) are allocated and de-allocated at every instant, so an efficient VMA has been a significant concern to enhance resource utilization and depreciate its wastage. Encouraged by the prodigious performance of the nature-inspired algorithm, the binary whale optimization approach has been eventuated to get to grips with the VMA issue with the focus on minimizing the resource waste and volume of servers working actively. The deliberate approach's accomplishment is assessed against the literature's well-known algorithms for VMA issues. The comparison results showed that the least resource wastage fitness of 15.68, minimum active servers of 216, and effective CPU and memory utilization of 88.31% and 88.79%, respectively, have been achieved.

KEYWORDS

Binary Whale Optimization Algorithm, Cloud Computing, Metaheuristic Algorithm, Resource Allocation, Resource Wastage, Swarm Optimization, Virtual Machine Allocation, Virtualization

1. INTRODUCTION

The progression in technology has led to the emergence and evolvement of cloud computing from the basic computing paradigm of distributed, parallel, and grid computing. It has revolutionized the procedure for managing the data or information and the resources which have impacted human society socially and economically. The services are offered to the users in the resources form (e.g., storage, CPU, servers, network, and application). These resources are organized at one central point, the data center, from where accessing these resources comes at ease. The overall management of data centers is the responsibility borne by the service providers. The service providers serve the users' interests with three basic services via the internet IaaS for hardware resources, PaaS for runtime environment,

DOI: 10.4018/IJSIR.317111

*Corresponding Author

This article published as an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>) which permits unrestricted use, distribution, and production in any medium, provided the author of the original work and original publication source are properly credited.

and SaaS for software resources. These former services are being accomplished with virtualization. Virtualization facilitates the creation and configuration of VMs possessing variant operating systems with varying resources (memory, CPU, and storage) that are then nailed on the host machine to furnish the services desired by users. It also expedites sharing of multiple VMs deployed on one distinct server and sharing the hardware resources. To accomplish the request or interest demanded by users, VMs are created and configured dynamically with variable configuration and resource demands. The key objective in this reference is to allocate resources in ways that they are effectively utilized, decreasing resource waste and resulting in low operational costs. Adopting an effective VMA algorithm is one way to accomplish this. VMA allows VMs to be placed on the servers such that computing resources must be efficiently utilized while also reducing the volume of active servers. With the decisive allocation techniques, there come challenges, including a reduction in QoS, reduction of performance with the curtailed energy usage, and then satisfying the user's QoS within the promised SLA. An effective VMs allocation narrows down the count of active and balanced multidimensional resource usage by servers, ultimately reducing the energy expenditure by the cloud. Figure 1 demonstrates the allocation of VMs in which all the VMs are allotted randomly to the server, while after optimization in figure 2, the VMs are consolidated in the first three servers, and the rest of the unutilized servers are turned off, thus saving the resources and energy. Identifying optimal VMs allocation belongs to the NP-complete problem (Lo, V. M. 1983), and achieving the optimum resolution of this is typically computationally infeasible, when the cloud involves multiple hosts and users (Widmer, T., Premm, M., & Karaenke, P. 2013). The issue can be addressed to minimize resource wastage and the number of active servers in the cloud data center. Various methods have been applied to resolve this issue in the literature. A theorem was given (Wolpert, D. H., & Macready, W. G. 1997), which stated that not a single meta-heuristic algorithm is available which can efficiently resolve all the optimization problems. In this regard, the WOA (Mirjalili & Lewis, 2016) has been successfully utilized for various optimization problems (Prakash, D.B. & Lakshminarayana, C., 2017; Sun, Wang, Chen, & Liu, 2018; Chen, H., Xu, Y., Wang, M., & Zhao, X., 2019; Too, J., Mafarja, M., & Mirjalili, S., 2021).

A meta-heuristic-based Genetic Algorithm (GA) has been used widely to resolve the allocation problem (Kaaouache, M. A., & Bouamama, S., 2018; Abohamama, A. S., & Hamouda, E., 2020). These methods are the victims of basically two issues. These algorithms converge when a large number of iterations are performed and claim for a large population size as opposed to other metaheuristic algorithms. Due to this, it involves a high processing time, and sometimes it also suffers from premature convergence. While other meta-heuristic algorithms like ant colony-based, grey-wolf-based algorithms converge with lower population size. This motivates exploiting different meta-heuristic algorithms that benefit from low population size and iterations or generations. The study introduces a novel state of the art for accomplishing VM allocation, which subsequently truncates the active servers' count and minimizes the undergoing resource wastage, thus enhancing the proportion of resource utilization.

Further, the study is assembled as section 2 background study for the associated work being performed, and section 3 demonstrates the proposed method. The algorithm for the given work is discussed in section 4, which is supported through simulation. The result analysis is done in section 5, and further discussion is performed in section 6, which is finally concluded in section 7.

2. LITERATURE REVIEW

Extensive advances in the cloud have led to a dramatic rise in the volume of data centers. Hence, energy consumption generated by these centers continues to upsurge the cost incurred by the cloud system. So, in this scenario mitigating the energy exhausted by these centers has now become a critical concern that needs urgent attention. Energy consumed through these data centers varies directly with resource usage, thus, to prevent energy consumption, one needs to be careful about resource consummation as it aids in reducing operating costs. The resources should be used proficiently, thus

Figure 1. VMA before optimization

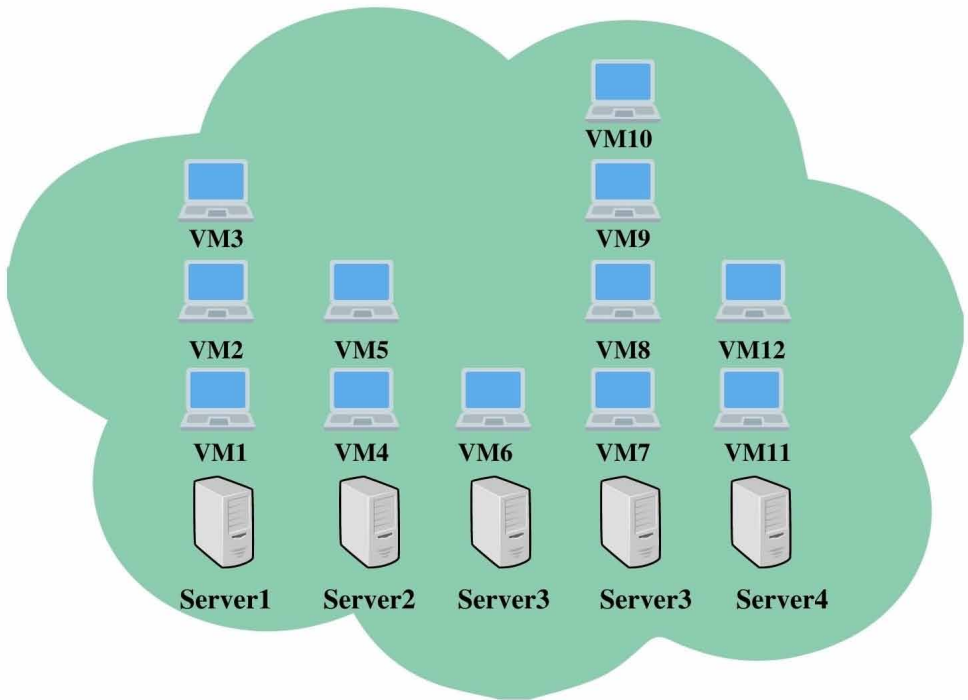
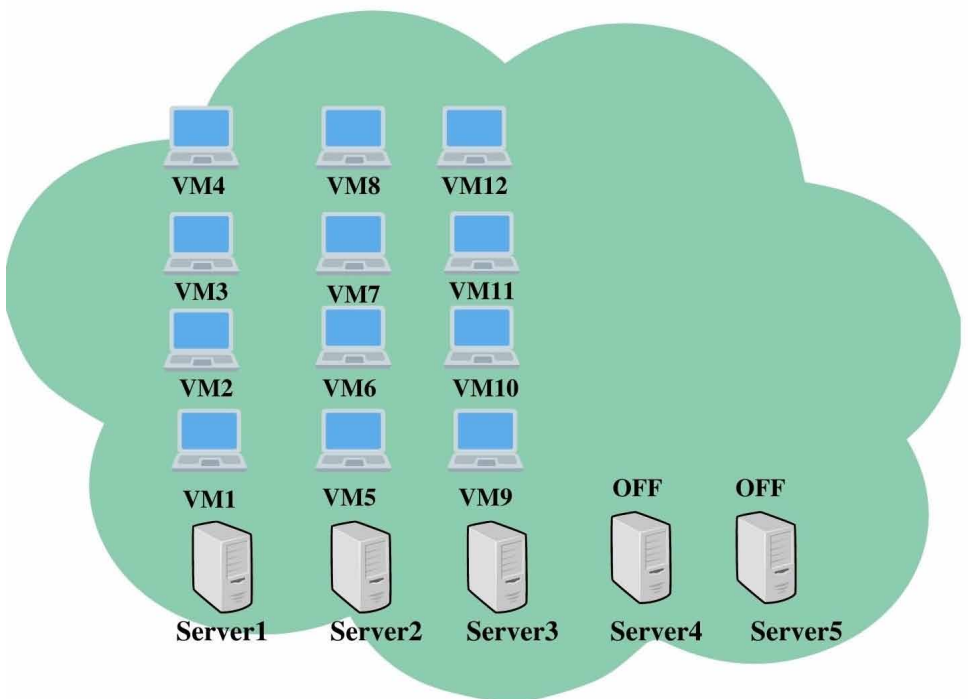


Figure 2. VMA after optimization



minimizing resource wastage. The most effective method to accomplish this is utilizing virtualization to allocate resources effectively.

Various learning has been put forward which are proficient in resource allocation. The most basic way to confront the VMA problem is by using linear programming. Stochastic Integer Programming is cast-off to reduce the price for hosting VMs in multiple cloud environments (Chaisiri, S., Lee, B. S., & Niyato, D., 2009). Heuristic algorithms encompass the Best Fit, or First Fit algorithms are also used. A Modified Best Fit Decrease algorithm is implemented to consolidate virtual machines (Beloglazov, A., & Buyya, R., 2010). The estimation module is applied to predict the future load of the system, and the scheduler is employed to schedule the expected and unpredicted loads. In turn, these schedulers enforce the technique of column generation to exploit integer linear or quadratic programming optimization problems (Vakiliinia, S., Heidarpour, B., & Cheriet, M., 2016). These methods produce the locally optimal solution, bringing about the imbalanced consumption of energy and resources and leading to greater SLA violations.

Some scholars have used heuristic methods that produce higher efficiency in tackling the VMA problem. The heuristic method is unified with the machine learning technique for conducting VMA which is correlated with traffic, energy, migration, QoS, and scalability (Pahlevan, A., Qu, X., Zapater, M., & Atienza, D. 2017). An improved analytic hierarchy procedure is presented, which processes the task before the allocation, and allocation is performed using optimization methods considering the given load and bandwidth (Gawali, M. B., & Shinde, S. K., 2018). Another approach that considers VMs to server allocation initially utilizes the renowned First Fit heuristic and MBFD and then performs the iterative search for optimal mapping (Zhang, X., et al., (2019). These methods have not taken the underuse of resources as their main focus. Another approach to tackle this issue is the assignment method. The author has utilized a modified assignment method to optimize the VMs allocation by minimizing SLA violation and resource wastage (Toutov, A., et al., 2021). This method does not take into account the active server's volume.

Some scholars adopted intelligent computation for VMs allocation, which includes genetic algorithms. An allocation technique is being contrived in virtualized network function at the data center, where two genetic algorithms are analyzed for VMA and scaling (Rankothge, W., et al., 2017). A multi-objective-based genetic algorithm is presented for VMA, which performs the encoding of chromosomes using group method, crossover, and further mutation operations whose length is varying. It decodes the chromosome using a three-dimensional split method (Qiang, L., 2011). Another approach is the Ant Colony Algorithm which has been elaborated, where a collection of servers is selected to allocate VMs and data on it (Shabeera, T. P., et al., 2017). Resource allocation has been resolved as a convex optimization problem. The major drawback of this algorithm was that it was only applicable in the case of homogenous VMs, and it didn't consider the migration aspect of the VMs to utilize the resources efficiently. A deep reinforcement framework has been presented for allocating resources, resulting in energy efficiency for cloud radio access networks. In this, a learning agent having deep reinforcement is defined through several parameters like reward method, action space, and state space. Further, the action-value function is approximated using deep neural networks, and resource allocation is formulated formally (Peng, Y., et al., 2017). One of the shortcomings of this work was the large population size was considered, which has increased the processing time of the algorithm. Combinatorial optimization is undertaken, which considers Quantum Genetic Algorithm (GA) & Ant Colony Optimization (ACO) algorithm with GA for resolving resource allocation issues. A mapping process is performed between the resource allocation matrix and chromosomes of the aforesaid algorithm and aids in refining the resource utility (Xia, W., & Shen, L. 2018). It worked significantly on the small data set, but it didn't discuss the large dataset. An enhanced cuckoo search procedure is familiarized to confront the VMA issue, which effectively narrows down the energy dissipation. It is motivated by the brooding behavior of cuckoo birds. In this exploration, the process is performed via Levy flights which strived to perceive the best local solution among the specified collection of different solutions. Eventually, fitness function assessment was performed to foster the

new list for VMA. This algorithm works efficiently with fewer VMs (Barlaskar, E., et al., 2018). It didn't give any idea about CPU and memory consumption, which is one of the important aspects of VM allocation. An ACO system embedded with a new heuristic is explored. In this methodology, four steps are there, which are initializing the pheromone, defining heuristic information, construction solution, and ultimately the last step is updating the pheromone. Firstly, the pheromones are initialized in which the assignment of VMs is being performed to servers through the First Fit Decrease process. Then the heuristic routine is designed considering CPU utilization, energy exhaustion, CPU capacity, and further VM is set down on the server in accordance with this information. After these, ants construct the solution with the extreme probability utilizing the pseudo-random proportional rule. Finally, the pheromone is updated through pheromone evaporation, which attempts to substantiate the global best solution (Alharbi, F., et al., 2019). It has not considered the resource utilization and the dynamic migration of the VMs across PMs. Another nature-inspired heuristic algorithm is the Flower Pollination algorithm. It is executed on a methodology termed dynamic switching probability. A methodology is enacted that quickly identifies the closest optimal solution and maintains a balance between exploitation of local search and exploration for the global search. It takes into deliberation the storage, processor, and memory restraint of a server and allocates VMs to it by minimizing and emphasizing the consummation of energy (Usman, M. J., et al., 2019). The algorithm works for a single objective, which restricts the technique's broad work scope. An Intelligent Water Drop technique is used to allocate VMs by optimizing the task execution and security consideration (Dubey, K., & Sharma, S. C., 2020). A metaheuristic approach grounded on the binary gravitational search process is propounded, which is established on the law of gravity to perform energy-efficient VMA. In this procedure, agents are entitled as entities, and their capability is quantified in relation to their masses. Agent position is set randomly, and, in every iteration, the fitness assessment is accomplished, which signifies the mass per agent, and the fitness value is updated accordingly. Then the mass is determined for each agent based on updated fitness value which ultimately assists in the assessment of acceleration. Then a new position is calculated beside the new velocity (Abdessamia, F., et al., 2020). This approach only considers energy, but CPU usage and availability are not considered, and the algorithm was analyzed for 100 VMs only. A Fair Fit algorithm is put forward, effectively increasing resource use by restraining its imbalance across multiple dimensions (Gohil, B. N., et al., 2021). Energy efficiency, resource wastage, and performance have been addressed using neural networks. The future demands have been forecasted using a forward feed neural network, and then auto-scaling is performed based on the clustering of resource requirements predicted, and further VMs are being allocated to the servers (Saxena, D., & Singh, A. K., 2021). It suffered the drawback of manually selecting the number of nodes in the input and output layers of the predictor. While a perceptive priority-aware VM allocation is presented, which prioritizes the application based on memory, computing, and bandwidth using a machine learning-based model. The work is identified in mitigating the power consumption, average response time, and execution time (Savitha, S., & Salvi, S. 2021). The work stated above does not account for resource wastage and active servers, which automatically reduce energy consumption and enhance resource utilization. Further, the technique didn't discuss its implementation in a real working environment. An integrated methodology of Artificial Bee Colony with Chicken Swarm optimization has been implemented to allocate VMs with lesser power consumption, load, and migration cost (Pushpa, R., & Siddappa, M., 2021). Another swarm-based technique is being identified to minimize power consumption and resource wastage (Saravanan, D., et al., 2021). It only considered the two performance indicators of allocation, didn't give any idea about CPU and RAM utilization and didn't observe the migration aspect of VMs while allocating. Another approach for allocation has been proposed using Monarch Butterfly optimization, which worked on the packaging efficiency and simultaneously reduced the number of active servers (Ghetas, M. 2021). This algorithm does not account for reducing the wastage of resources and VMs migration. The algorithms discussed above suffered from various limitations, and thus it motivates the author to analyze the issue and devise a technique that could resolve the issue efficiently.

The Whale Optimization Algorithm (WOA) has recently piqued the interest of several scholars. Being an algorithm motivated by nature, it is ingrained in the humpback hunting procedure of whales. In terms of function, it's similar to genetic algorithms, and it has many benefits: faster convergence, fewer parameters, and simplified operation. It provides an efficient result with lower population size and iterations, reducing the algorithm's time complexity and contributing to the reason for adopting this algorithm. Furthermore, WOA can't be directly applied to discrete problems, so a binary version of WOA is required to resolve the VMA issue.

3. PROPOSED RESOURCE-AWARE VMA BASED ON BWOA

3.1 An Introduction to BWOA

Being a swarm-based optimization algorithm, the WOA is driven by the humpback whales hunting behavior, origina

lly brought forward by Seyedali Mirjalili and Andrew Lewis (Mirjalili, S., & Lewis, A. 2016). It includes three operators which simulate prey's search, encircling prey, and bubble-net foraging behavior shown by humpback whales. Contrasting to other meta-heuristic functions, WOA is extremely competitive and far superior to conventional techniques.

Standard WOA is appropriate for solving continuous space problems. However, the VMA issue belongs to the class of discrete optimization problems, so it's unfeasible to apply this algorithm directly to VMA. A modified genre of WOA is the paramount requirement for being applied to problems with discrete space. An advanced binary genre of WOA was proposed (Hussien, A. G., et al., 2020) to resolve discrete optimization problems.

BWOA is parallel to WOA, except that the whale's position is hosted in binary code. A transfer method is added in BWOA, which maps the distance vector to probability which further aids in the position update of the agent. The target is anticipated to be the recent best solution furnished through the algorithm. Then the best search agent is defined, and other agents attempt to modify the position as:

$$\vec{P}(k+1) = \vec{P}^*(k) - \vec{B} \times \vec{G} \quad (1)$$

$$\vec{G} = \left| H \times \vec{P}^*(k) - \vec{P}(k) \right| \quad (2)$$

Here, k indicates the recent iteration, B and H represent the coefficient vectors, P^* denotes the best alternative found so far, P stands for the position vector, and $||$ finds absolute value and $\times$ is element by element multiplication. Vector B and H are evaluated as

$$B_{\rightarrow} = 2l_{\rightarrow} \times \lambda_{\rightarrow} - l_{\rightarrow} \quad (3)$$

$$H_{\rightarrow} = 2 \times \lambda_{\rightarrow} \quad (4)$$

Where l gets decelerated from 2 to 0 over the entire duration of looping and λ is a random vector in $[0,1]$.

The mathematical simulation of the exploitation phase is as follows:

1. *Shrinking Encircling*: It is incurred by decreasing l values which has a random value in $[-l, l]$ where l is decremented from 2 to 0 during the course of iteration.

Spiral Updating: It computes the distance across the prey and the agent, which is expressed as:

$$\vec{P}(k+1) = \vec{X}' \times e^{\mu\beta} \times \cos(2\pi\beta) + \vec{P}^*(k) \quad (5)$$

$$\vec{X}' = \left| \vec{P}^*(k) - \vec{P}(k) \right| \quad (6)$$

Where, β takes random values in the range $[-1,1]$ and μ has fixed value.

Search for prey (Exploration phase):

$$\vec{E} = \left| H \times \vec{P}_{rand} - \vec{P} \right| \quad (7)$$

$$\vec{P}(k+1) = \vec{P}_{rand}(k) - \vec{B} \times \vec{E} \quad (8)$$

Where, $\vec{P}_{rand}$ displays a random position vector considered from the population in the current iteration.

The sigmoid function is adopted to ensure the probability for every bit lies in $[0,1]$, the sigmoid function adopted is:

$$F(x_i^j(k)) = \frac{1}{1 + e^{-S_i^j(k)}} \quad (9)$$

Here, $x_i^j(k)$ stands for the distance across the agent, and α is the random real number, hence the position for the agent is given by:

$$x_i^j(k+1) = \begin{cases} 1 & \text{if } \alpha < F(x_i^j(k)) \\ 0 & \text{otherwise} \end{cases} \quad (10)$$

This is the transition of $x_i^j(k+1)$.

3.2 Resource Wastage Model

Various VMA algorithms seek to maximize the efficient exploitation of available resources by striving for the least volume of resource waste on the servers. Resource wastage is highly dependent on memory and CPU associated with the server, which is being formulated as follows:

$$R_p^{wastage} = \frac{|NF_p^{cpu} - NF_p^{mem}| + \eta}{U_p^{cpu} + U_p^{mem}} \quad (11)$$

Where $R_p^{wastage}$ denotes resource wastage, NF_p^{mem} , NF_p^{CPU} shows the remaining memory and CPU related to server p , in normalized form, U_p^{mem} (table 1) and U_p^{CPU} denotes memory and CPU usage by the server p , in the normalized form and η has a value 0.00001 which acquires generally a very small positive real number. Thus, the total-sum resource wastage can be identified as:

$$f(X) = \sum_{j=1}^m R_j^{wastage} = \frac{\sum_{j=1}^m y_j \times \left[\left(T_j^{cpu} - \sum_{i=1}^n (x_{ij} \times \alpha c_i) \right) - \left(T_j^{mem} - \sum_{i=1}^n (x_{ij} \times \alpha m_i) \right) \right] + \eta}{\sum_{i=1}^n (x_{ij} \times \alpha c_i) + \sum_{i=1}^n (x_{ij} \times \alpha m_i)} \quad (12)$$

3.3 Problem Formulation

The heart of the cloud is Virtualization. When the data center gets the query from the user, a VM has to create and hosted on a server in relation to the given CPU, OS, and memory request. Assigning the VM to a different server is governed by the allocation or placement strategy, which is eminently accountable for mitigating resource wastage. This paper tracks VMA for minimalizing resource wastage in association with the servers' volume which is active. It is presumed that there exist n VMs characterized through the set $V = \{1, 2, 3, \dots, n\}$ and m servers or servers characterized through the

Table 1. Relinquishes the description for the notations being used in problem formulation along with the proposed method

Notations	Descriptions
V	VMs List and $ V = n$
P	servers List and $ P = m$
U_p^{cpu}	CPU utilization by server p in normalized form
U_p^{mem}	Memory utilization by server p in normalized form
T_j^{cpu}	Total CPU utilization by server j
T_j^{mem}	Total memory utilization of server j
TU_j^{cpu}	Total CPU utilization by an active server j
TU_j^{mem}	Total memory utilization by an active server j
NF_p^{cpu}	Residual CPU of server p in normalized form
NF_p^{mem}	Residual memory of server p in normalized form
$R_{cpu}^{wastage}$	Resource wastage of server p
$\alpha m_i, \alpha c_i$	memory and CPU requirement for VM i
Tmj, Tcj	Thresholds for memory and CPU utilization corresponding to jth physical server
x_{ij}	A variable that defines whether or not the ith virtual machine is assigned to the jth physical server.
y_j	This variable shows whether or not the jth physical server is in use.

set $P=\{1,2,3,\dots,m\}$. Let αm_i and αc_i represent memory and CPU requirement by all generated VMs $i \in V$ and Tm_j and Tc_j represent memory and CPU of all servers $j \in P$. Let x_{ij} and y_j be the binary variables which are being defined as:

$$x_{ij} = \begin{cases} 1 & \text{if VM } i \forall i \in V \text{ is allocated to server } j \forall j \in P \\ 0 & \text{otherwise} \end{cases} \quad (13)$$

$$y_j = \begin{cases} 1 & \text{if the server } j \forall j \in P \text{ is in use} \\ 0 & \text{otherwise} \end{cases} \quad (14)$$

The main motive of this works is to

$$\text{Minimize } f(X) = \sum_{j=1}^m R_{cpu}^{wastage} \quad (15)$$

Subject to following:

$$\sum_{j=1}^m x_{ij} = 1 \quad \forall i \in V \quad (16)$$

$$\sum_{i=1}^n x_{ij} \times \alpha c_i \leq Tc_j \times y_j \quad \forall j \in P \quad (17)$$

$$\sum_{i=1}^n x_{ij} \times \alpha m_i \leq Tm_j \times y_j \quad \forall j \in P \quad (18)$$

The function in equation 15 is expressed to reduce resource waste. Equation 16 restricts VM i is to be deployed to one server only. The restriction in equations 17& 18 ensures that the allocated VM does not drive beyond the extent of the server.

3.4 Solution Representation

In BWOA each solution is represented by a decision variable x_{ij} which depicts the allotment of VM j on the server i . Each solution or whale in the VMA issue is represented by a binary matrix as represented through figure 3 with the condition as every VM should have been allotted to one and only one server.

3.5 Modification in BWOA for the VMA

A few modifications are required for applying the BWOA to VM allocation problems, which definitely needs to be performed.

1. **Initialization:** Initially, the whale population is initialized randomly within the binary search space so it's not necessary that the generated population will satisfy the constriction of Equations 16& 17. An initialization function is obligatory to initialize the desired population. At the commencement of the procedure, every VM is allocated to the server with the probability $1 / m$. To allot the VM i to the server or server j a random array of $1 \times n$ which contains the value between $[1, m]$ with no duplicates. The x_{ij} is set up as 1 in accordance with the random array, and if it satisfies the constraint, then it is allotted to the first server, otherwise to the second, and the process goes on.

Figure 3. Solution Representation

$$\begin{bmatrix} 0 & 1 & 0 & \dots\dots\dots & 0 & 1 \\ 1 & 1 & 0 & \dots\dots\dots & 0 & 1 \\ \vdots & & & & & \\ \vdots & & & & & \\ \vdots & & & & & \\ 1 & 1 & 0 & \dots\dots\dots & 1 & 0 & 0 & 1 \end{bmatrix}$$

2. *Position update of the whales:* Each whale updates its position with respect to Equations 1, 5, and 8 in accord with the given random probability. After the update, the value for each position associated to the agent is squashed using equation 18. In all these updates, it might be possible the solution becomes infeasible by violating equation 16. To keep track of VMs, a *vm_status* array is implemented to restrict the update of the elements in column *jth* column in x_{ij} which obtained the value 1 in the calculated solution so far. At the start of every iteration, the *vm_status* array is initialized by 0 and before the updation of x_{ij} it is checked whether *vm_status*[*j*] has any value. If some value exists with it, then it demonstrates the VM *j* has already been allotted to the server from x_{ij} to x_{i-1j} . Then, x_{ij} is not updated, and it is only updated when *vm_status*[*j*] = 0 and $\alpha < S(x_i^j(k))$ then x_{ij} is set up as 1. The *vm_status* array has substantially increased the algorithm's effectiveness.
3. *Function Mapping:* The redeployment of VMs does not necessarily change the status of the server. The original BWOA is unexpected to converge to a globally optimal solution. If it does, the whales' position will be zero, increasing the likelihood of modifications in the associated bit. The search will be more random and lacking in direction. The original mapping does not look out for the VM allocation because if the whale's position leads to VM reallocation, then it is not necessary that the status of the server also changes. If server *X* is hosting VM *a, b, c*, and server *Y* is hosting VM *d* then both have status 1. If VM *c* is migrated from server *X* to *Y* then also their status will remain the same. Thus, to confirm the allocation problem, the sigmoid function needs updating, which is as follows:

$$F(x_i^j(k)) = \begin{cases} 1 - \frac{2}{1 + e^{-S_i^j(k)}} & S_i^j(k) \leq 0 \\ \frac{2}{1 + e^{-S_i^j(k)}} - 1 & S_i^j(k) > 0 \end{cases} \quad (19)$$

4. *Redeployment of VM:* With the above information, one can determine the distance of whale in each generation and evaluate the modifications required in VM reallocation. When a bit of whale displacement transitions to 1, the related server remains unchanged; nevertheless, if a bit of whale displacement transitions to 0, the corresponding server requires some adjustment. A new list of VMs is being generated, which needs deployment on the server. In this work, BFD (Best Fit Decrease) is employed for the reallocation of VMs. This scheme helps in reducing the active server's volume to a great extent. Then the fitness is determined again, and the global best solution is obtained.

4. ALGORITHM

4.1 Data Structure

- **Whale[wnum]:** It specifies the total-sum size of whale's population present.
- **Iteration:** It defines the total iteration in BWOA.
- **Server_List:** It's the servers' listing that holds information about each one, such as its ID, CPU, and disc, including the volume of VMs assigned to it.
- **Vm_List:** It provides all the information associated with a VM.
- **RandWhale:** It stores the position linked with the given whale.
- **Wvalue[wnum]:** It provides the optimum resource wastage of the current whale in the given iteration.
- **G_best:** It stores the effective position of a whale in the population in the given iteration.
- **G_value:** It gives the minimum resource wastage fitness for the whale population in the given iteration.

The pseudocode for the suggested technique is displayed by algorithm 1 in figure 4. It depicts the overall procedure followed in allocation. The first step is to generate wnum number of whales, each of which has the same server list and VM list but each has its own deployment of VMs as their allocation on servers is random. Then among the listing given for various allocations, the global best allocation is determined through the given fitness function providing the minimum resource wastage. This provides initial preparation for the implementation of BWOA (lines 2-10). Next, the BWOA is utilized to update each whale's position to procure the optimal global allocation of VMs (lines 11-21). After that, the global optimal allocation of the whale is the desired VM allocation (line 22).

Algorithm 2 in figure 5 is responsible for the initial allocation of VMs on servers; first, a random array vmRand is generated which has a random value from 1 to m. Each VM selects the server with the probability $1/m$, a server is allocated to a VM only if it has enough resources to accommodate it. The server's displacement is set to 1 (from 0) if it does not already host any other VMs, indicating that the server has already been assigned to the VM. If the current server is insufficiently resourced, the algorithm will move on to the next server until it finds one.

Algorithm 3 in figure 6 depicts the process of upgrading each whale position during each iteration. It updates each bit of whale, which represents each server (lines 2-12). Then it applies the transition function on each server. The transition value obtained decides whether or not the provided server requires any adjustments. If the transition value is 0, the server is added to the list of servers that require VM redeployment (lines 13-18). The list of VMs is procured from the list and allocated with the BFD strategy on the servers. The information is updated in the serverList (lines 19-22).

5. SIMULATION

The proposed procedure's effectiveness has been assessed through simulation experiments. The simulation's output is contrasted with other algorithms. The simulation for the framework has been enforced on Cloudsim using the Eclipse platform and runs on a PC Intel Core CPU i3-7020U and 4.0 GB Ram and Windows 10 Home as the operating system. Each experiment is performed with a similar VMs number and servers. The start-up population of whales is initialized by 50, and maximum iterations are initialized to 100. To ensure the methodology's applicability, the proposed server configuration has been kept equivalent to the one available in the market. The approach is analyzed against the Genetic Algorithm (GA), First Fit Algorithm (FFA), and Best Fit Decrease algorithm (BFD). FFA is a VM allocation bin packing method in which a VM is designated to the first available server with all of the needed resources. BFD is the modified form of the Best First algorithm, which first arranges VMs in the decreasing order of their requisitioned resources, and then it is allocated to the most fitted available

Figure 4. Algorithm 1 for VMA according to BWOA

```

Algorithm-1

Input: vm_List, wnum, Iteration, server_List
Output: G_best
(1) G_best = Null, G_value =  $\infty$ 
(2) for j in wnum
(3)   server_List = RandomVMAllocation (vm_List)
(4)   Whale[j] = server_List
(5)   Wvalue[j] = fitness (server_List)
(6)   if (G_value > Wvalue[j]) then
(7)     G_best = Whale[j]
(8)     G_value = Wvalue[j]
(9)   end if
(10) end for
(11) while Iteration > 0
(12) for j in wnum
(13)   Whale[j] = BWOA (server_List, Whale[j])
(14)   Wvalue[j] = fitness (Whale[j])
(15)   if (G_value > Wvalue[j]) then
(16)     G_best = Whale[j]
(17)     G_value = Wvalue[j]
(18)   end if
(19) end for
(20)   Iteration = Iteration - 1
(21) end while
(22) return G_best

```

Figure 5. Algorithm 2 for Random initial allocation of VM

```

Algorithm-2

Input : vm_List
Output: server_List
(1) x=zeros(m,n)
(2) vmRand=random(n,1)//Generates a nx1 vector with all values distributed randomly between [1,m]
(3) for server i in server_List
(4)   for vm j in vm_List
(5)     if(server i has sufficient resources to deploy vm j) then
(6)       server.add(vm)
(7)       if(server is empty)
(8)         server.setDisplacement(1)
(9)       end if
(10)      x[vmRand[i],i]=1 // Allot vm j to server i
(11)     end if
(12)   end for
(13) end for

```

server with all the resources needed. Both algorithms follow the greedy approach to furnish the solution. GA is the metaheuristic population-based approach working closely with the concept of survival of the fittest and reproduction. To develop the solution, it selects the best individual from a given population and performs crossover and mutation. The execution of the algorithms on the simulator is represented in figure 7, figure 8, figure 9, and figure 10. The inference of the result is as follows:

5.1 Resource Wastage fitness

The proposed method is compared with the three available algorithms in correspondence with resource wastage. Figure 11 shows the resource wastage achieved from the given four algorithms under the condition that different numbers of VMs and servers are deployed. Table 2 represents the statistical data of the same. The iteration parameter is initiated to 40, and the BWOA population size is also set to 50. The outcome accomplishes that the anticipated approach can get the minimal resource

Figure 6. Algorithm 3 for BWOA

Algorithm-3

Input: vm_List, server_List, Whale
Output: server_List

- (1) Update B , $\bar{H}$, p and l
- (2) **if** ($p < 0.5$)
- (3) **if** ($|B| < 1$)
- (4) Whale's position is updated according to Eq. 1
- (5) **else if** ($|B| \geq 1$)
- (6) RandWhale, A random whale chosen from the list
- (7) Whale's position is updated according to Eq. 8
- (8) **end if**
- (9) **end if**
- (10) **else if** ($p \geq 0.5$)
- (11) Whale's position is updated according to Eq. 5
- (12) **end if**
- (13) Update the whale according to Eq. 18
- (14) **for** server in ServerList
- (15) $S = \text{Whale.getServer}(\text{server.getId()}).\text{getDisplacement}()$
- (16) $S' = S.\text{getTransition}()$
- (17) **if** ($S = 0$) **then**
- (18) ServerListRedeploy.add(server)
- (19) **endif**
- (20) **end for**
- (21) MigratedVmlist = RedeployServerList.getAllVms()
- (22) Reallocate VMs from MigratedVmlist to the servers in BFD order.
- (23) Update the vm_status in server_List according to the latest reallocation.
- (24) **return** server_List

Figure 7. FF representation

```

VMPWOA.java  VMAllocateFF.java
46  if (this.allocation){
47
48      if (vmList.isEmpty())
49          return Collections.EMPTY_MAP;
50
51      for (Vm vm : vmList) {
52          this.deallocateHostForVm(vm);
53      }
54
55      normalizeHostdata(); // using z score
56      normalizeVmdata(vmList);
57
58
59
60      List<Host> hostlist = this.getHostList();
61
62      for (Vm vm : vmList) {
63          for (Host h : hostlist) {
64              if (h.isSuitableForVm(vm)) {
65                  h.createTemporaryVm(vm);
66              }
67          }
68      }
69
70      // ... (rest of the code)
71
72      // ... (rest of the code)
73
74      // ... (rest of the code)
75
76      // ... (rest of the code)
77
78      // ... (rest of the code)
79
80      // ... (rest of the code)
81
82      // ... (rest of the code)
83
84      // ... (rest of the code)
85
86      // ... (rest of the code)
87
88      // ... (rest of the code)
89
90      // ... (rest of the code)
91
92      // ... (rest of the code)
93
94      // ... (rest of the code)
95
96      // ... (rest of the code)
97
98      // ... (rest of the code)
99
100     // ... (rest of the code)
101
102     // ... (rest of the code)
103
104     // ... (rest of the code)
105
106     // ... (rest of the code)
107
108     // ... (rest of the code)
109
110     // ... (rest of the code)
111
112     // ... (rest of the code)
113
114     // ... (rest of the code)
115
116     // ... (rest of the code)
117
118     // ... (rest of the code)
119
120     // ... (rest of the code)
121
122     // ... (rest of the code)
123
124     // ... (rest of the code)
125
126     // ... (rest of the code)
127
128     // ... (rest of the code)
129
130     // ... (rest of the code)
131
132     // ... (rest of the code)
133
134     // ... (rest of the code)
135
136     // ... (rest of the code)
137
138     // ... (rest of the code)
139
140     // ... (rest of the code)
141
142     // ... (rest of the code)
143
144     // ... (rest of the code)
145
146     // ... (rest of the code)
147
148     // ... (rest of the code)
149
150     // ... (rest of the code)
151
152     // ... (rest of the code)
153
154     // ... (rest of the code)
155
156     // ... (rest of the code)
157
158     // ... (rest of the code)
159
160     // ... (rest of the code)
161
162     // ... (rest of the code)
163
164     // ... (rest of the code)
165
166     // ... (rest of the code)
167
168     // ... (rest of the code)
169
170     // ... (rest of the code)
171
172     // ... (rest of the code)
173
174     // ... (rest of the code)
175
176     // ... (rest of the code)
177
178     // ... (rest of the code)
179
180     // ... (rest of the code)
181
182     // ... (rest of the code)
183
184     // ... (rest of the code)
185
186     // ... (rest of the code)
187
188     // ... (rest of the code)
189
190     // ... (rest of the code)
191
192     // ... (rest of the code)
193
194     // ... (rest of the code)
195
196     // ... (rest of the code)
197
198     // ... (rest of the code)
199
200     // ... (rest of the code)
201
202     // ... (rest of the code)
203
204     // ... (rest of the code)
205
206     // ... (rest of the code)
207
208     // ... (rest of the code)
209
210     // ... (rest of the code)
211
212     // ... (rest of the code)
213
214     // ... (rest of the code)
215
216     // ... (rest of the code)
217
218     // ... (rest of the code)
219
220     // ... (rest of the code)
221
222     // ... (rest of the code)
223
224     // ... (rest of the code)
225
226     // ... (rest of the code)
227
228     // ... (rest of the code)
229
230     // ... (rest of the code)
231
232     // ... (rest of the code)
233
234     // ... (rest of the code)
235
236     // ... (rest of the code)
237
238     // ... (rest of the code)
239
240     // ... (rest of the code)
241
242     // ... (rest of the code)
243
244     // ... (rest of the code)
245
246     // ... (rest of the code)
247
248     // ... (rest of the code)
249
250     // ... (rest of the code)
251
252     // ... (rest of the code)
253
254     // ... (rest of the code)
255
256     // ... (rest of the code)
257
258     // ... (rest of the code)
259
260     // ... (rest of the code)
261
262     // ... (rest of the code)
263
264     // ... (rest of the code)
265
266     // ... (rest of the code)
267
268     // ... (rest of the code)
269
270     // ... (rest of the code)
271
272     // ... (rest of the code)
273
274     // ... (rest of the code)
275
276     // ... (rest of the code)
277
278     // ... (rest of the code)
279
280     // ... (rest of the code)
281
282     // ... (rest of the code)
283
284     // ... (rest of the code)
285
286     // ... (rest of the code)
287
288     // ... (rest of the code)
289
290     // ... (rest of the code)
291
292     // ... (rest of the code)
293
294     // ... (rest of the code)
295
296     // ... (rest of the code)
297
298     // ... (rest of the code)
299
300     // ... (rest of the code)
301
302     // ... (rest of the code)
303
304     // ... (rest of the code)
305
306     // ... (rest of the code)
307
308     // ... (rest of the code)
309
310     // ... (rest of the code)
311
312     // ... (rest of the code)
313
314     // ... (rest of the code)
315
316     // ... (rest of the code)
317
318     // ... (rest of the code)
319
320     // ... (rest of the code)
321
322     // ... (rest of the code)
323
324     // ... (rest of the code)
325
326     // ... (rest of the code)
327
328     // ... (rest of the code)
329
330     // ... (rest of the code)
331
332     // ... (rest of the code)
333
334     // ... (rest of the code)
335
336     // ... (rest of the code)
337
338     // ... (rest of the code)
339
340     // ... (rest of the code)
341
342     // ... (rest of the code)
343
344     // ... (rest of the code)
345
346     // ... (rest of the code)
347
348     // ... (rest of the code)
349
350     // ... (rest of the code)
351
352     // ... (rest of the code)
353
354     // ... (rest of the code)
355
356     // ... (rest of the code)
357
358     // ... (rest of the code)
359
360     // ... (rest of the code)
361
362     // ... (rest of the code)
363
364     // ... (rest of the code)
365
366     // ... (rest of the code)
367
368     // ... (rest of the code)
369
370     // ... (rest of the code)
371
372     // ... (rest of the code)
373
374     // ... (rest of the code)
375
376     // ... (rest of the code)
377
378     // ... (rest of the code)
379
380     // ... (rest of the code)
381
382     // ... (rest of the code)
383
384     // ... (rest of the code)
385
386     // ... (rest of the code)
387
388     // ... (rest of the code)
389
390     // ... (rest of the code)
391
392     // ... (rest of the code)
393
394     // ... (rest of the code)
395
396     // ... (rest of the code)
397
398     // ... (rest of the code)
399
400     // ... (rest of the code)
401
402     // ... (rest of the code)
403
404     // ... (rest of the code)
405
406     // ... (rest of the code)
407
408     // ... (rest of the code)
409
410     // ... (rest of the code)
411
412     // ... (rest of the code)
413
414     // ... (rest of the code)
415
416     // ... (rest of the code)
417
418     // ... (rest of the code)
419
420     // ... (rest of the code)
421
422     // ... (rest of the code)
423
424     // ... (rest of the code)
425
426     // ... (rest of the code)
427
428     // ... (rest of the code)
429
430     // ... (rest of the code)
431
432     // ... (rest of the code)
433
434     // ... (rest of the code)
435
436     // ... (rest of the code)
437
438     // ... (rest of the code)
439
440     // ... (rest of the code)
441
442     // ... (rest of the code)
443
444     // ... (rest of the code)
445
446     // ... (rest of the code)
447
448     // ... (rest of the code)
449
450     // ... (rest of the code)
451
452     // ... (rest of the code)
453
454     // ... (rest of the code)
455
456     // ... (rest of the code)
457
458     // ... (rest of the code)
459
460     // ... (rest of the code)
461
462     // ... (rest of the code)
463
464     // ... (rest of the code)
465
466     // ... (rest of the code)
467
468     // ... (rest of the code)
469
470     // ... (rest of the code)
471
472     // ... (rest of the code)
473
474     // ... (rest of the code)
475
476     // ... (rest of the code)
477
478     // ... (rest of the code)
479
480     // ... (rest of the code)
481
482     // ... (rest of the code)
483
484     // ... (rest of the code)
485
486     // ... (rest of the code)
487
488     // ... (rest of the code)
489
490     // ... (rest of the code)
491
492     // ... (rest of the code)
493
494     // ... (rest of the code)
495
496     // ... (rest of the code)
497
498     // ... (rest of the code)
499
500     // ... (rest of the code)
501
502     // ... (rest of the code)
503
504     // ... (rest of the code)
505
506     // ... (rest of the code)
507
508     // ... (rest of the code)
509
510     // ... (rest of the code)
511
512     // ... (rest of the code)
513
514     // ... (rest of the code)
515
516     // ... (rest of the code)
517
518     // ... (rest of the code)
519
520     // ... (rest of the code)
521
522     // ... (rest of the code)
523
524     // ... (rest of the code)
525
526     // ... (rest of the code)
527
528     // ... (rest of the code)
529
530     // ... (rest of the code)
531
532     // ... (rest of the code)
533
534     // ... (rest of the code)
535
536     // ... (rest of the code)
537
538     // ... (rest of the code)
539
540     // ... (rest of the code)
541
542     // ... (rest of the code)
543
544     // ... (rest of the code)
545
546     // ... (rest of the code)
547
548     // ... (rest of the code)
549
550     // ... (rest of the code)
551
552     // ... (rest of the code)
553
554     // ... (rest of the code)
555
556     // ... (rest of the code)
557
558     // ... (rest of the code)
559
560     // ... (rest of the code)
561
562     // ... (rest of the code)
563
564     // ... (rest of the code)
565
566     // ... (rest of the code)
567
568     // ... (rest of the code)
569
570     // ... (rest of the code)
571
572     // ... (rest of the code)
573
574     // ... (rest of the code)
575
576     // ... (rest of the code)
577
578     // ... (rest of the code)
579
580     // ... (rest of the code)
581
582     // ... (rest of the code)
583
584     // ... (rest of the code)
585
586     // ... (rest of the code)
587
588     // ... (rest of the code)
589
590     // ... (rest of the code)
591
592     // ... (rest of the code)
593
594     // ... (rest of the code)
595
596     // ... (rest of the code)
597
598     // ... (rest of the code)
599
600     // ... (rest of the code)
601
602     // ... (rest of the code)
603
604     // ... (rest of the code)
605
606     // ... (rest of the code)
607
608     // ... (rest of the code)
609
610     // ... (rest of the code)
611
612     // ... (rest of the code)
613
614     // ... (rest of the code)
615
616     // ... (rest of the code)
617
618     // ... (rest of the code)
619
620     // ... (rest of the code)
621
622     // ... (rest of the code)
623
624     // ... (rest of the code)
625
626     // ... (rest of the code)
627
628     // ... (rest of the code)
629
630     // ... (rest of the code)
631
632     // ... (rest of the code)
633
634     // ... (rest of the code)
635
636     // ... (rest of the code)
637
638     // ... (rest of the code)
639
640     // ... (rest of the code)
641
642     // ... (rest of the code)
643
644     // ... (rest of the code)
645
646     // ... (rest of the code)
647
648     // ... (rest of the code)
649
650     // ... (rest of the code)
651
652     // ... (rest of the code)
653
654     // ... (rest of the code)
655
656     // ... (rest of the code)
657
658     // ... (rest of the code)
659
660     // ... (rest of the code)
661
662     // ... (rest of the code)
663
664     // ... (rest of the code)
665
666     // ... (rest of the code)
667
668     // ... (rest of the code)
669
670     // ... (rest of the code)
671
672     // ... (rest of the code)
673
674     // ... (rest of the code)
675
676     // ... (rest of the code)
677
678     // ... (rest of the code)
679
680     // ... (rest of the code)
681
682     // ... (rest of the code)
683
684     // ... (rest of the code)
685
686     // ... (rest of the code)
687
688     // ... (rest of the code)
689
690     // ... (rest of the code)
691
692     // ... (rest of the code)
693
694     // ... (rest of the code)
695
696     // ... (rest of the code)
697
698     // ... (rest of the code)
699
700     // ... (rest of the code)
701
702     // ... (rest of the code)
703
704     // ... (rest of the code)
705
706     // ... (rest of the code)
707
708     // ... (rest of the code)
709
710     // ... (rest of the code)
711
712     // ... (rest of the code)
713
714     // ... (rest of the code)
715
716     // ... (rest of the code)
717
718     // ... (rest of the code)
719
720     // ... (rest of the code)
721
722     // ... (rest of the code)
723
724     // ... (rest of the code)
725
726     // ... (rest of the code)
727
728     // ... (rest of the code)
729
730     // ... (rest of the code)
731
732     // ... (rest of the code)
733
734     // ... (rest of the code)
735
736     // ... (rest of the code)
737
738     // ... (rest of the code)
739
740     // ... (rest of the code)
741
742     // ... (rest of the code)
743
744     // ... (rest of the code)
745
746     // ... (rest of the code)
747
748     // ... (rest of the code)
749
750     // ... (rest of the code)
751
752     // ... (rest of the code)
753
754     // ... (rest of the code)
755
756     // ... (rest of the code)
757
758     // ... (rest of the code)
759
760     // ... (rest of the code)
761
762     // ... (rest of the code)
763
764     // ... (rest of the code)
765
766     // ... (rest of the code)
767
768     // ... (rest of the code)
769
770     // ... (rest of the code)
771
772     // ... (rest of the code)
773
774     // ... (rest of the code)
775
776     // ... (rest of the code)
777
778     // ... (rest of the code)
779
780     // ... (rest of the code)
781
782     // ... (rest of the code)
783
784     // ... (rest of the code)
785
786     // ... (rest of the code)
787
788     // ... (rest of the code)
789
790     // ... (rest of the code)
791
792     // ... (rest of the code)
793
794     // ... (rest of the code)
795
796     // ... (rest of the code)
797
798     // ... (rest of the code)
799
800     // ... (rest of the code)
801
802     // ... (rest of the code)
803
804     // ... (rest of the code)
805
806     // ... (rest of the code)
807
808     // ... (rest of the code)
809
810     // ... (rest of the code)
811
812     // ... (rest of the code)
813
814     // ... (rest of the code)
815
816     // ... (rest of the code)
817
818     // ... (rest of the code)
819
820     // ... (rest of the code)
821
822     // ... (rest of the code)
823
824     // ... (rest of the code)
825
826     // ... (rest of the code)
827
828     // ... (rest of the code)
829
830     // ... (rest of the code)
831
832     // ... (rest of the code)
833
834     // ... (rest of the code)
835
836     // ... (rest of the code)
837
838     // ... (rest of the code)
839
840     // ... (rest of the code)
841
842     // ... (rest of the code)
843
844     // ... (rest of the code)
845
846     // ... (rest of the code)
847
848     // ... (rest of the code)
849
850     // ... (rest of the code)
851
852     // ... (rest of the code)
853
854     // ... (rest of the code)
855
856     // ... (rest of the code)
857
858     // ... (rest of the code)
859
860     // ... (rest of the code)
861
862     // ... (rest of the code)
863
864     // ... (rest of the code)
865
866     // ... (rest of the code)
867
868     // ... (rest of the code)
869
870     // ... (rest of the code)
871
872     // ... (rest of the code)
873
874     // ... (rest of the code)
875
876     // ... (rest of the code)
877
878     // ... (rest of the code)
879
880     // ... (rest of the code)
881
882     // ... (rest of the code)
883
884     // ... (rest of the code)
885
886     // ... (rest of the code)
887
888     // ... (rest of the code)
889
890     // ... (rest of the code)
891
892     // ... (rest of the code)
893
894     // ... (rest of the code)
895
896     // ... (rest of the code)
897
898     // ... (rest of the code)
899
900     // ... (rest of the code)
901
902     // ... (rest of the code)
903
904     // ... (rest of the code)
905
906     // ... (rest of the code)
907
908     // ... (rest of the code)
909
910     // ... (rest of the code)
911
912     // ... (rest of the code)
913
914     // ... (rest of the code)
915
916     // ... (rest of the code)
917
918     // ... (rest of the code)
919
920     // ... (rest of the code)
921
922     // ... (rest of the code)
923
924     // ... (rest of the code)
925
926     // ... (rest of the code)
927
928     // ... (rest of the code)
929
930     // ... (rest of the code)
931
932     // ... (rest of the code)
933
934     // ... (rest of the code)
935
936     // ... (rest of the code)
937
938     // ... (rest of the code)
939
940     // ... (rest of the code)
941
942     // ... (rest of the code)
943
944     // ... (rest of the code)
945
946     // ... (rest of the code)
947
948     // ... (rest of the code)
949
950     // ... (rest of the code)
951
952     // ... (rest of the code)
953
954     // ... (rest of the code)
955
956     // ... (rest of the code)
957
958     // ... (rest of the code)
959
960     // ... (rest of the code)
961
962     // ... (rest of the code)
963
964     // ... (rest of the code)
965
966     // ... (rest of the code)
967
968     // ... (rest of the code)
969
970     // ... (rest of the code)
971
972     // ... (rest of the code)
973
974     // ... (rest of the code)
975
976     // ... (rest of the code)
977
978     // ... (rest of the code)
979
980     // ... (rest of the code)
981
982     // ... (rest of the code)
983
984     // ... (rest of the code)
985
986     // ... (rest of the code)
987
988     // ... (rest of the code)
989
990     // ... (rest of the code)
991
992     // ... (rest of the code)
993
994     // ... (rest of the code)
995
996     // ... (rest of the code)
997
998     // ... (rest of the code)
999
1000    // ... (rest of the code)

```

Console :

```

terminated> VMPWOA [Java Application] C:\Users\HP\p2\pool\plugins\org.eclipse.justi.openjdk.hotspot.jre.full.win32.x86_64_16.0.2.v20210721-1149\jre\bin\javaw.exe 03-Jan-2022, 1:18:50 pm - 1:18:57 pm
486 109 SUCCESS 2 27
487 106 SUCCESS 2 26
488 107 SUCCESS 2 26
489 104 SUCCESS 2 26
490 105 SUCCESS 2 26
491 102 SUCCESS 2 25
492 103 SUCCESS 2 25
493 116 SUCCESS 2 29
494 117 SUCCESS 2 29
495 114 SUCCESS 2 28
496 115 SUCCESS 2 28
497 112 SUCCESS 2 28
498 113 SUCCESS 2 28
499 110 SUCCESS 2 27
500 111 SUCCESS 2 27
The global resource wastage fitness is=41.67375984126794
The number of active host is= 256
The CPU Utilization is= 67.20%
The memory Utilization is= 70.29%

```

wastage against the three other algorithms in all the servers and VMs. It can be observed from the figure the FF algorithm shows the maximum resource wastage, and as the VM is scaled up to 500, the wastage is maximum. The BFD algorithm has performed better than FF but has as compared to GA and BWOA its efficiency is less. It can be seen from the table the BWOA has reduced the resource wastage by 62% approximately as compared to FF and this reduction is 36% as against BFD. The GA has also shown a comparable result to BWOA; unfortunately, its efficiency is less than BWOA. The control of the resource wastage is mainly because the adaptive variation of the search agent has led to a smooth transition between the exploration and exploitation phase. Moreover, the approach proposed globally optimizes, ensuring the minimum resource wastage. Whereas the FFA and BFD

Figure 8. BFD representation

```

110  @Override CPU, mem, vmlcpu, vmware;
111  boolean isEmpty = false;
112
113
114
115  Map<Host, Set<Vm>> mapping= new HashMap();
116  Set<Vm> vmset ;
117
118  for(Map.Entry entry: allocatedvm.entrySet()){
119      Host h = (Host)entry.getValue();
120      Vm vm = (Vm)entry.getKey();
121      vmset = new HashSet();
122      vmset = mapping.get(h);
123      if(vmset!=null){
124          { vmset.add(vm); }
125      }
126      else {
127          { vmset = new HashSet(); }
128          mapping.put(h, vmset);
129      }
130  }
131
132
133
134
135
136

```

```

-terminated- VMPWOA [Java Application] C:\Users\HP\p2\pool\plugins\org.eclipse.justi.openjdk hotspot\re.full.win32.x86_64_16.0.2.v20210721-1149\jre\bin\javaw.exe (03-Jan-2022, 1:17:07 pm - 1:17:14 pm)
486  109  SUCCESS  2  234
487  106  SUCCESS  2  231
488  107  SUCCESS  2  232
489  104  SUCCESS  2  229
490  105  SUCCESS  2  230
491  102  SUCCESS  2  227
492  103  SUCCESS  2  228
493  116  SUCCESS  2  241
494  117  SUCCESS  2  242
495  114  SUCCESS  2  239
496  115  SUCCESS  2  240
497  112  SUCCESS  2  237
498  113  SUCCESS  2  238
499  110  SUCCESS  2  235
500  111  SUCCESS  2  236
The global resource wastage fitness is=24.53458972178363
The number of active host is= 229
The CPU Utilization is= 81.23%
The memory Utilization is= 85.99%

```

Figure 9. GA representation

```

487  List<Host> sortedhostlist = arrange(1);
488  for(int migratevm : migratelist ) {
489      Vm vm = vmListById.get(migratevm);
490      vmallocated = false;
491      for(int l =0 ; l< hostlistsize ; l++ ) {
492          if(hoststatusList[l]==0 || tempprevioushoststatus[l][1]==0)
493          {
494              inactivehost.add(l);
495          }
496          else {
497              Host h = hostListById.get(l);
498              if(issuitable(h,vm,i) && !vmallocated)
499              {
500                  vmstatusList[migratevm]=l;
501                  vmallocated=true;
502                  updatehostconfig(h, vm,i,false);
503              }
504          }
505      }
506  }

```

```

-terminated- VMPWOA [Java Application] C:\Users\HP\p2\pool\plugins\org.eclipse.justi.openjdk hotspot\re.full.win32.x86_64_16.0.2.v20210721-1149\jre\bin\javaw.exe (03-Jan-2022, 1:14:28 pm - 1:14:40 pm)
487  487  SUCCESS  2  390
488  488  SUCCESS  2  139
489  489  SUCCESS  2  448
490  490  SUCCESS  2  319
491  491  SUCCESS  2  435
492  492  SUCCESS  2  62
493  493  SUCCESS  2  32
494  494  SUCCESS  2  297
495  495  SUCCESS  2  330
496  496  SUCCESS  2  158
497  497  SUCCESS  2  19
498  498  SUCCESS  2  54
499  499  SUCCESS  2  371
500  500  SUCCESS  2  207
The global resource wastage fitness is=17.80127854695384
The number of active host is= 221
The CPU Utilization is= 86.20%
The memory Utilization is= 86.68%

```

heuristic algorithms perform the operation in a greedy manner that lacks such global optimization, leading to greater resource wastage.

5.2 Number of Active Servers

Besides resource utilization, the other major parameter which defines the efficaciousness of VMA is mostly the volume of working servers. The VMA has shown more effectiveness while there are lesser active servers. The system is configured in a similar way as before, only the iteration parameter is 50,

Figure 10. BWOA representation

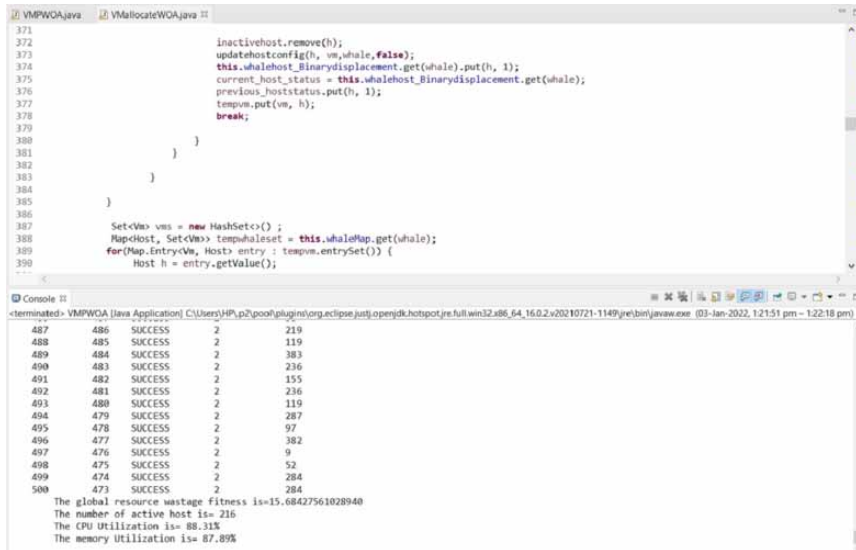


Table 2. Comparison of Resource Wastage Fitness

	Number of Virtual Machines				
Techniques	100	200	300	400	500
FF	11.0015	16.1356	25.1247	34.4372	41.6737
BFD	6.6825	9.8961	14.3863	19.8849	24.5345
GA	4.6637	7.4562	12.6731	14.9631	17.8013
BWOA	3.3410	5.0827	9.8816	12.4574	15.6842

and the population is 50. As can be observed from figure 12, the suggested methodology based on BWOA creates lesser active servers than other algorithms. It can be inferred from the figure in the case of 100VMs that all the methods have created a comparable number of active servers, and as the number of VMs increases, FF is consuming more servers compared to the other approaches. BFD and GA have shown quite comparable consumption of servers. Table 3 shows that when the number of VMs is less, BWOA creates 19% less active servers compared to FF, and when the number of VMs is increased to 500, it creates 15% less active servers. In BWOA, 216 active servers are utilized for allocating the VMs, which is the least compared to BFD and GA, which used 229 and 221 servers, respectively. It is also inferred that as VMs grow in number, so does the algorithm's performance.

5.3 Memory & CPU Utilization

Memory and CPU utilization is the memory or CPU utilized by the VMs of the working server. It is given as follows:

$$TU^{CPU} = \frac{1}{\eta} \sum_{\lambda=1}^{\eta} \left[\frac{\sum_{i=1}^n (x_{i\lambda} \cdot \alpha c_i)}{U_p^{cpu}} \right] \quad (20)$$

Figure 11. Resource Wastage Fitness

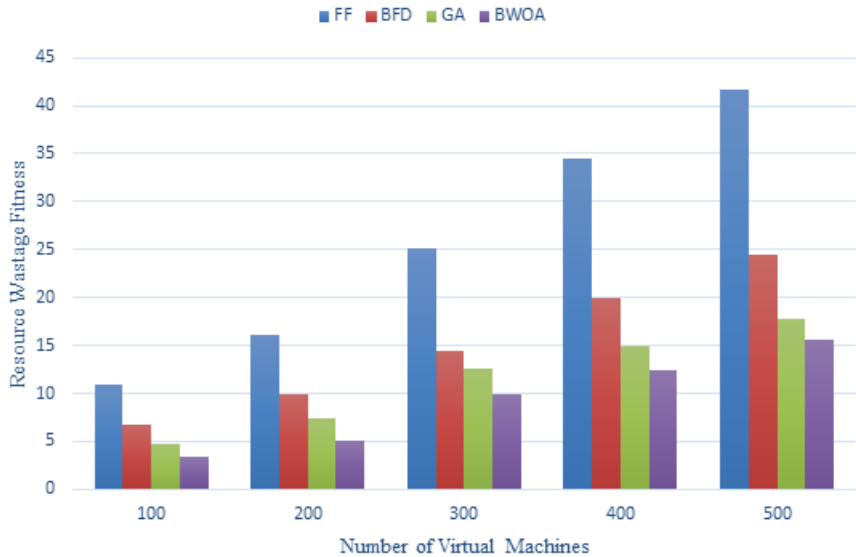
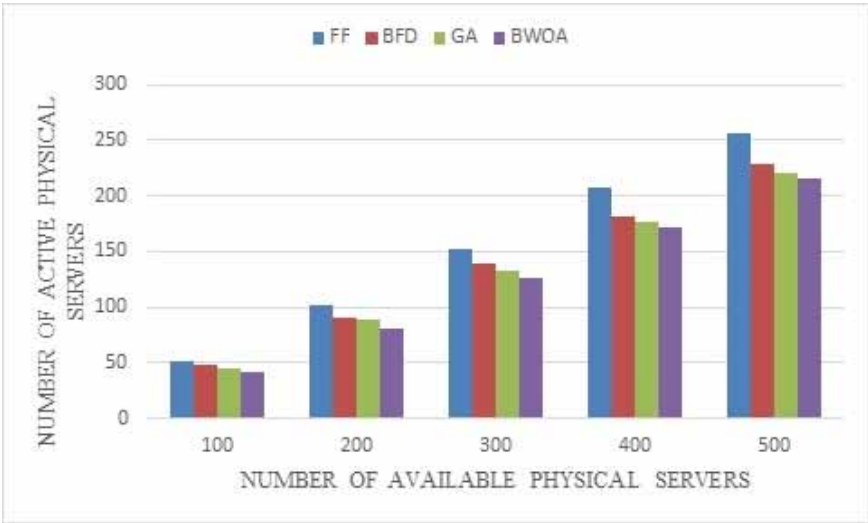


Table 3. Comparison of the Number of Active Servers

	Number of Virtual Machines				
Techniques	100	200	300	400	500
FF	52	102	153	208	256
BFD	48	91	139	181	229
GA	45	89	133	177	221
BWOA	42	81	126	171	216

Figure 12. Number of Active Physical Servers



$$TU^{mem} = \frac{1}{\eta} \sum_{\lambda=1}^{\eta} \left[\frac{\sum_{i=1}^n (x_{i\lambda} \cdot \alpha m_i)}{U_p^{mem}} \right] \quad (21)$$

Here η is the total sum volume of servers running VMs on it.

It is evident from figures 13 and 14 and table 4 and table 5 the method proposed has distributed the resources to VMs in a balanced manner making the uttermost CPU utilization and memory utilization. It can be concluded from figure 13 that FF algorithm's memory utilization is the least, while BFD and GA have shown a quite comparable results in terms of memory utilization. BWOA has shown the maximum utilization of memory. Table 4 shows that 87.89% of memory is utilized through the proposed approach, which is the maximum, while FF has consumed 70.29%, which is the least. It can also be inferred that as the number of VMs increases, memory utilization has also improved.

Figure 14 shows the CPU utilization by the proposed approach and other approaches. The least utilization of CPU can be observed in FF algorithm, while BFD and GA utilization of CPU is greater than FF but less than BWOA. BWOA utilizes the CPU most efficiently as compared to others. It can be concluded from table 5, BWOA uses the CPU by 88.31%, which is the maximum as compared to others. FF does the minimum utilization of 67.24% in the 500 VMs scenario.

5.4 Iterations and Population

The following analysis is performed on the iteration scale. With a similar scenario as discussed before and the VMs volume as 50 and the population as 50, resource wastage efficiency is observed. Figure 15 demonstrates that as the iteration value escalates, resource wastage fitness becomes more stable, and as the iteration goes beyond 200, the fitness resource wastage becomes almost consistent.

The whale population size is another crucial element that significantly affects the algorithm's efficacy. The algorithm's efficiency, convergence, and resource waste value are all affected because of the population's small size. Figure 16 illustrates that resource wastage fitness remains stable as

Table 4. Comparison of Memory utilization

Techniques	Number of Virtual Machines				
	100	200	300	400	500
FF	70.11	70.23	70.52	70.14	70.29
BFD	84.01	85.21	85.32	85.11	85.99
GA	85.23	85.56	85.51	86.89	86.68
BWOA	86.11	86.27	86.58	87.45	87.89

Table 5. Comparison of CPU utilization

Techniques	Number of Virtual Machines				
	100	200	300	400	500
FF	65.40	65.61	65.86	66.52	67.24
BFD	79.12	80.23	80.68	80.99	81.47
GA	85.10	85.51	85.57	85.99	86.20
BWOA	87.23	87.56	87.89	88.20	88.31

Figure 13. Memory Utilization

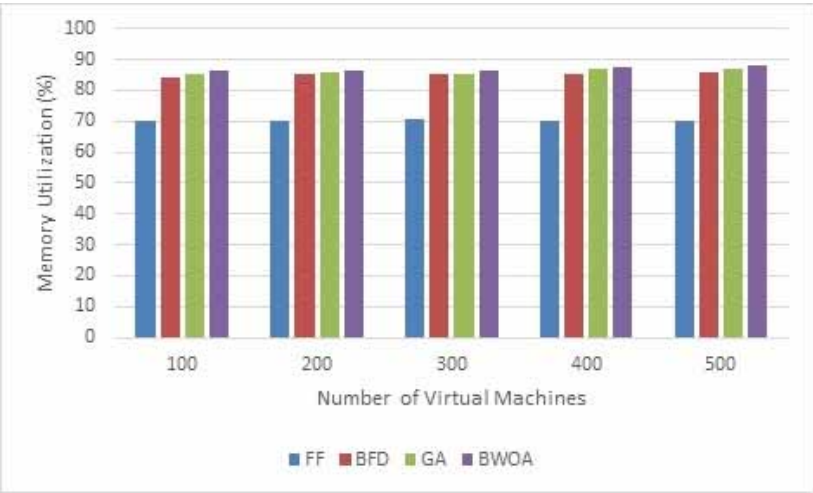
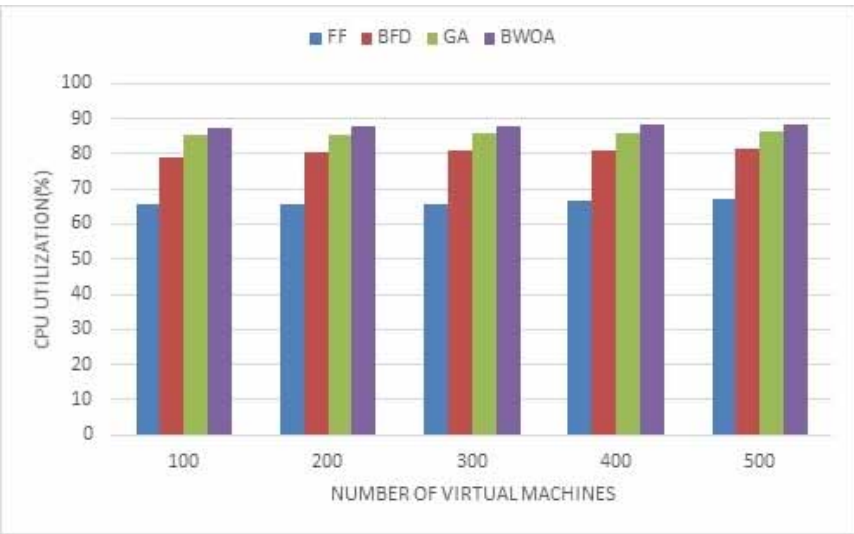


Figure 14. CPU Utilization



the population grows. Thus, it can be inferred that on the larger scale of the population, consistency in the result can be obtained.

6. DISCUSSION

Cloud computing is a technology that has gained widespread acceptance in industries ranging from business to research. The field requires a lot of research for its full and enhanced utilization. In such an infrastructure, VM allocation plays a crucial role in efficiently using all resources. Several existing policies have considered energy, bandwidth, memory and CPU utilization, Response time, and execution time. Very few works have been done, taking resource wastage and active servers' volume as the main consideration. Resource wastage and active servers' volume have been taken

Figure 15. Number of Iterations (50 VMs)

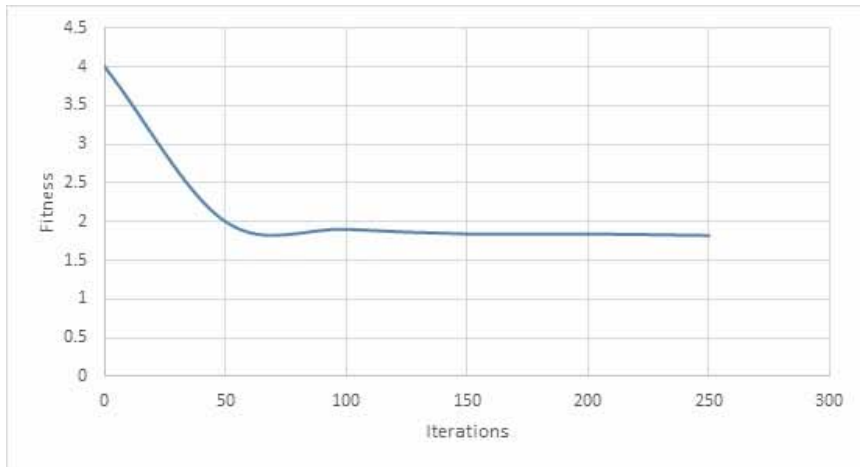
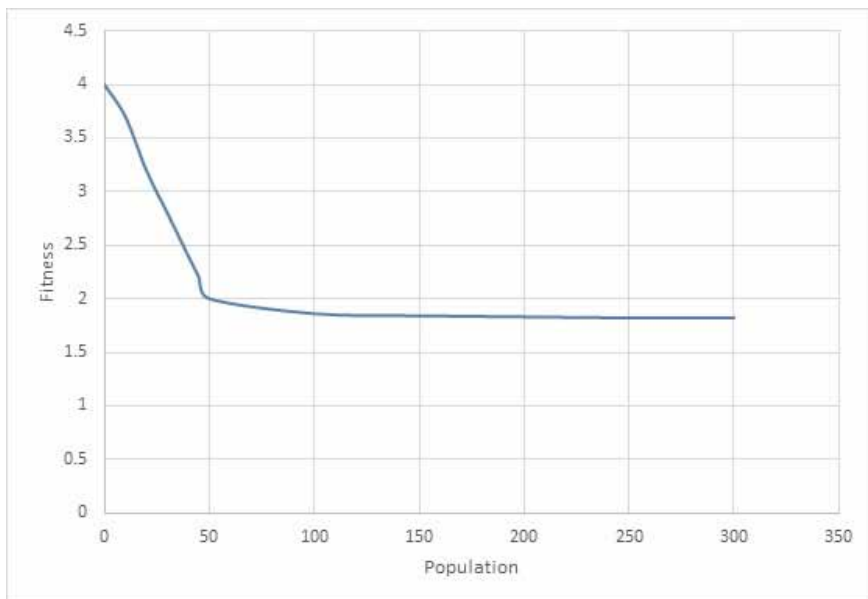


Figure 16. Population Size (50 VMs)



into account to perform an effective allocation in the study. Furthermore, the computed results were compared to FF, BFD, and GA, and it was discovered that the computed results were more efficient.

The work's significance can indeed be summarized as:

- The VM's allocation is heavily influenced by resource use, such as CPU usage, the number of active servers, and memory usage. This has the effect of lowering the cost and energy of the allocation policy.
- The work has adopted BWOA to allocate the VM to the server. This strategy is discrete and applies exploration and exploitation abilities. It easily transitions between exploration and exploitation owing to its adaptive variation, resulting in faster convergence and better results.

- In addition, BFD was employed to accomplish the migration to minimize the number of active servers drastically.
- The outcome of the research helps in efficiently allocating VMs on servers.
- The findings of the study will aid in allocating the VMs such that wastage of resources is highly reduced and the servers' active volume is also minimized to a large extent.
- This further helps in the efficient usage of resources, reducing the energy consumption and amount of resources required and thus reducing the cost of resources and infrastructure.

The overall contribution illustrates the assessment of VMs allocation in the cloud environment and is significant. There are still certain limitations that can be addressed in future efforts. The following are some of them:

- More parameters can be included in addition to the identified ones.
- The results may vary with the number of factors evaluated in the fitness function.
- The cloud metric contemplated in the research proposed is the servers' quantity in active mode and resource wastage, this might be stretched to other metrics like execution time, traffic diversion, SLA violation, and the impact of their collocation on the efficiency.

The overall objective of the research initiative is to minimize resource wastage and active servers' volume with respect to CPU and memory configuration. The outcome of this work will help researchers and cloud practitioners address the allocation issue. The obtained result is verified by comparing them with other conventional approaches.

7. CONCLUSION

This paper enumerates an innovative strategy for optimizing VM allocation with diverse data centers, aiming to lower the volume of servers in active mode and mitigate system resource waste. A modified BWOA is applied where VMs allocation has been performed in the binary space. Firstly, the population is initialized randomly, and the search agent's position is updated for every iteration, which is transformed into binary using a sigmoid function. Gradually, it converges to deliver a globally optimized solution. The framework's performance suggested is validated through simulation. Besides, it provides an optimized solution for memory and CPU consumption. The implication of other relevant parameters on the algorithm's efficiency is also discussed. The feasibility of the effort is asserted based on the implementation analysis and results obtained. The algorithm's capabilities can be extended by including more metrics like execution time, bandwidth, SLA violation, and security as future work. The algorithm can be applied to other cloud issues like task scheduling, load balancing, and dynamic migration.

CONFLICTS OF INTEREST

The authors of this publication declare there is no conflicts of interest.

FUNDING STATEMENT

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors. Funding for this research was covered by the author(s) of the article.

REFERENCES

- Abdessamia, F., Zhang, W. Z., & Tian, Y. C. (2020). Energy-efficiency virtual machine placement based on binary gravitational search algorithm. *Cluster Computing*, 23(3), 1577–1588. doi:10.1007/s10586-019-03021-0
- Abohamama, A. S., & Hamouda, E. (2020). A hybrid energy-aware virtual machine placement algorithm for cloud environments. *Expert Systems with Applications*, 150, 113306. doi:10.1016/j.eswa.2020.113306
- Alharbi, F., Tian, Y. C., Tang, M., Zhang, W. Z., Peng, C., & Fei, M. (2019). An ant colony system for energy-efficient dynamic virtual machine placement in data centers. *Expert Systems with Applications*, 120, 228–238. doi:10.1016/j.eswa.2018.11.029
- Barlaskar, E., Singh, Y. J., & Issac, B. (2018). Enhanced cuckoo search algorithm for virtual machine placement in cloud data centres. *International Journal of Grid and Utility Computing*, 9(1), 1–17. doi:10.1504/IJGUC.2018.090221
- Beloglazov, A., & Buyya, R. (2010). Adaptive threshold-based approach for energy-efficient consolidation of virtual machines in cloud data centers. *MGC@ Middleware*, 4(10.1145), 1890799-803.
- Chaisiri, S., Lee, B. S., & Niyato, D. (2009, December). Optimal virtual machine placement across multiple cloud providers. In 2009 IEEE Asia-Pacific Services Computing Conference (APSCC) (pp. 103-110). IEEE. doi:10.1109/APSCC.2009.5394134
- Chen, H., Xu, Y., Wang, M., & Zhao, X. (2019). A balanced whale optimization algorithm for constrained engineering design problems. *Applied Mathematical Modelling*, 71, 45–59. doi:10.1016/j.apm.2019.02.004
- Dubey, K., & Sharma, S. C. (2020). An extended intelligent water drop approach for efficient VM allocation in secure cloud computing framework. *Journal of King Saud University-Computer and Information Sciences*.
- Gawali, M. B., & Shinde, S. K. (2018). Task scheduling and resource allocation in cloud computing using a heuristic approach. *Journal of Cloud Computing*, 7(1), 1–16.
- Ghetas, M. (2021). A multi-objective Monarch Butterfly Algorithm for virtual machine placement in cloud computing. *Neural Computing & Applications*, 33(17), 1–15. doi:10.1007/s00521-020-05559-2
- Gohil, B. N., Gamit, S., & Patel, D. R. (2021). Fair Fit—A Load Balance Aware VM Placement Algorithm in Cloud Data Centers. In *Advances in communication and computational technology* (pp. 437–451). Springer. doi:10.1007/978-981-15-5341-7_35
- Hussien, A. G., Hassanien, A. E., Houssein, E. H., Amin, M., & Azar, A. T. (2020). New binary whale optimization algorithm for discrete optimization problems. *Engineering Optimization*, 52(6), 945–959. doi:10.1080/0305215X.2019.1624740
- Kaaouache, M. A., & Bouamama, S. (2018). An energy-efficient VM placement method for cloud data centers using a hybrid genetic algorithm. *Journal of Systems and Information Technology*, 20(4), 430–445. doi:10.1108/JSIT-10-2017-0089
- Li, Q., Hao, Q. F., Xiao, L. M., & Li, Z. J. (2011). Adaptive management and multi-objective optimization for virtual machine placement in cloud computing. *Jisuanji Xuebao*, 34(12), 2253–2264. doi:10.3724/SP.J.1016.2011.02253
- Lo, V. M. (1983). *Task assignment in distributed systems* [Doctoral dissertation]. University of Illinois at Urbana-Champaign.
- Mirjalili, S., & Lewis, A. (2016). The whale optimization algorithm. *Advances in Engineering Software*, 95, 51–67. doi:10.1016/j.advengsoft.2016.01.008
- Pahlevan, A., Qu, X., Zapater, M., & Atienza, D. (2017). Integrating heuristic and machine-learning methods for efficient virtual machine allocation in data centers. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 37(8), 1667–1680. doi:10.1109/TCAD.2017.2760517
- Peng, Y., Kang, D. K., Al-Hazemi, F., & Youn, C. H. (2017). Energy and QoS aware resource allocation for heterogeneous sustainable cloud datacenters. *Optical Switching and Networking*, 23, 225–240. doi:10.1016/j.osn.2016.02.001

- Prakash, D. B., & Lakshminarayana, C. (2017). Optimal siting of capacitors in radial distribution network using whale optimization algorithm. *Alexandria Engineering Journal*, 56(4), 499–509. doi:10.1016/j.aej.2016.10.002
- Pushpa, R., & Siddappa, M. (2021). An Optimal Way of VM Placement Strategy in Cloud Computing Platform Using ABCS Algorithm. *International Journal of Ambient Computing and Intelligence*, 12(3), 16–38. doi:10.4018/IJACI.2021070102
- Rankothge, W., Le, F., Russo, A., & Lobo, J. (2017). Optimizing resource allocation for virtualized network functions in a cloud center using genetic algorithms. *IEEE eTransactions on Network and Service Management*, 14(2), 343–356. doi:10.1109/TNSM.2017.2686979
- Saravanan, D., Rajakumar, R., Sreedevi, M., Dinesh, K., Sudha, S. V., Anguraj, D. K., & Mubarakali, A. (2021). Multi-objective swarm-based model for deploying virtual machines on cloud physical servers. *Distributed and Parallel Databases*, 1–19. doi:10.1007/s10619-021-07341-2
- Savitha, S., & Salvi, S. (2021, April). Perceptive VM Allocation in Cloud Data Centers for Effective Resource Management. In *2021 6th International Conference for Convergence in Technology (I2CT)* (pp. 1-5). IEEE. doi:10.1109/I2CT51068.2021.9417960
- Saxena, D., & Singh, A. K. (2021). A proactive autoscaling and energy-efficient VM allocation framework using online multi-resource neural network for cloud data center. *Neurocomputing*, 426, 248–264. doi:10.1016/j.neucom.2020.08.076
- Shabeera, T. P., Kumar, S. M., Salam, S. M., & Krishnan, K. M. (2017). Optimizing VM allocation and data placement for data-intensive applications in cloud using ACO metaheuristic algorithm. *Engineering Science and Technology, an International Journal*, 20(2), 616–628.
- Too, J., Mafarja, M., & Mirjalili, S. (2021). Spatial bound whale optimization algorithm: An efficient high-dimensional feature selection approach. *Neural Computing & Applications*, 33(23), 16229–16250. doi:10.1007/s00521-021-06224-y
- Toutov, A., Toutova, N., Vorozhtsov, A., & Andreev, I. (2021, January). Multicriteria Optimization of Virtual Machine Placement in Cloud Data Centers. In *2021 28th Conference of Open Innovations Association (FRUCT)* (pp. 482-487). IEEE. doi:10.23919/FRUCT50888.2021.9347607
- Usman, M. J., Ismail, A. S., Chizari, H., Abdul-Salaam, G., Usman, A. M., Gital, A. Y., & Aliyu, A. (2019). Energy-efficient virtual machine allocation technique using flower pollination algorithm in cloud datacenter: A panacea to green computing. *Journal of Bionics Engineering*, 16(2), 354–366. doi:10.1007/s42235-019-0030-7
- Vakiliinia, S., Heidarpour, B., & Cheriet, M. (2016). Energy efficient resource allocation in cloud computing environments. *IEEE Access: Practical Innovations, Open Solutions*, 4, 8544–8557. doi:10.1109/ACCESS.2016.2633558
- Widmer, T., Premm, M., & Karaenke, P. (2013, February). Energy-aware service allocation for cloud computing. In *Proceedings of the International Conference on Wirtschaftsinformatik* (pp. 1147-1161). Academic Press.
- Wolpert, D. H., & Macready, W. G. (1997). No free lunch theorems for optimization. *IEEE Transactions on Evolutionary Computation*, 1(1), 67–82. doi:10.1109/4235.585893
- Xia, W., & Shen, L. (2018). Joint resource allocation using evolutionary algorithms in heterogeneous mobile cloud computing networks. *China Communications*, 15(8), 189–204. doi:10.1109/CC.2018.8438283
- Zhang, X., Wu, T., Chen, M., Wei, T., Zhou, J., Hu, S., & Buyya, R. (2019). Energy-aware virtual machine allocation for cloud with resource reservation. *Journal of Systems and Software*, 147, 147–161. doi:10.1016/j.jss.2018.09.084

Ankita Srivastava completed her B.Tech in IT and M.Tech in CS from Dr. A.P.J. Abdul Kalam Technical University, Lucknow, UP, India. Currently, she is pursuing her PhD in Computer Science from Babasaheb Bhimrao Ambedkar University (A Central University), Lucknow, UP, India. Her research interests are Cloud Computing, Nature-Inspired Metaheuristics, Optimization Techniques.

Narander Kumar received his Post Graduate Degree and Ph. D. in CS & IT, from the Department of Computer Science and Information Technology, Faculty of Engineering and Technology, M. J. P. Rohilkhand University, Bareilly, Uttar Pradesh, INDIA in 2002 and 2009, respectively. His current research interest includes Quality of Service (QoS), Cloud Computing, Computer Networks, Resource Management Mechanism, in the networks, Performance Evaluation. Presently he is working as Assistant Professor, in the Department of Computer Science, Babasaheb Bhimrao Ambedkar University (A Central University), Lucknow, UP, India.