

RESEARCH ARTICLE

Artificial Intelligence and Applications
2026, Vol. 4(3) 385–394
DOI: [10.47852/bonviewAIA52026353](https://doi.org/10.47852/bonviewAIA52026353)

BON VIEW PUBLISHING

AI-Driven Adaptive VM Placement Using Performance-to-Power Ratio for Sustainable Data Center Management

Abdelhadi Amahrouch¹ , Youssef Saadi¹ and Said El Kafhali^{2,*} ¹Faculty of Sciences and Techniques, Sultan Moulay Slimane University, Morocco²Faculty of Sciences and Techniques, Hassan First University of Settat, Morocco

Abstract: Cloud data centers provide essential scalable computing resources but often suffer from inefficient resource allocation, resulting in excessive energy consumption and increased carbon emissions. This paper proposes a Q-Learning-driven adaptive virtual machine placement strategy that simultaneously optimizes thermal performance and energy efficiency in heterogeneous data centers. The proposed approach explicitly considers the physical location of servers within racks as well as their processor performance-to-power ratio (PPR). By dynamically adjusting CPU utilization thresholds according to the servers' rack positions, the method ensures that servers operate near their optimal PPR. The algorithm formally classifies servers into “best gear,” “high preferred gear,” and “low preferred gear” states. A reinforcement learning framework based on Q-Learning learns optimal placement policies to minimize energy consumption, maintain stable service-level agreements (SLAs), and reduce the risk of thermal hotspots. Experimental results show that our approach reduces energy usage by 18.43% compared to particle swarm optimization, 20% compared to genetic algorithms, and 13% compared to predictive thermal-aware cloud optimization. Furthermore, it significantly lowers SLA violations and hotspot occurrences. By improving both thermal and energy management, this work contributes to more sustainable, efficient, and environmentally responsible cloud data center operations.

Keywords: cloud computing, VM placement, performance-to-power ratio, thermal risk, Q-Learning, thermal management, resource optimization

1. Introduction

Cloud computing's explosive growth has resulted in a sharp rise in data center energy consumption, posing serious infrastructure, cost, and sustainability challenges for service providers. Approximately one million physical machines and five to six million virtual machines (VMs) are housed in contemporary data centers. Consequently, it is anticipated that within the next years, their overall energy consumption could account for around 3% to 13% of the world's total electricity production by 2030 [1, 2]. Significant economic impacts, such as higher operating costs, and environmental impacts, including increased greenhouse gas emissions, result from this continuous rise in energy consumption.

Today, data centers already consume about 1% or more of all electricity generated worldwide. For example, in China alone, data centers are projected to use about 6.4% of the country's total electricity consumption by 2030, reflecting the growing dependence on digital infrastructures and cloud services [3, 4]. With approximately 40% of total energy consumption, servers remain the largest energy consumers in data centers, powering computing processes and hosting VMs. Storage devices and communication equipment, such as network switches and routers, account for approximately 5% of total energy use. Power supply systems, including uninterruptible power supplies and power distribution units, contribute about 10%. It should also be noted that cooling systems, such as air conditioning and ventilation, continue to use nearly 40% of the energy consumed in data centers.

Maintaining optimal operating temperatures is essential to prevent equipment overheating and to ensure reliability [5, 6].

A range of solutions has been introduced to address the energy and thermal challenges in data centers. Classical numerical methods such as reduced-order models, flow network modeling, and computational fluid dynamics are used to optimize airflow and improve cooling efficiency. Intelligent systems, including adaptive fan controls, have also demonstrated significant potential in reducing energy consumption. In addition, the use of machine learning or deep reinforcement learning is increasing to anticipate thermal fluctuations and optimize the distribution of VMs in real time. Together, these solutions contribute to more efficient energy and thermal management in data centers.

Our contribution presents a model for optimizing the placement of VMs in heterogeneous cloud data centers, with a focus on both energy efficiency and thermal management. It employs the servers' performance-to-power ratio (PPR) along with their physical position in racks to determine the optimal CPU load thresholds for each server. It takes into account factors such as server power consumption, temperature fluctuations, and heat recirculation within the data center. Using reinforcement learning techniques, specifically Q-Learning (QL), our model dynamically adapts the placement of VMs based on thermal and energy characteristics. By strategically distributing workloads across servers with favorable thermal and energy profiles, the model improves overall energy efficiency, reduces cooling requirements, and prevents hot spots from forming. This innovative approach improves both thermal management and resource allocation, resulting in significant reductions in operating costs and carbon footprint.

In addition to operational efficiency and cost savings, the proposed thermal and energy-aware placement strategy directly

*Corresponding author: Said El Kafhali, Faculty of Sciences and Techniques, Hassan First University of Settat, Morocco. Email: said.elkafhali@uhp.ac.ma

contributes to mitigating the environmental impact of data centers. By reducing overall energy consumption and cooling demands, our approach lowers the carbon footprint associated with fossil fuel-based electricity generation. This is increasingly important as the global ICT sector is expected to become a significant source of greenhouse gas emissions if not properly managed. Therefore, optimizing VM placement with awareness of thermal dynamics and energy use supports not only sustainable data center operations but also broader climate goals and environmental commitments.

The remainder of this article is organized as follows. Section 2 reviews related work, focusing on the main contributions in this field. Section 3 offers a detailed description of the system model. Section 4 presents our proposed method. Section 5 presents the results and analysis, offering a detailed assessment of the model's performance as well as comparisons with existing methods. Finally, Section 6 summarizes the paper's findings and suggests possible avenues for future research.

2. Related Work

Saadi et al. [7] explore how combining task clustering methods with scheduling algorithms can improve the execution of scientific workflows on cloud platforms. These distributed scientific applications generate complex workflows with multiple tasks that require significant computing power in cloud data centers. The resulting high workload may result in increased energy consumption and longer runtimes. Job clustering, the consolidation of multiple jobs into a single job, is an effective technique for reducing resource allocation time and minimizing runtime. Using the open-source Workflow-Sim simulator, the study demonstrates that grouping techniques have a major impact on energy usage and lead times. Specific combinations of planning and grouping methods, such as vertical grouping and the Max-Min algorithm, reveal particular energy-saving efficiencies. The research highlights the essential role of careful integration of scheduling algorithms in improving the efficiency of cloud-based scientific workflows.

Monshizadeh Naeen et al. [8] highlight the increasing energy consumption in cloud data centers due to the growing adoption of cloud computing, motivating the exploration of techniques that balance energy savings with quality of service requirements. Server consolidation is widely recognized as a key method to maximize data center productivity by minimizing the number of physical servers in operation while meeting service-level agreements (SLAs). The main challenge is to allocate new VMs to the most suitable hosts to avoid resource waste and unnecessary server activation. This study examines current SC algorithms using CloudSim as a simulation platform and real workload data for evaluation.

Ghasemi et al. [9] highlight the major importance of virtualization in cloud systems, which are among the world's largest energy consumers. This technology plays a central role in resource management, offering effective solutions to many of the challenges faced in this domain. Resource consolidation, particularly through the grouping of VMs, is a widely adopted method to improve the efficiency of cloud datacenters. The rise of container-based systems and containerized workloads has opened up new opportunities for consolidation. In their study, the authors focus specifically on consolidating data centers within distributed cloud environments. They offer a comprehensive overview of computational consolidation at various levels of cloud services, including virtualized data centers, and examine multiple consolidation strategies. Additionally, they introduce a thematic classification and review existing solutions from the literature. The paper concludes by highlighting key research challenges and outlining promising directions for future work, underlining their significance and potential impact.

Amahrouch et al. [10] introduce a model that combines reinforcement learning, a Firefly-based optimization algorithm, and

machine learning-based classification to optimize the placement of VMs in cloud datacenters. This method categorizes VMs into two groups: sensitive VMs requiring strict SLA compliance, and insensitive VMs such as batch workloads. The authors demonstrate the effectiveness of their model in VM consolidation, achieving significant improvements in energy efficiency and reduction of live VM migrations. Their work underscores the importance of combining classification, optimization, and adaptive decision-making to enhance energy efficiency, minimize SLA violations, and enhance overall system performance. This method is in line with current trends in intelligent resource management for cloud environments and provides valuable insights into sustainable cloud infrastructure design.

Tabrizchi et al. [11] address the pressing issue of thermal management in large-scale data centers, emphasizing the growing need to control power consumption due to increased cooling requirements. As data centers grow in size and complexity, it becomes crucial to be able to accurately predict temperature in order to avoid hot spots, optimize energy consumption, and ensure system reliability. The proposed model combines convolutional neural networks (CNNs) with bidirectional LSTM networks (BiLSTMs) to form a hybrid deep learning framework. This model effectively captures the nonlinear and dynamic behavior of temperature variations in data centers, resulting in highly accurate predictions. Their approach demonstrates a significant improvement in prediction accuracy, with an R2 value of 97.15% and lower error rates compared to existing methods. The CNN-BiLSTM model is particularly effective in data centers, where temperature variations are influenced by factors such as server location, processor load, and airflow. By leveraging real sensor data, the model reduces cooling costs and improves energy efficiency, providing a robust solution for cloud service providers seeking to improve sustainability and reliability. This approach also paves the way for future research on thermal anomaly detection and optimization of thermal management strategies in industrial environments.

Kumar et al. [12] review studies conducted between 2015 and 2019 with particular emphasis on the power performance of data centers and the increase in power consumption caused by the widespread use of cloud services. Data centers face several challenges, including heat loss, power factor issues, and inefficient cooling systems, all of which contribute to excessive energy consumption. The study identifies airflow, heat dissipation, and ambient temperature as critical factors for improving energy efficiency. Machine learning algorithms, particularly reinforcement learning techniques, have proven useful in optimizing cooling systems. The article proposes exploring hybrid models combining SVM and Ant Colony Optimization (ACO) algorithms to optimize energy consumption. The study emphasizes the importance of the thermal environment in determining overall efficiency and highlights the need for better automation of cooling processes. Current cooling systems lack intelligent automation, resulting in inefficiencies between consumed energy and expected performance levels. The authors advocate the development of automated systems based on machine learning and hybrid SVM-ACO approaches to optimize cooling parameters, which may contribute to substantial efficiency gains in data center operations.

Akbari et al. [13] suggest a genetic algorithm (GA) to maximize cooling energy efficiency in a cloud data center. The three primary parts of this framework were a GA-based optimization mechanism, a thermodynamics model to estimate cooling power, and an artificial neural network to predict temperature. In order to maximize computing load and minimize air conditioning energy requirements, the static portion of the framework focuses on figuring out how cooling setpoints and computing load are distributed at various levels (room, rack, and row). According to the study, rack-level distribution is the most efficient, lowering cooling energy consumption by 21% to 50%

during dynamic optimization, contingent on operating conditions. The general framework is flexible and useful, even though the temperature prediction model is specific to a particular laboratory setup, which limits its direct applicability to other settings. By using this method, data centers can save a significant amount of energy and enhance the performance of their cooling infrastructure.

Athavale et al. [14] address the challenge of reducing total power usage in heterogeneous data centers by optimizing the distribution of VMs, with a particular focus on thermal considerations. This problem, known as MITEC, is formulated as a nonlinear integer optimization problem that takes into account both computing and cooling energy consumption. The authors propose a VM allocation heuristic based on a GA (MITEC-GA). This approach iteratively refines allocation strategies to achieve near-optimal results, balancing energy efficiency and thermal considerations. Experimental results show that the MITEC-GA heuristic achieves more than 30% energy savings compared to traditional heat-aware greedy algorithms and power-aware VM allocation heuristics. The study highlights the importance of VM consolidation and migration in reducing energy consumption and eliminating thermal hotspots. The MITEC-GA heuristic offers a promising path toward more energy-efficient data center management by improving VM placement while taking thermal dynamics into account. Future research will explore dynamic VM migration and fan speed control to further improve real-time thermal management.

In the interest of improving the thermal performance and energy efficiency of cloud data centers, we propose a new optimization model based on QL, which leverages the PPR of processors and the physical location of servers in racks. Inspired by existing approaches, such as task grouping, VM consolidation, and thermal management algorithms, our method dynamically adjusts processor load and optimizes VM allocation. The goal is to reduce total energy consumption for both computing and cooling systems while minimizing thermal risks. By focusing solely on optimization, our QL model improves temperature regulation and overall data center efficiency while ensuring high performance.

3. System Model

3.1. Data center description

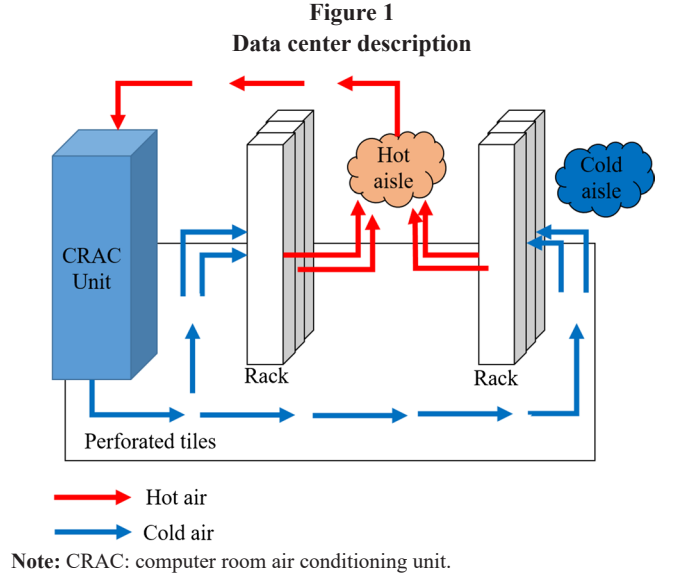
In a typical data center, server racks are positioned on a raised floor with perforated floor slabs, which are arranged into hot aisles and cold aisles. Cool air from the floor, supplied by the cooling system—also known as the computer room air conditioning unit (CRAC)—rises through the perforated tiles into the cold aisles. As it moves through the servers, this air absorbs heat from their internal components before exiting at the back. To complete the cooling cycle, hot aisles are created, and the hot air is sent back to the CRAC, which is often located above the hot aisles. High-density deployment and effective thermal management are enabled by the multiple chassis contained in each rack, as well as the multiple servers housed within each chassis. As illustrated in Figure 1, this configuration ensures a continuous, controlled airflow that facilitates efficient heat dissipation in data centers.

3.2. Power consumption by cooling system

The majority of power consumption in a data center comes from the computing and cooling systems [15–17].

When the processor operates at high performance levels, the host consumes more energy, which is converted into heat. This heat must be dissipated by the CRAC cooling system to prevent the host temperature from exceeding critical limits, thereby avoiding thermal hot spots and potential hardware malfunctions.

The data center's thermal behavior is controlled by the cooling systems (CRACs). The power consumption of the air conditioning



system was estimated using the approximation from the study by Lin et al. [18]. The cooling system's efficiency is measured by the coefficient of performance (CoP), which is the ratio of the power used by the data center's IT components to the power used by the cooling system. The CoP is a quadratic function of the supply air temperature (T_{sup}).

$$CoP(T_{sup}) = \frac{P_{IT}}{P_{cooling}} \quad (1)$$

In Equation (1), P_{IT} and $P_{cooling}$ represent the computing and cooling power consumption in the data center, respectively.

The CoP depends on the data center and can be modeled using regression techniques with various parameters, including workload and supply air temperature. To estimate the CoP, we adopt the model from the HP Utility data center, as shown in Equation (2), where T_{sup} presents the supply air temperature.

$$CoP(T_{sup}) = 0.0068T_{sup}^2 + 0.0008T_{sup} + 0.458 \quad (2)$$

According to Equation (2), by increasing (T_{sup}), we increase CoP (T_{sup}) and consequently the cooling power is reduced.

3.3. IT components power model

The primary IT components in a data center are servers and networking devices. However, the power consumed by network devices is generally considered negligible. Therefore, the total IT power consumption mainly corresponds to the sum of the power consumed by all active servers.

In this study, we use a power model based on CPU utilization, assuming a linear correlation between a server's power consumption and its CPU load. This relationship is expressed in Equation (3) as follows:

$$P(u) = P_{idle} + (P_{max} - P_{idle}) \cdot u \quad (3)$$

where $P(u)$ is the power consumption of the server when its current CPU utilization reaches the level u . P_{idle} and P_{max} represent the power consumption of an idle and a fully utilized server, respectively.

The overall power usage in the data center comprises both the energy consumed by cooling systems and that used by IT components (refer to Equation (4)).

$$P_{total} = P_{cooling} + P_{IT} \quad (4)$$

Using Equations (3) and (4), we formulate the total power consumption P_{total} as follows:

$$P_{total} = \left(1 + \frac{1}{CoP(T_{sup})}\right) P_{IT} \quad (5)$$

3.4. Host temperature (emplacement)

The thermal distribution within server racks is strongly influenced by vertical placement. Tang et al. [19] demonstrated that when all servers are idle, the inlet temperature is not uniform across rack levels. Chassis located near the bottom of the rack benefit from direct access to cold air from floor vents, resulting in lower inlet temperatures. In contrast, upper-level chassis receive less cold airflow and thus operate at higher inlet temperatures due to the reduced cooling supply.

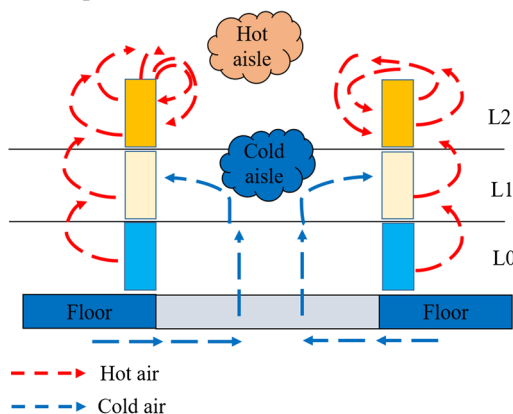
In a related study, Tang et al. [20] examined the impact of hot air recirculation on the upper rack levels. Their findings indicate that hot air expelled by the lower chassis tends to rise and affect the thermal conditions of the upper ones. However, because the top chassis are close to ceiling-level fans or exhaust vents, much of this hot air is quickly directed back to the cooling system, partially mitigating its thermal impact.

Xu et al. [21] focused on the dependency between rack position and heat dissipation performance. Their results indicate that the higher a server is placed in the rack, the lower its cooling efficiency. Specifically, each higher rack level experiences an average temperature increase of between 2.7 °C and 3.4 °C, demonstrating the importance of thermally-aware VM placement strategies in energy-efficient cloud data centers.

Based on these observations, we categorize rack levels into three thermal zones: Level L0, closest to the cold air sources, offers the lowest inlet temperature; Level L1, situated in the middle, receives less cold air and experiences more hot air recirculation; and Level L2, at the top, is most affected by hot air accumulation, leading to the highest operating temperatures. Figure 2 illustrates this vertical temperature distribution across the rack levels.

Because of the complex nature of the airflow inside the data center, a part of the hot air exhausted from the server outlets recirculates into the inlet of other servers. One of the reasons for the energy overhead of the cooling system is mainly due to the mixing of cold air with the hot air recirculated into the inlet. To provide an acceptable inlet temperature

Figure 2
Temperature distribution across rack levels



for all servers, the air temperature must be lowered, which increases the energy cost of cooling. To avoid hot spots in the data center, the position of the servers in the racks must be taken into consideration when setting the maximum load threshold.

3.5. Performance power ratio (PPR)

PPR represents the energy efficiency of a server, indicating the number of operations executed per watt of power used. This metric provides useful information on the efficiency of a system's energy use, particularly under varying workload conditions. According to results published on the SPEC website [22], several manufacturers have evaluated their servers using the SPECpower_ssj2008 benchmark. Figure 3 [22] illustrates the corresponding graphs for two types of servers, showing the evolution of power consumption and PPR as a function of increasing processor load. In the case of the Dell PowerEdge R830 (Intel Xeon E5-4669 v4 2.20 GHz), the highest PPR is reached at around 80% CPU utilization, whereas for the HP Apollo XL225n Gen 10 Plus (AMD EPYC 7702 2.0 GHz), the maximum PPR is observed at full CPU load (100%). These curves highlight the importance of identifying the optimum utilization point for each server model to maximize energy efficiency.

The highest PPR varies by server type; for example, it peaks at 80% CPU utilization for the Dell PowerEdge R830 and at 100% for the HP Apollo XL225n Gen10 Plus. Power consumption rises nearly linearly with CPU utilization.

To enhance energy efficiency, we propose an allocation strategy that determines the best PPR gear for each server and establishes it as

Figure 3
PPR and power usage as CPU utilization rises (SPECpower_ssj2008): (a) Dell PowerEdge R830 (Intel Xeon E5-4669 v4 2.20 GHz) and (b) HP Apollo XL225n Gen 10 Plus (AMD EPYC 7702 2.0 GHz)

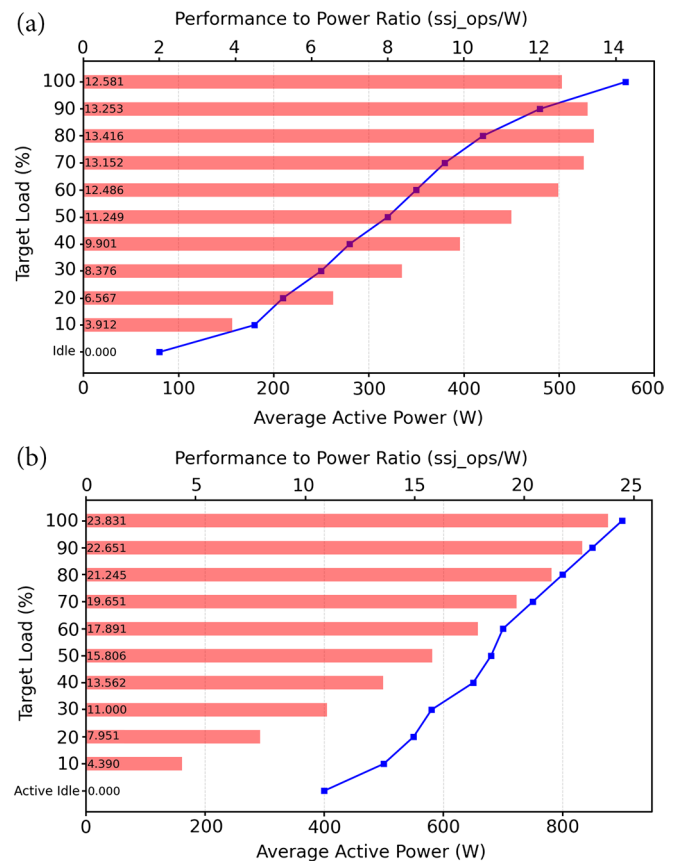
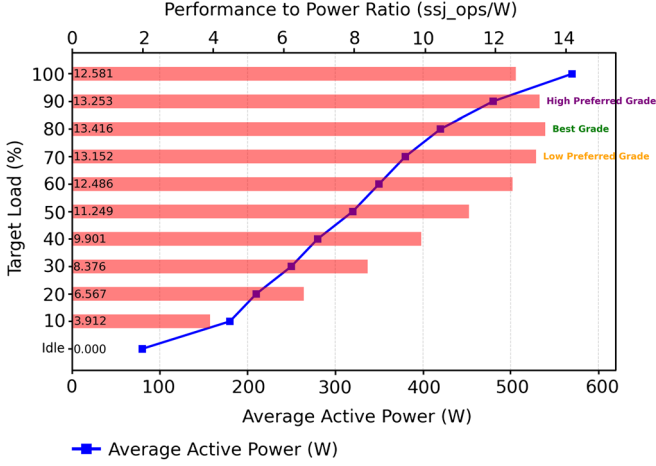


Figure 4
PPR and power usage as CPU utilization rises: Dell PowerEdge R830 (Intel Xeon E5-4669 v4 2.20 GHz)



a CPU load threshold. In a heterogeneous data center, this CPU load threshold varies per server.

To ensure energy efficiency, servers are loaded close to their optimal PPR—that is, where each server executes the highest number of instructions per watt consumed. The exact CPU load threshold is defined based on the server's position within the rack. Specifically, servers located in Level L0 are assigned the High Preferred Gear (HPG), those in L1 the Best Gear (BG), and those in L2 the Low Preferred Gear (LPG).

Figure 4 illustrates this concept for the Dell PowerEdge R830, where the BG corresponds to approximately 80% CPU utilization, the HPG to 90%, and the LPG to 70%. This stratified thresholding enables a thermally-aware and energy-efficient VM placement strategy, tailored to both server characteristics and their rack positions.

The BG corresponds to the CPU utilization that maximizes the PPR for each server. The HPG and LPG are then defined as upper and lower operating ranges around this optimum, typically within $\pm 10\%$ of the BG point. This ensures that each server operates close to its most energy-efficient state while considering thermal constraints and rack position.

4. Proposed Approach

4.1. Problem modeling

The placement of VMs with respect to temperature control and energy consumption is an NP-hard problem. An optimization problem that aims to minimize overall energy consumption while respecting all operational constraints can be used to formulate the best location for VMs. The optimization problem that this work attempts to solve is formally defined as follows:

- 1) Minimize the total energy consumption of both the IT equipment and the cooling system.
- 2) Ensure satisfactory machine performance by respecting CPU load constraints based on the PPR and the host's position in the racks.

In this study, the user workload requests to the data center are abstracted as VMs. A data center consists of t racks, each rack contains r chassis, and each chassis holds n servers. The objective is to place m VMs across these servers optimally. The power consumed for all servers (PIT) is calculated by Equation (6).

$$P_{IT} = \sum_{i=1}^n x_i P_i, \quad x_i = \{0, 1\} \quad (6)$$

x_i is a binary variable that takes 1 when the host i is active and zero if it is switched off.

The final aim is to minimize P_{total} , the sum of power consumed by IT equipment and cooling systems, represented mathematically as follows:

$$\min P_{total} = \min \left(1 + \frac{1}{CoP(T_{sup})} \right) P_{IT} = \min \sum_{i=1}^n x_i \left(1 + \frac{1}{CoP(T_{sup})} \right) P_i \quad (7)$$

4.2. Threshold determination

For balanced thermal management and optimum energy efficiency in heterogeneous data centers, dynamically setting the CPU load threshold for each server based on its physical rack location is crucial. Algorithm 1 outlines a straightforward but effective method for determining the appropriate CPU load threshold according to the server's surrounding temperature.

Servers located in the lower part of the rack (L0 level), which benefit from better cold air circulation, can support a higher load (HPG). Servers located in the middle (L1) operate at an optimum threshold (BG), whereas those at the top (L2), exposed to recirculated hot air, are limited to a lower load (LPG) to avoid overheating.

- 1) L0: Servers in the lower part of the rack, where there's more cool air. CPU load threshold = HPG.
- 2) L1: Intermediate servers, where the temperature is moderate. CPU load threshold = BG.
- 3) L2: Servers located at the top part of the rack, where hot air tends to concentrate. CPU load threshold = LPG.

Algorithm 1: Server threshold determination

Input: hostList

Output: CPU Load Threshold (CPUTh) assigned to each server based on its rack position

```

foreach server in hostList do
    If server.getPosition() is L0 then
        setCpuTh (HPG)
    else if server.getPosition() is L1 then
        setCpuTh (BG)
    else if server.getPosition() is L2 then
        setCpuTh (LPG)
    end if
end foreach
    
```

4.3. Proposed VM placement method

Q-Learning is a type of reinforcement learning that enables systems to learn the best actions to take through experience. Instead of relying on a predefined model, it learns by trying different actions and observing the results, gradually improving its decisions based on accumulated feedback.

Applied to the placement of VMs in a data center, QL makes real-time decisions by observing current system conditions, such as energy consumption, temperature levels, and resource utilization. Based on this information, it chooses actions such as moving a VM to another server or starting or stopping a server. Each action is evaluated using a reward signal that reflects how much it contributes to improving energy efficiency, maintaining thermal stability, or improving overall system performance. Over time, as the system gains experience, it develops an increasingly effective strategy for VM placement. The result is a smarter distribution of workloads that helps reduce energy consumption, prevent overheating, and make better use of available resources, all without the need for an exact mathematical model of data center behavior.

- 1) State Space (S): Represents the current configuration of VMs and hosts in the data center.

$$S = [\text{CPU utilization, temperature, host states}] \quad (8)$$

- 2) Action Space (A): Represents the possible actions for VM placement or migration.

$$A = \{\text{Place VM from Host}_i \text{ to Host}_j, \text{Turn on Host, Turn Off Host}\} \quad (9)$$

- 3) Reward Function (R (s, a)): Evaluates the impact of an action a in states.

- a. Energy efficiency

$$R_{\text{energy}} = -\text{Total Energy Consumption (E)} \quad (10)$$

- b. Thermal management

$$R_{\text{thermal}} = \begin{cases} -T_{\text{max}} & \text{if } T_{\text{max}} \leq \text{Threshold,} \\ -\infty & \text{if } T_{\text{max}} > \text{Threshold} \end{cases} \quad (11)$$

- c. Load balancing

$$R_{\text{balance}} = -\text{Variance of CPU utilization across hosts} \quad (12)$$

- d. Combined reward function

$$R(s, a) = w_1 \cdot R_{\text{energy}} + w_2 \cdot R_{\text{thermal}} + w_3 \cdot R_{\text{balance}}, \quad (13)$$

where w_1 , w_2 , and w_3 are weights to prioritize objectives.

- 4) $Q(s, a)$: is the cumulative reward of taking action a in state s :

$$Q(s, a) \leftarrow Q(s, a) + \alpha [R(s, a) + \gamma \max_{a'} Q(s', a') - Q(s, a)] \quad (14)$$

- a. s : state.
 b. a : Action.
 c. $R(s, a)$ Immediate reward.
 d. s' : Next state (After action a).
 e. α : Learning rate.
 f. γ : Discount factor.

- 5) Exploration Strategy (ϵ_{Greedy}): Balances exploration (try new actions) and exploitation (use learned knowledge):

- a. With probability ϵ (epsilon), select a random action.
 b. Otherwise, select the action a with the highest Q -value

The QL procedure for VM placement optimization is summarized in Algorithm 2. This algorithm iteratively updates the Q -table based on the observed state, selected action, and received reward, balancing exploration and exploitation through the ϵ -greedy strategy.

Algorithm 2: Q-Learning for VM Placement Optimization

- Input:** S (State space), A (Action space), Learning rate, Discount factor, Max iterations.
Output: Optimal Q -table for VM placement decisions
 1. **Initialize** Q -table with arbitrary values for all (s in S) and all (a in A).
 2. **Repeat** for (I) iterations or until convergence:
 a. **Observe** the state (s).
 b. **Select** action a using an exploration strategy (e.g., (epsilon-greedy):
 - With probability (epsilon), choose a random action.
 - Otherwise, choose ($a = \arg(\max Q(s, a))$).
 c. **Execute** action a , observe the next state s' and reward ($R(s, a)$)
 d. Update the Q -value using Equation (12):
 3. **Return** Q -table $Q(s, a)$.

4.4. Algorithms for comparison

In this study, the performance of the QL-based model is evaluated against three established VM placement techniques. Approaches such as particle swarm optimization (PSO), GA, and predictive thermal-aware cloud optimization (PTACO)-MMT are employed.

- 1) PSO: is a “swarm intelligence” algorithm that uses the collective behavior of particles to find the optimal position for virtual machines, minimizing power consumption and improving resource utilization [23].
 2) GA: is an evolutionary algorithm inspired by natural selection. It is used to generate optimal placement solutions through crossover and mutation operations, aiming to reduce overall energy consumption [24].
 3) PTACO-MMT [25]: is a power- and thermal-aware VM placement scheme that reduces the total energy consumption of both IT and cooling systems. It combines a VM selection algorithm based on the Minimization Method and a VM placement algorithm based on an improved ACO approach.

Although more recent techniques have been proposed, these methods remain widely used as benchmark algorithms for performance comparison, providing a robust framework to evaluate the effectiveness of the QL model based on the power performance ratio. To ensure robustness, two models of host behavior described by Wang et al. [26] are also considered: *lr_1.2_mmt* and *igr_1.5_mc*, which are based on the following methods:

- 1) IqrMc: This approach uses a VM allocation strategy that relies on the Inter Quartile Range technique for detecting hotspots, combined with a policy that prioritizes migrating VMs exhibiting the highest correlation.
 2) LrMmt: Based on Cleveland’s local regression method, LrMmt determines whether a host is overloaded and if its VMs should be migrated. When overload is detected, it selects the VM with the shortest migration time, applying a minimum migration time strategy.

5. Result and Analysis

5.1. Simulation setup

This study utilized real workload traces from the Microsoft Azure 2019 dataset, focusing specifically on CPU utilization data collected over ten randomly selected days. The dataset encompasses approximately 2.6 million VMs and 1.9 billion utilization records, captured at 5-minute intervals [27].

For the simulation setup, CloudSim was used to model a cloud data center consisting of 800 identical host machines. The data center is organized into 10 zones, each containing 10 racks arranged in a 4×2 layout. Each rack houses 10 hosts, resulting in a total of 800 hosts. Four types of single-core VMs, based on Amazon EC2 specifications, were

Table 1
VMs specifications

VM type	MIPS	Core (processing elements)	RAM (MB)	Bandwidth (Mbps)
High-CPU medium instance	2,500	1	870	100
Extra-large instance	2,000	1	1,740	100
Small instance	1,000	1	1,740	100
Micro instance	500	1	613	100

Note: VM: virtual machine.

Table 2
Experimental parameters

Item	Parameters	Valeur
Datacenter	Number of zones	10
	Number of racks in zone	4*2
	Number of hosts in the rack	10
CRAC	Supply air temperature from CRAC	25 °C
Server	Heat capacity	340 J/K
	Thermal resistance	0,34 K/W
	Initial CPU temperature	318 K
	Threshold of CPU temperature	75 °C
	Upper utilization threshold	80%

Note: CRAC: computer room air conditioning unit.

Table 3
Number of virtual machines per day

Date	Number of VMs
Day 1	1,052
Day 2	898
Day 3	1,061
Day 4	1,516
Day 5	1,078
Day 6	1,463
Day 7	1,358
Day 8	1,233
Day 9	1,054
Day 10	1,033

Note: CRAC: computer room air conditioning unit.

defined with processing capacities ranging from 500 to 2,500 MIPS, RAM between 613 and 1,740 MB, and a standardized bandwidth of 100 Mbps. Detailed VM specifications are provided in Table 1.

The thermal parameters are detailed in Table 2.

The initial allocation of VMs is based on the resource requirements defined by the VM type. However, during their lifetime, VMs consume fewer resources due to workload variations, allowing dynamic VM consolidation. Real system workload traces were employed to conduct the experiments. Data from Microsoft Azure 2019 nodes, including CPU utilization from 10 randomly selected days and over 1,000 VMs, were used. Utilization measurements were recorded at 5-minute intervals. Table 3 summarizes the number of VMs observed each day, ranging from 898 to 1,516 VMs per day.

5.2. Results and discussion

To evaluate the effectiveness of our model, four key performance indicators are used, each providing insight into different aspects of system behavior. Energy consumption, measured in kilowatt-hours (kWh), quantifies the total energy used by the data center and serves as a primary indicator of energy efficiency. SLA violation captures the degradation in service quality caused by dynamic VM consolidation, taking into account both the impact of VM migrations—which can momentarily reduce application performance—and the occurrence of hotspots, which represent servers that have exceeded a critical temperature threshold, potentially threatening hardware reliability. The

VM migration count directly reflects the number of live VM relocations executed during operation, providing a measure of system reactivity and consolidation cost. Finally, hotspots are used to monitor thermal stability, with a higher number indicating poor thermal distribution and greater risk of overheating, which may necessitate cooling

Figure 5
Energy consumption (kWh) under different workload datasets

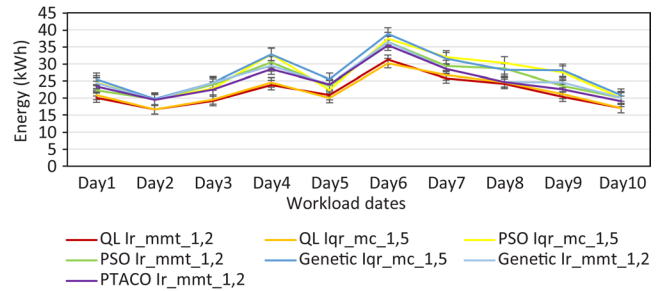


Figure 6
Average energy consumption (kWh)

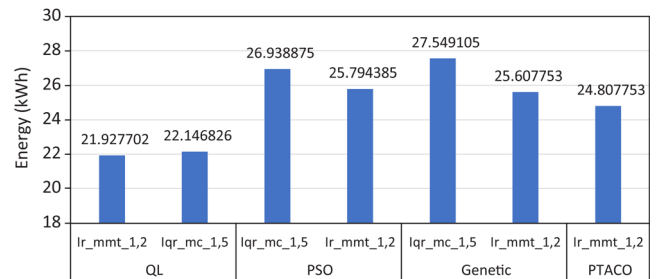
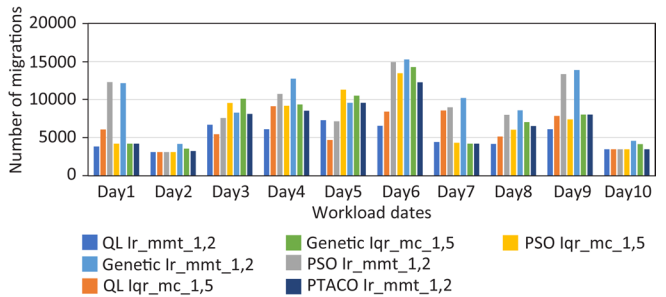
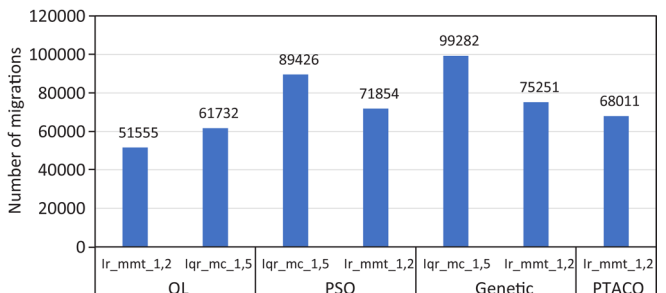


Figure 7
Comparison of VM migration counts under different loads



Note: VM: virtual machine.

Figure 8
Average number of migrations



Note: PSO: particle swarm optimization. PTACO: predictive thermal-aware cloud optimization. QL: Q-Learning.

intervention or trigger thermal throttling. Combined, these metrics provide a thorough evaluation of the balance between energy efficiency, operational performance, and thermal safety.

To evaluate our methodology and analyze the results, we performed a comparative study including PSO, GA, and our QL model. We considered two types of host behavior models: lr_1.2_mmt and iqr_1.5_mc.

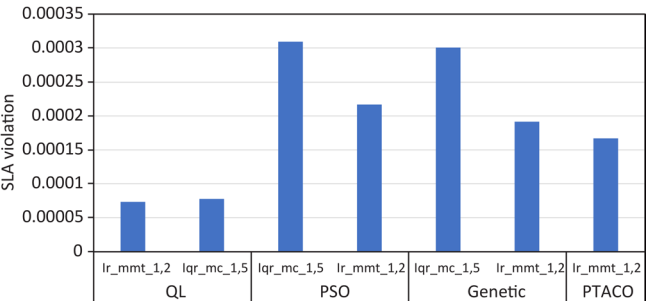
From the analysis of Figures 5 and 6, it is clear that VM placement using QL yields the lowest energy consumption among the methods tested. Specifically, under the lr_1.2_mmt strategy, our Q-Learning model achieves an energy reduction of 18.43% compared to PSO, 20% compared to the GA, and 13% compared to PTACO. Similarly, with the iqr_1.5_mc strategy, our model reduces energy consumption by 14.14% compared to PSO and 13.51% compared to the GA.

Observing Figures 7 and 8, we observe that the QL-based model consistently achieves the lowest number of VM migrations across both host behavior strategies. Under the lr_1.2_mmt configuration (Figure 7), our model significantly outperforms the others with only 51,555 migrations, compared to 71,854 for PSO, 75,251 for the GA, and 68,011 for PTACO.

Similarly, under the iqr_1.5_mc, the QL approach results in 61,732 migrations, whereas GA and PSO record 89,426 and 99,282 migrations, respectively.

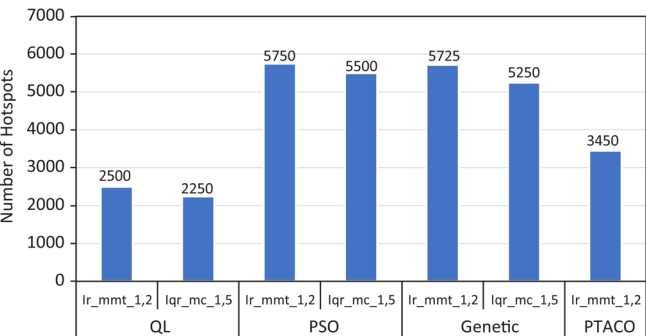
These results demonstrate the effectiveness of the proposed approach in reducing the number of migrations, which is vital for ensuring service stability and lowering energy costs during dynamic VM consolidation.

Figure 9
Average SLA violation



Note: SLA: service-level agreement.

Figure 10
Average hotspots



Note: PSO: particle swarm optimization. PTACO: predictive thermal-aware cloud optimization. QL: Q-Learning.

Figure 9 shows the percentage of SLA violations. Our model outperforms the other approaches under both the iqr_1.5_mc and lr_1.2_mmt strategies, maintaining a consistently low violation rate over the ten-day evaluation period. This performance is primarily attributed to the reduced number of migrations achieved by the QL method, which helps preserve service quality.

Our QL model demonstrates a significant reduction in thermal hotspots, recording only 2,500 and 2,250 occurrences under the lr_1.2_mmt and iqr_1.5_mc strategies, respectively. In contrast, PSO registers 5,750 and 5,500, whereas Genetic records 5,725 and 5,250 under the same strategies. PTACO also performs better than PSO and Genetic, but remains superior to QL, with 3,450 hotspots under lr_1.2_mmt. These results highlight the superior effectiveness of QL in managing thermal conditions, thanks to its adaptive learning capability, whereas PSO and Genetic, which are less responsive to dynamic variations in workload, exhibit a higher number of hot spots (see Figure 10).

6. Conclusion

In this study, we introduced a new strategy for placing VMs in heterogeneous cloud data centers, focusing simultaneously on energy efficiency and thermal management. Our method considers each server's PPR and its physical placement within the rack, allowing dynamic adjustment of CPU load thresholds to maximize efficiency and minimize thermal hotspots.

By leveraging QL, the system learns optimal placement decisions over time, resulting in smarter resource allocation that balances thermal risks and power usage patterns. Experimental results confirm the effectiveness of the model, showing clear reductions in overall energy consumption and a lower risk of overheating compared to traditional methods such as PSO and GAs.

These improvements directly contribute to lowering the carbon footprint of data centers by reducing the energy required for both computation and cooling, supporting a more sustainable digital infrastructure.

In the future, integrating real-time monitoring data could further enhance the adaptability of the placement strategy under dynamic workloads. Combining QL with other advanced machine learning methods or hybrid optimization approaches also shows promise for improving both energy and thermal performance. Finally, future work could investigate the impact of network traffic on heat generation and energy use to develop holistic solutions that jointly optimize computing, cooling, and communication in modern cloud environments.

Ethical Statement

This study does not contain any studies with human or animal subjects performed by any of the authors.

Conflicts of Interest

The authors declare that they have no conflicts of interest to this work.

Data Availability Statement

The data that support the findings of this study are openly available in the Microsoft Azure 2019 dataset at <https://github.com/Azure/AzurePublicDataset>, reference number [27].

Author Contribution Statement

Abdelhadi Amahrouch: Conceptualization, Methodology, Software, Investigation, Resources, Data curation, Writing – original draft, Writing – review & editing, Visualization. **Youssef Saadi:** Conceptualization, Methodology, Validation, Formal analysis, Investigation, Resources, Supervision, Project administration. **Said El Kafhali:** Conceptualization, Methodology, Validation, Formal analysis, Investigation, Resources, Writing – review & editing, Supervision, Project administration.

References

- [1] Wen, J., Chen, Z., Jin, X., & Liu, X. (2023). Rise of the planet of serverless computing: A systematic review. *ACM Transactions on Software Engineering and Methodology*, 32(5), 131. <https://doi.org/10.1145/3579643>
- [2] Buyya, R., Ilager, S., & Arroba, P. (2024). Energy-efficiency and sustainability in new generation cloud computing: A vision and directions for integrated management of data centre resources and workloads. *Software: Practice and Experience*, 54(1), 24–38. <https://doi.org/10.1002/spe.3248>
- [3] Rasoulpour Shabestari, E., & Shameli-Sendi, A. (2025). An intelligent VM placement method for minimizing energy cost and carbon emission in distributed cloud data centers. *Journal of Grid Computing*, 23(1), 12. <https://doi.org/10.1007/s10723-025-09798-2>
- [4] Qazi, F., Kwak, D., Khan, F. G., Ali, F., & Khan, S. U. (2024). Service Level Agreement in cloud computing: Taxonomy, prospects, and challenges. *Internet of Things*, 25, 101126. <https://doi.org/10.1016/j.iot.2024.101126>
- [5] Kotteswari, K., Dhanaraj, R. K., Balusamy, B., Nayyar, A., & Sharma, A. K. (2025). EELB: An energy-efficient load balancing model for cloud environment using Markov decision process. *Computing*, 107(3), 81. <https://doi.org/10.1007/s00607-025-01439-6>
- [6] Singh, J., & Walia, N. K. (2023). A comprehensive review of cloud computing virtual machine consolidation. *IEEE Access*, 11, 106190–106209. <https://doi.org/10.1109/ACCESS.2023.3314613>
- [7] Saadi, Y., Jounaidi, S., El Kafhali, S., & Zougagh, H. (2023). Reducing energy footprint in cloud computing: A study on the impact of clustering techniques and scheduling algorithms for scientific workflows. *Computing*, 105(10), 2231–2261. <https://doi.org/10.1007/s00607-023-01182-w>
- [8] Monshizadeh Naeen, M. A., Ghaffari, H. R., & Monshizadeh Naeen, H. (2024). Cloud data center cost management using virtual machine consolidation with an improved artificial feeding birds algorithm. *Computing*, 106(6), 1795–1823. <https://doi.org/10.1007/s00607-024-01267-0>
- [9] Ghasemi, A., Toroghi Haghighat, A., & Keshavarzi, A. (2024). Enhancing virtual machine placement efficiency in cloud data centers: A hybrid approach using multi-objective reinforcement learning and clustering strategies. *Computing*, 106(9), 2897–2922. <https://doi.org/10.1007/s00607-024-01311-z>
- [10] Amahrouch, A., Saadi, Y., & El Kafhali, S. (2025). Optimizing energy efficiency in cloud data centers: A reinforcement learning-based virtual machine placement strategy. *Network*, 5(2), 17. <https://doi.org/10.3390/network5020017>
- [11] Tabrizchi, H., Razmara, J., & Mosavi, A. (2023). Thermal prediction for energy management of clouds using a hybrid model based on CNN and stacking multi-layer bi-directional LSTM. *Energy Reports*, 9, 2253–2268. <https://doi.org/10.1016/j.egy.2023.01.032>
- [12] Kumar, R., Khatri, S. K., & José Diván, M. (2020). Effect of cooling systems on the energy efficiency of data centers: Machine learning optimisation. In *2020 International Conference on Computational Performance Evaluation*, 596–600. <https://doi.org/10.1109/ComPE49325.2020.9200088>
- [13] Akbari, A., Khonsari, A., & Ghoreyshi, S. M. (2020). Thermal-aware virtual machine allocation for heterogeneous cloud data centers. *Energies*, 13(11), 2880. <https://doi.org/10.3390/en13112880>
- [14] Athavale, J., Yoda, M., & Joshi, Y. (2021). Genetic algorithm-based cooling energy optimization of data centers. *International Journal of Numerical Methods for Heat & Fluid Flow*, 31(10), 3148–3168. <https://doi.org/10.1108/HFF-01-2020-0036>
- [15] Chen, R., Liu, B., Lin, W., Lin, J., Cheng, H., & Li, K. (2023). Power and thermal-aware virtual machine scheduling optimization in cloud data center. *Future Generation Computer Systems*, 145, 578–589. <https://doi.org/10.1016/j.future.2023.03.049>
- [16] Mongia, V. (2024). EMaC: Dynamic VM consolidation framework for energy-efficiency and multi-metric SLA compliance in cloud data centers. *SN Computer Science*, 5(5), 643. <https://doi.org/10.1007/s42979-024-02982-3>
- [17] Feng, H., Deng, Y., & Li, J. (2021). A global-energy-aware virtual machine placement strategy for cloud data centers. *Journal of Systems Architecture*, 116, 102048. <https://doi.org/10.1016/j.sysarc.2021.102048>
- [18] Lin, W., Yu, T., Gao, C., Liu, F., Li, T., Fong, S., & Wang, Y. (2021). A hardware-aware CPU power measurement based on the power-exponent function model for cloud servers. *Information Sciences*, 547, 1045–1065. <https://doi.org/10.1016/j.ins.2020.09.033>
- [19] Tang, Q., Gupta, S. K. S., & Varsamopoulos, G. (2008). Energy-efficient thermal-aware task scheduling for homogeneous high-performance computing data centers: A cyber-physical approach. *IEEE Transactions on Parallel and Distributed Systems*, 19(11), 1458–1472. <https://doi.org/10.1109/TPDS.2008.111>
- [20] Tang, Q., Gupta, S. K. S., & Varsamopoulos, G. (2007). Thermal-aware task scheduling for data centers through minimizing heat recirculation. In *2007 IEEE International Conference on Cluster Computing*, 129–138. <https://doi.org/10.1109/CLUSTER.2007.4629225>
- [21] Xu, X., Zhang, Q., Maneas, S., Sotiriadis, S., Gavan, C., & Bessis, N. (2019). VMSAGE: A virtual machine scheduling algorithm based on the gravitational effect for green cloud computing. *Simulation Modelling Practice and Theory*, 93, 87–103. <https://doi.org/10.1016/j.simpat.2018.10.006>
- [22] Standard Performance Evaluation Corporation. (2008). *SPECpower_ssj 2008 results*. https://www.spec.org/power_ssj2008/results
- [23] Rozekkhani, S. M., Mahan, F., & Pedrycz, W. (2024). Efficient cloud data center: An adaptive framework for dynamic Virtual Machine Consolidation. *Journal of Network and Computer Applications*, 226, 103885. <https://doi.org/10.1016/j.jnca.2024.103885>
- [24] Alourani, A., Khalid, A., Tahir, M., & Sardaraz, M. (2024). Energy efficient virtual machines placement in cloud datacenters using genetic algorithm and adaptive thresholds. *PLOS One*, 19(1), e0296399. <https://doi.org/10.1371/journal.pone.0296399>
- [25] Wang, J. V., Cheng, C.-T., & Tse, C. K. (2015). A power and thermal-aware virtual machine allocation mechanism for cloud data centers.

- In *IEEE International Conference on Communication Workshop*, 2850–2855. <https://doi.org/10.1109/ICCW.2015.7247611>
- [26] Wang, J., Gu, H., Yu, J., Song, Y., He, X., & Song, Y. (2022). Research on virtual machine consolidation strategy based on combined prediction and energy-aware in cloud computing platform. *Journal of Cloud Computing*, 11(1), 50. <https://doi.org/10.1186/s13677-022-00309-2>
- [27] Cortez, E. (n.d.). *AzurePublicDataset* [Data set]. <https://github.com/Azure/AzurePublicDataset>

How to Cite: Amahrouch, A., Saadi, Y., & El Kafhali, S. (2026). AI-Driven Adaptive VM Placement Using Performance-to-Power Ratio for Sustainable Data Center Management. *Artificial Intelligence and Applications*, 4(3), 385–394. <https://doi.org/10.47852/bonviewAIA52026353>