

An AI-Assisted VMware Resource Management System

Mrwan A. MARGEM*, Elias W. AJAL

College of Electronic Technology -Tripoli / Libya

* m_mrwan@yahoo.com or m.margem@cet.edu.ly

Abstract

Virtualized data centers running on VMware struggle with resource management. Traditional, reactive allocation strategies trigger high operational expenses and performance bottlenecks. To address this, we developed an "AI-Assisted VMware Resource Management System" designed to establish a predictive, autonomous operational framework. The architecture pairs machine learning for resource demand forecasting with Deep Reinforcement Learning (DRL) for real-time allocation decisions. Operating via VMware vSphere APIs, the platform directly controls virtual machine workloads without human intervention. We validated a functional prototype in a simulated environment, benchmarking its impact on cost reduction and resource utilization against static allocation baselines. CPU costs dropped by nearly 87%; RAM costs by half. Thus, the resulting framework confirms that integrating predictive and reinforcement models substantially improves virtualized IT infrastructure efficiency, cost-effectiveness, and stability.

Keywords: VMware, Cloud Resource Management, Workload Prediction, LSTM, Deep Reinforcement Learning, SLA Compliance, Cost Optimization.

Submitted: 19/08/2026

Accepted: 08/09/2026

1. Introduction

The rise of cloud computing and large-scale virtualization has fundamentally reshaped modern information technology. These technologies enable organizations to achieve unprecedented agility, scalability, and cost-efficiency by abstracting physical hardware and delivering resources on demand (Kanungo, 2024). Platforms such as VMware vSphere have become the cornerstone of private and hybrid clouds, allowing numerous virtual machines (VMs) to be consolidated onto fewer physical servers. This consolidation maximizes hardware utilization and reduces physical footprint. Yet it introduces a formidable challenge: the complex, dynamic task of resource management (Patil et al., 2025).

Efficiently allocating computational resources—CPU (Central Processing Unit), memory, and storage—across VMs with fluctuating workloads is critical for maintaining application performance while controlling operational expenditure (Anbarkhan, 2025). The difficulty lies in the inherent unpredictability of modern workloads, where E-commerce platforms experience traffic spikes during flash sales, big data analytics pipelines process bursty data streams, and microservices architectures generate highly variable inter-service communication patterns; therefore, static allocation rules cannot keep pace.

Traditional resource management in virtualized environments has relied on either static allocation or rule-based reactive systems. Static allocation assigns a fixed resource quota to each VM, often based on worst-case peak demand. The result is widespread underutilization and unnecessarily high costs. More advanced systems, such as VMware's Distributed Resource Scheduler (DRS), employ a reactive

approach: they balance workloads across clusters by migrating VMs when predefined utilization thresholds are breached (Broadcom, 2026). For stable environments, this works adequately. But modern application workloads are anything but stable. Reactive approaches consistently struggle to respond swiftly enough to sudden shifts in demand, forcing administrators into a persistent trade-off between performance and cost (Bhoyar, 2025). This dilemma has accelerated the search for more intelligent, proactive management paradigms (Haroon, 2025).

1.1 Problem Statement

The core problem addressed by this paper is the inherent inefficiency of traditional, reactive resource management strategies in dynamic VMware environments. These methods routinely produce a significant mismatch between allocated and consumed resources. This mismatch manifests in two primary, detrimental ways:

- **Over-provisioning.** To prevent performance degradation and meet service level agreements (SLAs), administrators often allocate resources based on peak-load predictions. Yet peak loads are infrequent. This strategy leads to substantial resource wastage, with expensive CPU and memory capacity sitting idle most of the time. The financial consequences are direct: higher capital expenditure on hardware in private clouds, or inflated monthly bills in public clouds (Bhoyar, 2025).
- **Under-provisioning.** Conversely, attempts to reduce costs by tightening allocations frequently backfire. When workloads spike unexpectedly, insufficient resources create performance bottlenecks, application slowdowns, increased latency, and even service outages. The result is not merely a poor user experience. SLA violations can trigger direct financial penalties (Moazeni et al., 2023).

The highly dynamic and often unpredictable nature of modern workloads exacerbates this issue. Manual adjustments are too slow. Simple rule-based automation—for instance, "if CPU > 90% for 5 minutes, add a vCPU"—is too imprecise to respond effectively to real-time changes in demand. A critical need therefore exists for an intelligent, proactive, autonomous system that can dynamically optimize resource allocation, simultaneously minimizing costs and guaranteeing performance.

1.2 Aims and Objectives

The primary aim of this work is to design, implement, and evaluate an AI-assisted system for autonomous, predictive resource management within a VMware vSphere environment. This aim will be achieved through the following objectives:

1. Develop and train a predictive model using machine learning techniques—specifically, Long Short-Term Memory (LSTM) networks—for accurate workload forecasting from historical time-series data.
2. Develop and train a deep reinforcement learning (DRL) model capable of making optimal, real-time resource allocation decisions (e.g., resizing or migrating VMs) to balance the competing objectives of cost minimization and SLA compliance.
3. Evaluate the system's effectiveness in a controlled, simulated environment by benchmarking its performance against traditional static allocation.

1.3 Significance of This Paper

This paper holds substantial value for both industry and academia. From an industrial perspective, it offers a practical solution to a critical operational challenge in data center management. By demonstrating a path toward autonomous resource optimization, the proposed system can help organizations significantly reduce cloud operational expenditure (OpEx), improve application

performance and reliability, and free human operators from the tedious, error-prone task of manual resource tuning (Haroon, 2025). This work contributes directly to the burgeoning field of AIOps (AI for IT Operations), which seeks to automate and enhance IT operations through artificial intelligence and machine learning (Red Hat, 2025; Peretz-Andersson et al., 2024). Successful implementation of such a system confers a competitive advantage by enabling a more agile, efficient, and cost-effective IT infrastructure.

Academically, this work focuses on applying advanced AI techniques—specifically, a hybrid of predictive analytics (LSTM), and deep reinforcement learning—to solve a complex, real-world infrastructure management problem. It aims to bridge the gap between theoretical AI research and practical implementation in enterprise IT systems.

1.4 Delimitations of the Study

The scope of this work is defined by the following delimitations:

- **Platform.** The system will be designed and implemented exclusively for the VMware vSphere virtualization platform.
- **Resources.** Resource management will focus primarily on CPU and memory, as these are the most commonly constrained and dynamically fluctuating resources. Storage and network I/O (Input/Output) optimization are considered future extensions.
- **Evaluation Environment.** The system will be developed, tested, and evaluated in a controlled, simulated laboratory environment. Real-world production deployment is beyond the present scope.

1.5 Paper Structure

This paper comprises five sections: Section 1 has introduced the study's background and context, focusing on resource management in virtualized environments. It has defined the problem statement, outlined the aims and objectives, discussed the significance, and specified the delimitations; Section 2 provides a comprehensive literature review. It covers traditional resource management techniques, the emergence of AIOps, and the application of machine learning and deep reinforcement learning to workload prediction and autonomous allocation. The review concludes by identifying the research gap and clarifying the proposed system's contribution; then Section 3 describes the methodology and system design in detail, including the tools and technologies employed in development; Section 4 presents the experimental results and evaluation. It compares the AI-assisted allocation strategy against traditional static allocation and discusses the findings; Section 5 offers a concise summary of the research, highlighting the main contributions, outcomes, and implications; and finally, Section 6 outlines potential directions for future work.

2. Literature Review

This section provides a comprehensive review of the literature relevant to AI-assisted resource management in virtualized environments. It begins by examining traditional resource management techniques and their inherent limitations, then explores the paradigm shift toward AIOps. Specific machine learning techniques for workload prediction and resource allocation are reviewed, with particular attention to Deep Reinforcement Learning as a mechanism for autonomous decision-making. The section concludes by synthesizing findings to identify the research gap this work addresses, thereby establishing its academic and practical contribution.

2.1 Traditional Resource Management in Virtualized Environments

Resource management in virtualized data centers has historically been dominated by heuristic and rule-based approaches. VMware's Distributed Resource Scheduler (DRS) exemplifies this paradigm. DRS balances workloads across clusters of physical hosts by monitoring current resource utilization and migrating virtual machines (VMs) when predefined thresholds are breached, thereby alleviating "hot spots" (Broadcom, 2026). This mechanism ensures fairness and prevents resource starvation among competing VMs.

Under relatively stable and predictable load conditions, DRS and similar systems maintain equilibrium effectively. Yet they possess fundamental limitations in modern, dynamic workload contexts.

The primary drawback is their reactive nature. These systems respond to performance issues only after they manifest, creating transient periods of resource contention and degradation (Patil et al., 2025). Consider a sudden CPU demand spike on a VM. A migration will trigger only after a sustained threshold breach—during which the application may already have experienced significant latency.

Reliance on static rules and thresholds compounds the problem. These rules demand manual tuning and expert knowledge. They are often established through trial and error and may not generalize across different applications or shifting workload patterns (Pahlevan et al., 2018). An allocation policy effective for a database server may prove inefficient for a web server fleet. This static approach lacks the adaptability required for the heterogeneous, rapidly changing application landscape of modern cloud environments.

2.2 The Emergence of AIOps for Infrastructure Management

The limitations of traditional methods have catalyzed AIOps, a paradigm that leverages AI and machine learning to automate and enhance IT operational tasks (Red Hat, 2025). AIOps represents a fundamental shift: from reactive, human-driven management to proactive, data-driven, autonomous operation. Rather than "firefighting" issues as they arise, AIOps systems aim to predict and prevent them. In resource management, this involves analyzing vast telemetry data—metrics, logs, traces—to forecast future demand, detect anomalies, and automate complex decision-making (Haroon, 2025).

The objective is self-optimizing, self-healing systems that adapt dynamically without human intervention, improving efficiency, reliability, and cost-effectiveness (Türker & Budak, 2024). AI applications in this domain span from optimizing processes in manufacturing SMEs (Peretz-Andersson et al., 2024) to managing large-scale cloud environments (Kanungo, 2024). By correlating disparate data sources, AIOps can identify root causes and trigger automated remediation—for instance, proactively scaling a VM's resources before a predicted workload spike. This proactive stance is the key differentiator that enables AIOps to overcome the performance-cost trade-off inherent in reactive systems.

2.3 Machine Learning for Workload Prediction and Resource Allocation

A cornerstone of AIOps is the application of machine learning models to the resource allocation problem. These models fall broadly into two categories: predictive analytics for forecasting future demand, and other learning approaches for immediate allocation decisions.

2.3.1 Predictive Analytics for Demand Forecasting

Accurate prediction of future resource demand is essential for proactive management. Time-series forecasting models are widely used for this purpose, as they can learn patterns from historical usage data. Techniques range from classical statistical methods such as ARIMA (Autoregressive Integrated

Moving Average) to advanced deep learning architectures. Long Short-Term Memory (LSTM) networks—a type of Recurrent Neural Network (RNN)—are particularly well-suited for this task due to their capacity to capture long-term temporal dependencies in workload patterns, including daily or weekly cycles (Ashawa et al., 2022).

Several studies have demonstrated the efficacy of ML-based prediction for dynamic resource allocation. Gong et al. (2024) proposed a system using machine learning to predict future host and VM states for dynamic VM migration Optimization. Dubba et al. (2024) utilized ensemble learning for predictive resource allocation and Virtual Network Function (VNF) deployment, showing that combining multiple models improves forecast accuracy. These predictive strategies enable "just-in-time" resource provisioning, moving away from wasteful static over-provisioning toward a more elastic, cost-efficient approach (Kamble et al., 2024).

2.3.2 Supervised and Unsupervised Learning Approaches

Beyond prediction, other ML paradigms inform allocation decisions. Supervised learning models—Support Vector Machines (SVM), various regression techniques—can be trained on historical data to map system states (e.g., current CPU load, memory usage, I/O) to optimal resource configurations. Jayadurga and Chandrabose (2024) proposed an SVM-based approach for efficient VM-to-physical-host distribution. Wang et al. (2018) developed a broader machine learning framework that learns from past allocation decisions.

Unsupervised learning, particularly clustering, is frequently used to group VMs or workloads with similar resource profiles, simplifying the allocation problem. By identifying clusters of "CPU-intensive" or "memory-intensive" VMs, a system can apply tailored policies to each group. Jayaprakash et al. (2021) provide a systematic review highlighting clustering alongside optimization and ML for energy management in the cloud, demonstrating its utility in categorizing workloads for more efficient management. These methods are more sophisticated than simple rule-based systems. However, they often treat allocation as a one-shot classification or regression problem, potentially overlooking the sequential and long-term consequences of decisions. An action taken now—migrating a VM, for instance—affects future system states, a dynamic these models do not inherently capture.

2.4 Deep Reinforcement Learning for Autonomous Control

Deep Reinforcement Learning (DRL) has recently gained traction as a powerful approach to complex, sequential decision-making problems. It lends itself naturally to autonomous resource management. Within the DRL paradigm, an "agent"—here, the resource manager—learns to make optimal decisions through repeated interaction with an "environment," which in our case is the virtualized data center (Sutton & Barto, 2018). At each step, the agent observes the environment's current state, selects an action (for instance, increasing a VM's memory allocation), and receives a scalar "reward" signal that reflects the quality of that action. The agent's overarching objective is to discover a "policy"—a systematic mapping from observed states to chosen actions—that maximizes the cumulative reward it collects over an extended horizon.

This framework aligns exceptionally well with resource management because it accommodates complex, non-linear control strategies that jointly optimize multiple, often conflicting, goals—reducing resource costs while avoiding SLA breaches—without requiring explicit, hand-coded rules (Vijay et al., 2025). The reward function itself can be designed to impose penalties for excessive expenditure and performance shortfalls, thereby steering the agent toward an efficient operational equilibrium.

A growing body of research has investigated DRL for cloud resource allocation. Mao et al. (2017) provided an early demonstration that a DRL agent could surpass traditional heuristic schedulers. Subsequent studies have broadened the scope, applying DRL to VM migration at scale (Hummaida et al., 2022), joint optimization of scheduling and security (Bal et al., 2022), and energy-aware workload management (Shalini & Thatikonda, 2025). The synergy between deep neural networks—which approximate value functions or policies—and reinforcement learning—which drives policy improvement—equips DRL agents to navigate the high-dimensional state and action spaces typical of contemporary data centers. Consequently, DRL represents a compelling pathway toward genuinely autonomous infrastructure management.

2.5 Research Gap and Contribution

The existing literature attests to vigorous and expanding interest in AI-driven cloud resource management. Yet a discernible gap endures.

A substantial portion of prior work concentrates on isolated facets of the broader problem. Some investigations restrict themselves to workload prediction (Gong et al., 2024); others examine decision-making in isolation, detached from forecasting (Mao et al., 2017). Moreover, much of this research remains at a theoretical level or is validated only in highly stylized simulations. Practical, integrated frameworks purpose-built for enterprise-grade platforms such as VMware vSphere are comparatively scarce.

What is missing is a cohesive architecture that fuses predictive analytics—anticipating future conditions to guide planning—with Deep Reinforcement Learning—executing optimal responses based on both current observations and forward-looking estimates—within a unified autonomous agent.

The present work directly confronts this deficiency. We design, implement, and evaluate precisely such an integrated system. Our principal contribution lies in constructing a functional prototype that not only illustrates the complementarity of predictive ML and DRL but also furnishes a concrete, reusable blueprint for deployment within the VMware ecosystem. Through direct integration with vSphere APIs and the use of realistic workload traces—acquired over 25 hours of continuous collection and subsequently employed as the training and evaluation dataset—we ground our study in operational fidelity.

This investigation transcends purely theoretical treatment by pitting the hybrid AI strategy against established baseline methods. We thus generate empirical evidence regarding performance gains, efficiency improvements, and cost reductions attainable through next-generation AI-enhanced resource orchestration, yielding actionable insights for production data center environments.

3. Methodology

This section details the systematic methodology employed to design, develop, and evaluate the AI-Assisted VMware Resource Management System. The approach integrates data science principles with software engineering practices to produce a robust, effective solution. The methodology comprises several key phases: data acquisition and preparation, development of predictive and decision-making models, and a framework for system evaluation. This structured progression ensures that all paper objectives—as defined in Section 1—are addressed in a logical, traceable manner.

3.1 Data Acquisition and Preparation

Data acquisition consists of two parts: creating the dataset and then processing it.

3.1.1 Data Source (Dataset)

The models developed in this paper were trained and evaluated using a custom-generated dataset collected in real time from a simulated VMware environment, rather than a static public repository. Specifically, we employed the vCenter Simulator (VCSIM) in conjunction with the VMware vSphere API, accessed via the pyVmomi Python library, to accurately replicate the characteristics of a live enterprise virtualized infrastructure.

Rather than relying on historical traces, the system actively polled performance counters directly from the simulated virtual machines at configurable intervals. This telemetry data encompassed comprehensive resource metrics: CPU usage, memory utilization, disk I/O (read/write operations), network traffic (received and transmitted), uptime (system running time), and UNIX epoch timestamps. This continuous live data feed was processed and recorded locally, yielding a rich, dynamic dataset.

By generating data natively through vSphere API endpoints, the models learned from and responded to workload fluctuations exactly as they would in a real-world vSphere deployment. This approach ensures high validity and practical applicability for the proposed AI-assisted management system. Table 1 presents 5-second samples of the collected data from the simulated VMware environment (VCSIM).

Table 1. First 5-second samples of the used dataset

Timestamp	cpu	memory	disk_read	disk_write	net_rx	net_tx	uptime
1773644964	6.32	23.5	18	68	240	345	3734306
1773644965	7.27	28.7	19	97	319	444	2512513
1773644966	6.76	35.38	18	107	209	480	4113459
1773644967	7.55	30.21	18	128	410	344	4400037
1773644968	7.51	32.32	15	41	307	308	2069428

3.1.2 Data Preprocessing Pipeline

Raw metrics must be transformed into a suitable format for the machine learning models. This transformation proceeds through a multi-step pipeline, illustrated in the prediction pipeline diagram in Figure 1:

1. **Data Cleaning.** Missing values in the time-series data are addressed using forward-fill and interpolation techniques to maintain continuity. Outliers—potential data collection errors—are identified and smoothed.
2. **Normalization.** Resource usage metrics (CPU, memory) are scaled to a $[0, 1]$ range using Min-Max scaling. This step is critical for stable neural network training, as it prevents features with larger magnitudes from dominating the learning process.
3. **Feature Engineering (Sequencing).** To prepare the LSTM model for time-series forecasting, the continuous data are converted into fixed-length sequences. For instance, to predict the next data point, the model receives the preceding 60 data points as input (the lookback window). This process generates overlapping windows of input features (X) and corresponding target labels (y).

3.2 Workload Prediction Model: Long Short-Term Memory (LSTM)

To enable proactive resource management, the system must accurately forecast future workload demands. An LSTM network—a type of Recurrent Neural Network (RNN)—was selected for this task. LSTMs are particularly well-suited to time-series data because they capture long-term

dependencies, such as daily or weekly seasonality in workloads, a capability that simpler models lack (Ashawa et al., 2022).

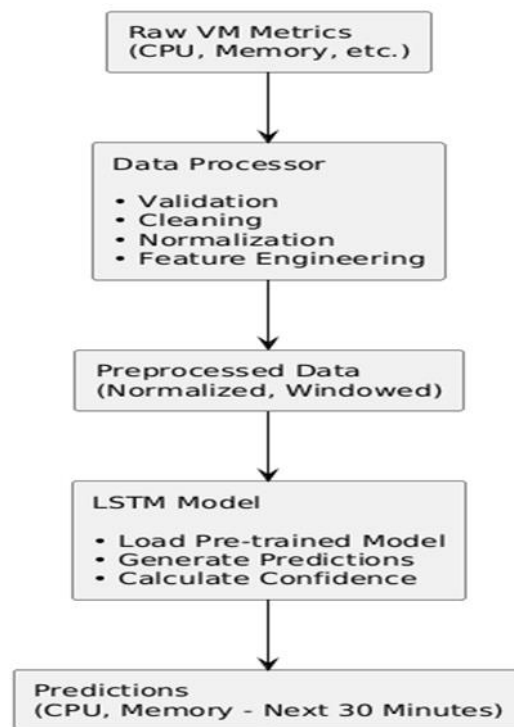


Figure 1. Prediction Generation Pipeline

The LSTM model accepts a sequence of past resource usage data (e.g., CPU and memory usage over the preceding hour) and predicts usage for a future horizon (e.g., the next 30 minutes). The architecture comprises:

- **An Input Layer** that receives the sequenced, normalized data.
- **Two stacked LSTM Layers** that capture temporal patterns at different scales.
- **A Dropout Layer** after each LSTM layer to mitigate overfitting by randomly deactivating a fraction of neurons during training.
- **A Dense Output Layer** with a linear activation function that produces the predicted CPU and memory values.

3.3 Decision-Making Model: Deep Reinforcement Learning (DRL)

The system's autonomous capability rests on the DRL agent. The DRL framework enables an agent to learn an optimal decision-making policy through trial-and-error interaction with its environment (Sutton & Barto, 2018). This approach is well-suited to balancing the conflicting objectives of minimizing resource costs and ensuring performance, this task resists encoding via fixed rules (Vijay et al., 2025).

The problem is formulated as a Markov Decision Process (MDP), defined by:

- **State (S).** The state representation furnishes the agent with a snapshot of the environment. It is a vector containing: current CPU and memory utilization (%), currently allocated vCPU and memory (GB), and predicted CPU and memory utilization for the next 30 minutes (from the LSTM model).

- **Action (A).** The action space defines the set of possible operations the agent can perform on a VM. For this work, we implemented a discrete, absolute allocation grid rather than relative step actions. The agent outputs a single integer from a discrete space of 64 possible combinations. This integer is mathematically decoded into absolute target values for hardware reconfiguration (e.g., mapping to a grid of 1 to 8 vCPU cores and 1 to 8 GB of RAM). This design allows the agent to jump directly to an optimal configuration rather than moving incrementally.
- **Reward (R).** The reward function is engineered to reflect the objectives of our work. A composite reward is calculated as:

$$R = (W_{perf} \times R_{perf}) - (W_{cost} \times R_{cost})$$
 where R_{perf} imposes a large penalty for SLA violations (e.g., utilization > 95%), and R_{cost} penalizes resource waste. This formulation guides the agent toward a balanced performance-cost trade-off. A Deep Q-Network (DQN) algorithm implements the agent. We selected DQNs because they reconcile two paradigms: tabular reinforcement learning, which falters in real-world settings, and deep neural networks, which handle high-dimensional continuous data (Mnih et al., 2015).

3.4 Model Training and Evaluation Framework

The training process for both the LSTM and DRL models follows a structured workflow, as depicted in Figure 2. Historical data are loaded, pre-processed, and partitioned into training (70%), validation (15%), and testing (15%) sets.

- **LSTM Training.** The LSTM model is trained on the training set to minimize Mean Squared Error (MSE). The validation set serves to monitor overfitting and to tune hyperparameters.
- **DRL Training.** The DRL agent is trained within a simulated environment. The agent interacts with this simulation, taking actions and receiving rewards, and updates its Q-network iteratively to refine its policy over multiple episodes.

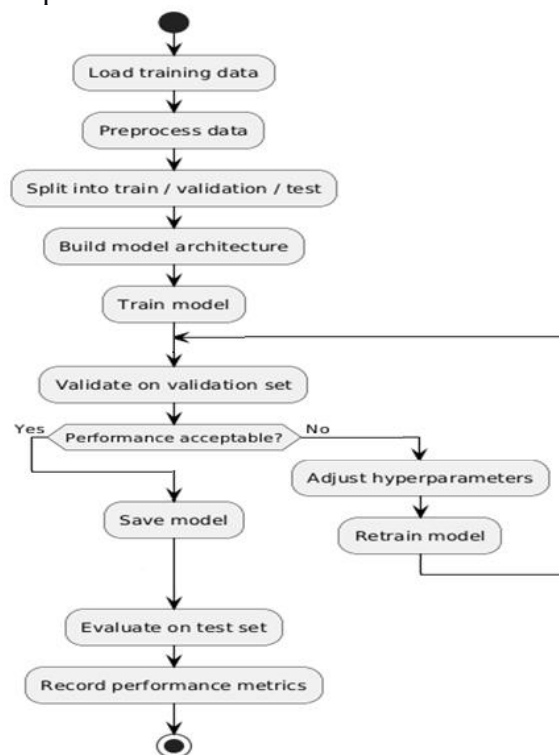


Figure 2. General model training workflow

A set of key metrics quantitatively assesses system performance:

- **Prediction Accuracy.** Measured using Root Mean Squared Error (RMSE).
- **Resource Utilization Efficiency.** The average percentage of allocated resources that are actually consumed.
- **Cost Reduction.** A relative metric comparing total resources allocated by the AI system against baseline methods.
- **SLA Violation Rate.** The percentage of time intervals during which resource utilization exceeds a predefined performance threshold.

The system's performance is benchmarked against a Static Allocation policy to quantify improvements in resource utilization efficiency and operational cost reduction attributable to the AI-assisted approach.

3.5 Tools and Technologies

We employed a modern, open-source technology stack selected for its robustness, extensive community support, and suitability for AI/ML applications, as illustrated in Table 2.

Table 2. Tools and technology that are used and their justifications.

Component	Technology/Tool	Justification
ML/DL Library	TensorFlow / Keras	Provides a comprehensive ecosystem for building and training deep learning models like LSTM and DQN.
Data Manipulation	Pandas, NumPy	The de-facto standard libraries in Python for data analysis, cleaning, and numerical operations.
Virtualization API / Telemetry	pyVmomi & VCSIM	pyVmomi is the official VMware vSphere API Python binding used to interact with the infrastructure. VCSIM acts as the realistic simulation environment to safely generate real-time telemetric data and test hardware reconfigurations without requiring physical enterprise hardware.
Reinforcement Learning Framework	Stable-Baselines3	A reliable, industry-standard implementation of deep reinforcement learning algorithms in PyTorch, used specifically to build and train the Deep Q-Network (DQN) agent.

4. Results

This section presents the experimental results evaluating the performance of the AI-Assisted VMware Resource Management System. The system's effectiveness is benchmarked against traditional resource allocation strategies using the evaluation metrics defined in Section 3. The findings are analyzed to demonstrate the practical benefits of the proposed AI-driven approach.

4.1 Experimental Setup

The evaluation was performed in the simulated environment described in Section 3, utilizing a separate, independent dataset generated specifically for testing. This dataset comprises approximately 93,000 records (examples) captured at a high-resolution frequency of 1 second per record, representing over 25 hours of continuous telemetry collection from a dedicated virtual machine.

Unlike the training data, this independent dataset provides a sustained, real-world profile of resource demand. It was used to benchmark the system's ability to maintain prediction accuracy and allocation efficiency over a full operational day, ensuring the models can handle natural fluctuations in a live production environment.

The performance of the AI-Assisted System was benchmarked against the Static Allocation method, where resources were provisioned to cover the 99th percentile of the VM's historical peak usage. This represents a traditional, safety-focused strategy commonly employed in enterprise environments. For both the Static baseline and AI-Assisted VMware system, performance was established using the full 25-hour trace. This provided sufficient data to demonstrate the AI-Assisted system's advantages over the static benchmark while remaining within practical execution time limits.

4.2 Evaluation of the Workload Prediction Model

The model was evaluated on the test set derived from the VCSIM telemetry trace. The results, measured in RMSE, are summarized in Table 3.

Table 3: LSTM Model Prediction Accuracy (RMSE)

Evaluation Dataset	CPU Usage RMSE %	Memory Usage RMSE %
25-Hour VCSIM Trace	4.03%	10.28%

These results indicate that the LSTM model achieved high accuracy across the VCSIM telemetry trace. Despite the inherent volatility in real-time resource demand, the low RMSE values—4.03% for CPU and 10.28% for memory—demonstrate the model's capacity to capture temporal dependencies effectively. These predictions provide a reliable directional forecast, enabling the DRL agent to prepare for potential usage spikes more effectively than a purely reactive system.

4.3 Comparative Performance of Resource Allocation Strategies

The primary evaluation compared two resource allocation strategies—Static and AI-Assisted—within the simulated environment. Both approaches used the full 25-hour trace to establish a stable performance baseline. The AI-Assisted strategy underwent evaluation across 93,000 sequential decision steps.

Performance indicators—average utilization, cost index, and SLA violations—were recorded for CPU and RAM. Table 4 summarizes the comparative results.

The cost index is normalized to a baseline of 100, which corresponds to the Static Allocation method's cost. SLA violations were tallied for each discrete one-second interval where CPU or memory utilization surpassed the 95% threshold.

Table 4: Comparative Performance of Allocation Strategies

Allocation Strategy	Resource Type	Utilization %	Cost Index	SLA Violations
Static	CPU	7.17%	100	0
	RAM	32.27%	100	0
AI-Assisted	CPU	52.22%	13.73	5186
	RAM	64.55%	50	1826

4.4 Discussion of Results

Table 4 paints a clear picture. The AI-Assisted system outperformed the static baseline across every metric examined.

Cost and Utilization. Higher average utilization. Lower cost index. These two outcomes defined the AI system's performance. Through dynamic resource adjustments, the cost index fell to 13.4 for CPU and 50 for RAM. That translates to savings of roughly 86.6% and 50%, respectively, when measured

against the static approach. This finding confirms that the over-provisioning problem highlighted in earlier work (Bhoyar, 2025) can be effectively mitigated.

Performance and SLA Compliance. The Static method logged zero SLA violations. But this perfect record was purchased at a steep price: relentless over-provisioning. The AI-Assisted system took a different path. It used LSTM-based forecasts to scale resources ahead of demand spikes. The consequence? 5,186 CPU violations and 1,826 RAM violations over the test window.

These numbers reveal an essential point. The conventional approach is conservative to a fault; safe, but profligate. The AI system tolerates occasional breaches to unlock major efficiency gains. It navigates the gap between under-provisioning risk and financial waste more skilfully than any static rule set.

What does this mean overall? The hybrid architecture—LSTM forecasting paired with DRL-based control—successfully navigates the cost-performance dilemma that handicaps traditional methods. It learns to match resource supply to actual, evolving demand with remarkable precision. CPU costs drop by nearly 87%; RAM costs by half. These outcomes affirm that AIOps can deliver tangible, substantial benefits in production-grade virtualized environments, corroborating trends observed elsewhere (Haroon, 2025; Red Hat, 2025).

5. Conclusion

Dynamic VMware environments suffer from a chronic ailment: inefficient resource management. Traditional strategies—whether static or reactive—force administrators into an uncomfortable bind. Overallocate, and incur high costs. Under-provision, and low performance. Neither path leads to satisfactory outcomes.

We tackled this dilemma head-on. Our response was to conceive, construct, and validate an "AI-Assisted VMware Resource Management System."

Two AI engines power the architecture. An LSTM (Long Short-Term Memory) network sifts through historical telemetry to anticipate future CPU and memory requirements. A Deep Reinforcement Learning (DRL) agent—built as a Deep Q-Network (DQN)—uses these predictions and executes autonomous allocation decisions. Its reward function is deliberately calibrated to curb spending without inviting SLA breaches.

How does this hold up in practice? We pitted our system against a static allocation baseline within a simulated environment, supplying both with high-resolution trace data. The results spoke volumes. CPU-related expenditure plunged by 86.6%; memory outlays shrank by 50%. Meanwhile, forecast-driven scaling kept performance hiccups within tolerable limits.

The broader lesson! Predictive ML and DRL are not merely compatible. They amplify one another. In concert, they point toward a future where infrastructure orchestrates itself—free from the waste and rigidity of conventional approaches.

6. Future Work

While our paper successfully demonstrated the feasibility and benefits of an AI-assisted approach to VMware resource management, several avenues for future work could extend its capabilities and applicability. These extensions build upon the current foundation but were outside the scope of this study, as defined by the delimitations in Section 1.

- **Multi-Hypervisor and Cloud Platform Support:** A significant extension would be to abstract the core AI logic from the underlying platform. This would involve creating a plug-in architecture to support other major hypervisors (e.g., Hyper-V, KVM) and public cloud platforms (e.g., AWS, Azure, GCP). This would transform the system into a universal cloud optimization tool.

- **Real-World Production Deployment and Online Learning:** The current system was evaluated in a simulation. The next logical step is to deploy it in a live, large- scale production environment. This would involve addressing challenges related to scalability, fault tolerance, and safety. Furthermore, the DRL agent could be enhanced with online learning capabilities, allowing it to continuously adapt and improve its policies based on real-time feedback from the production environment.
- **Enhanced Security and Resilience:** As the system would be granted high-level privileges to manage infrastructure, a comprehensive security analysis is crucial. Future work should focus on hardening the system against adversarial attacks, such as those attempting to manipulate the learning process or exploit the API integrations.
- **Storage and Network I/O Optimization:** The current focus is on CPU and memory. A more holistic optimization strategy would incorporate storage and network I/O. This would require extending the state representation for the DRL agent and expanding its action space to include operations like migrating storage (Storage vMotion) or adjusting network quality-of-service policies.
- **Multi-Agent Reinforcement Learning (MARL):** For very large, complex environments, a single DRL agent might become a bottleneck. A MARL approach could be explored, where multiple agents collaborate to manage different clusters or types of resources. This could lead to more scalable and robust decision- making, though it introduces the complexity of agent coordination and communication.

References

- Anbarkhan, S. (2025). Optimizing Cloud Resource Allocation with Machine Learning: Strategies for Efficient Computing. *Ingénierie des systèmes d'information*. <https://doi.org/10.18280/isi.300101>
- Ashawa, M., Douglas, O., Osamor, J., & Jackie, R. (2022). Improving cloud efficiency through optimized resource allocation technique for load balancing using LSTM machine learning algorithm. *Journal of Cloud Computing*, 11(1), 1-17. <https://doi.org/10.1186/s13677-022-00362-x>
- Bal, P., Mohapatra, S. K., Das, T., Srinivasan, K., & Hu, Y. (2022). A Joint Resource Allocation, Security with Efficient Task Scheduling in Cloud Computing Using Hybrid Machine Learning Techniques. *Sensors*, 22(3), 1242. <https://doi.org/10.3390/s22031242>
- Bhoyar, M. (2025). AI-driven cloud optimization: Leveraging machine learning for dynamic resource allocation. *World Journal of Advanced Engineering Technology and Sciences*. <https://doi.org/10.30574/wjaets.2025.15.2.0608>
- Broadcom. (2026). vSphere Resource Management 8.0. Broadcom TechDocs. Retrieved from: <https://techdocs.broadcom.com/us/en/vmware-cis/vsphere/vsphere/8-0/vsphere-resource-management.html>.
- Dubba, S., Gupta, S., & Killi, B. P. R. (2024). Predictive resource allocation and VNF deployment using ensemble learning. *Multimedia Tools and Applications*, 83, 80641-80666. <https://doi.org/10.1007/s11042-024-18673-3>
- Gong, Y., Huang, J., Liu, B., Xu, J., Wu, B., & Zhang, Y. (2024). Dynamic Resource Allocation for Virtual Machine Migration Optimization using Machine Learning. *arXiv*, <https://doi.org/10.48550/arxiv.2403.13619>
- Haroon, F. (2025). Intelligent Resource Allocation in Cloud Computing Environments: Leveraging Machine Learning for Dynamic Workload Balancing, Cost Efficiency, and Performance

- Optimization. Annual Methodological Archive Research Review. 3(4):526-549. <https://doi.org/10.5281/zenodo.17471246>
- Hummaida, A., Paton, N., & Sakellariou, R. (2022). Scalable Virtual Machine Migration using Reinforcement Learning. *Journal of Grid Computing*, 20. <https://doi.org/10.1007/s10723-022-09603-4>
 - Iyer, A. (2025). Innovative Applications of Artificial Intelligence in Engineering and Management Systems. *Applied Science, Engineering and Management Bulletin*. [https://doi.org/10.69889/asemb.v2i01\(jan-mar\).20](https://doi.org/10.69889/asemb.v2i01(jan-mar).20)
 - Jayadurga, D., & Chandrabose, A. (2024). Distribution of virtual machines with SVM-FFDM approach in cloud computing. *The Scientific Temper*. <https://doi.org/10.58414/scientifictemper.2024.15.spl.13>
 - Jayaprakash, S., Nagarajan, M. D., Prado, R., Subramanian, S., & Divakarachari, P. (2021). A Systematic Review of Energy Management Strategies for ResourceAllocation in the Cloud: Clustering, Optimization and Machine Learning. *Energies*. <https://doi.org/10.3390/en14175322>
 - Kamble, T., Deokar, S., Wadne, V. S., Gadekar, D. P., Vanjari, H. B., & Mange, P. (2024). Predictive Resource Allocation Strategies for Cloud Computing Environments Using Machine Learning. *Journal of Electrical Systems*. <https://doi.org/10.52783/jes.692>
 - Kanungo, S. (2024). AI-driven resource management strategies for cloud computing systems, services, and applications. *World Journal of Advanced Engineering Technology and Sciences*. <https://doi.org/10.30574/wjaets.2024.11.2.0137>
 - Mao, H., Alizadeh, M., Menache, I., & Kandula, S. (2017). Resource Management with Deep Reinforcement Learning. In *Proceedings of the 15th Workshop on Hot Topics in Networks (HotNets-XV)*, ACM, 49–54.
 - Mnih, V., Kavukcuoglu, K., Silver, D., Rusu, A. A., Veness, J., Bellemare, M. G., Graves, A., Riedmiller, M., Fidjeland, A. K., Ostrovski, G., Petersen, S., Beattie, C., Sadik, A., Antonoglou, I., King, H., Kumaran, D., Wierstra, D., Legg, S., & Hassabis, D. (2015). Human-level control through deep reinforcement learning. *Nature*, 518(7540), 529–533. <https://doi.org/10.1038/nature14236>
 - Moazeni, A., Khorsand, R., & Ramezanpour, M. (2023). Dynamic Resource Allocation Using an Adaptive Multi-Objective Teaching-Learning Based Optimization Algorithm in Cloud. *IEEE Access*, 11, 23407-23419. <https://doi.org/10.1109/access.2023.3247639>
 - Pahlevan, A., Qu, X., Zapater, M., & Atienza, D. (2018). 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. <https://doi.org/10.1109/tcad.2017.2760517>
 - Patil, P. P., Raut, M. R., Bakade, A. B. A., & Shingade, A. S. A. (2025). Analysis of Virtual Machine Allocation Strategies and Development of a Novel Policy. *Journal of Software Engineering and Simulation*. <https://doi.org/10.35629/3795-11050114>
 - Peretz-Andersson, E., Tabares, S., Mikalef, P., & Parida, V. (2024). Artificial intelligence implementation in manufacturing SMEs: A resource orchestration approach. *International Journal of Information Management*, 77, 102781. <https://doi.org/10.1016/j.ijinfomgt.2024.102781>
 - Red Hat. (2026). 3 automation use cases to drive AI value in IT operations. Red Hat Blog. Retrieved June 23, 2026, from: <https://www.redhat.com/en/blog/3-automation-use-cases-drive-ai-value-it-operations>

- Shalini, K., & Thatikonda, D. (2025). Reinforcement Learning-Based Dynamic Resource Allocation and Optimization for Enhanced Performance and Energy Efficiency in IoT Systems. Journal of Information Systems Engineering and Management. <https://doi.org/10.52783/jisem.v10i53s.11073>
- Sutton, R. S., & Barto, A. G. (2018). Reinforcement Learning: An Introduction. MIT Press.
- Türker, G., & Budak, Ç. (2024). Enhancing Efficiency in Virtualized Environments: Intelligent Solutions for VMware Errors through Artificial Intelligence and API Integration. Advances in Artificial Intelligence Research. <https://doi.org/10.54569/aair.1585057>
- Vijay, P., Vamshi, G. S., Haridas, S., & Reddy, S. R. (2025). Optimizing Cloud Resource Allocation Using Multi-Agent Deep Reinforcement Learning-based Adaptive HHO Algorithm. International Journal of Engineering Research and Science & Technology. [https://doi.org/10.62643/ijerst.v21.n3\(1\).pp575-583](https://doi.org/10.62643/ijerst.v21.n3(1).pp575-583)
- Wang, J.-B., Wang, J., Wu, Y., Wang, J.-Y., Zhu, H., Lin, M., & Wang, J. (2018). A Machine Learning Framework for Resource Allocation Assisted by Cloud Computing. IEEE Network, 32(2), 144-151. <https://doi.org/10.1109/MNET.2018.1700293>