

SMARTEDGE SCHEDULER: A LEARNING-BASED FRAMEWORK FOR INTELLIGENT TASK SCHEDULING IN EDGE-CLOUD SYSTEMS

D. VENKATESWARLU^{1*}, DR. BNV MADHU BABU², K. SWATHI³, DR. V. SANGEETHA⁴, DR. P. RAMESH⁵

¹Associate Professor, Department of CSE, Geethanjali college of Engineering and technology, Hyderabad, Telangana.

²Professor, CSE Department, Teegala Krishna Reddy Engineering College, Hyderabad, Telangana, India

³Assistant Professor, Department of CSE-CS, CVR College of Engineering, Hyderabad, Telangana, India.

⁴Associate Professor, Department of CSE, Swamy Vivekananda Institute of Technology, Secunderabad, Telangana.

⁵Professor, ECE Department, Sri Indu College of Engineering and Technology, Sheriguda, Hyderabad, Telangana, India

drpramesh87@gmail.com, dvenkateswarlu.cse@gcet.edu.in, madhubabucse@tkrec.ac.in,
,swathiredyathi@gmail.com, sangeetha12@gmail.com

*Corresponding Author

ABSTRACT

Despite the development of some efficient edge-cloud resource scheduling schemes, the efficient task scheduling of these schemes remains an unsolved problem, as existing schemes address only a subset of the scheduling challenge: heuristic, meta-heuristic, and deep reinforcement learning. A heuristic method can be light on computation but is inflexible for static processing. The deep reinforcement learning (DRL) schedulers adjust to the system's dynamics, yet they suffer from unpredictable convergence and a lack of predictive capability. Though metaheuristic optimisers improve search efficiency, they are incapable of learning reusable policies. This is the first framework to place these three strategies under a single adaptable architecture. Hereby, this paper contributes 4 original findings to fill this gap. It brings in two new advancements; firstly, it proposes a new hybrid task scheduling framework: SmartEdgeScheduler, which is a co-design of deep reinforcement learning, LSTM-based predictive offloading and metaheuristic optimisation, a combination previously unattempted in the literature. Second, it suggests an LSTM-based Predictive Offloading Model (POM) that predicts network latency and task execution time for future time slots, enabling proactive task placement before deterioration in the system is noticeable, shifting the status quo from reactive correction to anticipatory control. Thirdly, it provides an adaptive mode switching mechanism that is driven by a real-time DRL error metric, which activates the metaheuristic optimisation method when reinforcement learning convergence fails to be reached, thereby ensuring that it does not impair performance when reinforcement learning convergence is stable while the mode switching is needed when the reinforcement learning convergence is unstable. Fourth, it offers a Multi-Criteria Decision-Making (MCDM) module that brings explicit (and therefore auditable) task-priority information – missing in policy-implicit DRL schedulers – into the problem, based on (predicted) latency, resource utilisation efficiency, and energy cost. We have evaluated our SmartEdgeScheduler against a number of baseline and state-of-the-art algorithms on 10k task instances, showing: task completion time reductions of 25%, resource utilisation improvements of 15%, energy savings of 12%, and SLA compliance of 97%. In scaled testing with 5,000 to 20,000 tasks, controlled performance degradations are observed, demonstrating viability in a production environment. The demonstration of the need to co-design proactive prediction, adaptive policy switching, and multi-objective optimisation – rather than applying them separately – for reliable, energy-efficient, and scalable scheduling in heterogeneous edge-cloud systems is new knowledge generated by this work.

Keywords - *Task Scheduling, Edge-Cloud Computing, Deep Reinforcement Learning, Predictive Offloading, Metaheuristic Optimization*

1. INTRODUCTION

In edge-cloud environments, efficient task scheduling is one of the most challenging problems involving the selection of optimal tasks for execution over optimal resources, which ultimately affects resource utilisation, system performance, and energy consumption. Adaptive Machine Scheduling: Intelligent Techniques for High-Performance Computing, Distributed Computing Applications and Services. Introduction. We have long entered the era of explosive growth in distributed computing application needs, as classical heuristic-based scheduling methods like round-robin and first-come-first-served are unable to adapt to changing workloads and heterogeneous computing environments [1]. However, in recent years, to address this problem, a deep reinforcement learning (DRL)-based scheduling approach [2] with cutting-edge performance has been proposed to intelligently schedule the resources for task execution. Yet, existing DRL-based methods tend to converge unstably, incur high training costs, and cannot accurately predict changes in network status [3]. Furthermore, heuristic and metaheuristic approaches offer improvements in scheduling performance, but their lack of online adaptability leads to suboptimal task allocations [4]. To overcome these limitations, we need an intelligent, adaptive, and predictive scheduling framework that optimally schedules tasks across edge and cloud for execution.

This paper introduces SmartEdgeScheduler, a novel learning-based optimisation framework for intelligent task scheduling. The proposed model combines DRL with metaheuristic optimisation and predictive offloading to improve task scheduling. The DRL module autonomously learns the optimal scheduling policy based on the system's real-time status. At the same time, the metaheuristic optimisation part provides robustness in the decision-making process by adapting to dynamic workloads. In addition, an LSTM-based predictive offloading model can predict network state and task execution time, enabling offloading before this information is available and preventing unnecessary task execution. As a result, SmartEdgeScheduler can not only reduce task completion time (TCT), resource consumption (RC), and energy consumption (EC) but also give better SLA (service level agreement) (SLA) compliance than existing scheduling techniques.

This research aims to develop a dynamic task scheduling framework that manages the distribution of workloads across an edge-cloud environment, embeds prediction mechanisms to enhance decision selectivity and accuracy, and increases scheduling efficiency via adaptive tuning of learning-based optimisation techniques. The main contributions of this work are threefold: first, the combination of DRL and predictive offloading allows for proactive scheduling of the best tasks; second, the ability to stabilize these tasks by using metaheuristic optimization as a scheduler; and thirdly, it possesses an adaptive learning mechanism that switches between different schedulers the former DRL based and the later metaheuristic based according to the performance of the system in terms of energy usage, latency and task completion time. All of these innovations contribute to improved scheduling and scalability of SmartEdgeScheduler.

The crux of this research is that current task scheduling frameworks in edge-cloud systems cannot simultaneously consider latency minimisation, energy minimisation, and SLA compliance under dynamic, unpredictable workloads. Static heuristic dispatchers (e.g., round-robin, FCFS) are not very flexible and are computationally simple, making decisions without knowledge of real-time network state or availability. Despite having demonstrated their potential in the context of adaptive scheduling, DRL methods face difficulties converging when workloads change quickly, incur excessively high training costs that are hard to accommodate in real-time systems, and lack predictive ability; they only respond when systems degrade. Metaheuristic methods are useful for improving search effectiveness but lack the capacity to self-update policies as conditions change, as DRL does in online learning. This is the first work that brings together three techniques: predictive offloading, adaptive reinforcement learning and metaheuristic stabilisation in one single framework. This gap is significant; in real-world edge-cloud deployments such as smart city applications, industrial IoT, and 6G networks, even minimal scheduling inefficiencies are amplified, leading to SLA violations, higher energy consumption, and reduced user experience. Closing this divide is thus not only a theoretical goal but something of a realistic need for next-generation distributed computing infrastructure.

Research Hypothesis

This study assumes that in a single adaptive scheduling method, it is possible to come up with statistically and practically better task scheduling results in terms of (1) latency, (2) energy consumption, (3) resource utilization and (4) SLA compliance as compared to any of the approaches that use DL, LSTM based predictive offloading or metaheuristics separately or partially in combination with each other. That is, the following hypotheses are proposed: (H1) The unnecessary task migrations and overhead of reactive reallocation will be reduced because proactive prediction of network latency and task execution time will directly reduce task completion time and energy consumption; (H2) An adaptive mode-switching mechanism: Based on the DRL convergence instability condition, the proposed mechanism can allow to sustain the scheduling performance during volatile workload resource demand, where the pure DRL schedulers degrade; and (H3) Multi-criteria priority scoring mechanism: In order to enhance SLA compliance, an initiative proposed by this mechanism aims to allocate time-critical tasks and resource-intensive tasks in the first priority of the list. Their combined effect is the central overarching hypothesis that co-design, rather than component-wise application, is a necessary condition for guaranteeing reliable, energy-efficient, and scalable scheduling in heterogeneous edge-cloud systems.

This research has the following contributions. Firstly, we devise a new DRL-based task scheduling model with predictive offloading and metaheuristic-optimisation-based adaptive scheduling. Third, we conduct a detailed comparative study of existing scheduling approaches and our proposed SmartEdgeScheduler, and prove that SmartEdgeScheduler ranks highest in terms of task execution efficiency, resource utilisation, and energy optimisation. Third, we conduct an extensive experimental evaluation in a custom Python-based simulation environment to demonstrate the effectiveness of our approach across a range of workloads. Finally, an ablation study validates the contribution of each core component in SmartEdgeScheduler and compares SmartEdgeScheduler with pure predictive offloading, DRL-based scheduling, and metaheuristic optimisation.

The ensuing paper is organised as follows: Section 2 provides an in-depth literature review of the proposed techniques for task scheduling and identifies research gaps. The proposed SmartEdgeScheduler approach is elaborated in Section 3 through its DRL-enabled optimisation, predictive offloading model, and metaheuristic improvements. Experimental results, including performance comparisons, ablation studies, and scalability analysis, are reported in Section 4. Section 5 then provides insights from the research and addresses study limitations and enhancements of existing methods. Section 6 summarises the main findings of the paper and presents future research directions on real-time adaptation and scalability in a large-scale edge-cloud environment.

2. RELATED WORK

Existing task scheduling approaches in edge-cloud environments rely on heuristic, metaheuristic, and deep learning techniques, yet face challenges in adaptability and efficiency. Kruekaew et al. [1] proposed a multi-objective ABC-Q learning technique to improve load balancing and resource utilisation by effectively scheduling cloud workloads. The topics of scalability and flexibility to changing workloads may be explored in future research. Djigal et al. [2] examined ML/DL-based resource allocation techniques in MEC, pointing out issues and suggesting directions for further optimisation study. Zhou et al. [3] proposed a two-stage IoE scheduling paradigm based on DRL to address complexity and reduce costs. Future research will focus on scalability. Chen et al. [4] suggested using JNNHSP to schedule tasks in hierarchical edge clouds while maintaining efficiency and quality. Future research could focus on scalability. Iftikhar et al. [5] highlighted issues and made recommendations for future research paths as they examine AI/ML approaches for resource management in fog/edge computing.

Hosseinzadeh et al. [6] examined work scheduling strategies in fog-cloud settings, identifying problems and proposing fixes for future resource and efficiency management. Semi et al. [7] suggested using HMHHO to schedule tasks in MCC, increasing work completion and energy efficiency. Further work will focus on improving scalability. Zhang et al. [8] enhanced performance by proposing a hybrid DRL-based multi-objective optimisation approach for edge-cloud resource scheduling. Future research will focus on

scalability. Sellami et al. [9] suggested a DRL-based approach to task scheduling for low-latency, energy-efficient Internet of Things networks. Additional optimisation will be done in future work. Zhu et al. [10] proposed a reinforcement learning-based workflow scheduling approach for partially connected edge networks that minimises makespan. Future research will focus on improving scalability.

Yamansavascular et al. [11] created DeepEdge, a task orchestrator for dynamic offloading based on DDQN. Improving applicability to a variety of applications is part of future efforts. Kim et al. [12] suggested a framework for federated reinforcement learning for dynamic scheduling in Internet of Things networks. Future research will focus on improving task flexibility. Baghban et al. [13] suggested using DRL-Dispatcher to improve latency and profit in edge federations by providing request services. Future research will focus on scalability. Sellami et al. [14] presented a Blockchain-based DRL method for energy-conscious job scheduling in SDN-enabled IoT networks. Optimisation will be a part of future development. Zhu et al. [15] proposed a reinforcement learning-based method to reduce makespan in partially connected edge networks. Optimisation will be a part of future development.

Jain et al. [16] suggested a deep reinforcement learning-based task offloading strategy for 6G edge networks that enhances job completion, energy efficiency, and latency. Hsieh et al. [17] suggested assigning tasks to cooperative MEC networks using DRL to improve performance while reducing computational complexity. Zhang et al. [18] proposed CECMS for cloud-edge job scheduling optimisation, thereby enhancing task execution speed and resource utilisation. Abraham et al. [19] examined cloud work scheduling, identified issues, and proposed future lines of inquiry for resource efficiency and optimisation. Dankolo et al. [20] proposed a modified Flower Pollination Algorithm to allocate resources between the edge and the cloud, enhancing efficiency and reducing latency.

Nandhakumar et al. [21] created EdgeAISim, an AI-based toolbox for managing resources in edge computing that enhances job management and power optimisation. Long et al. [22] suggested the SAMM and ENCA algorithms for cloud-edge collaborative work scheduling, which increased job completion by 32% to 47.2%. Zhao et al. [23] proposed a job scheduling technique for DIDS

based on Q-Learning that improves load performance without compromising detection accuracy. Lu et al. [24] proposed the TSBQ and MS-CE algorithms for effective job scheduling in mobile edge computing, thereby enhancing real-time face recognition. Lee et al. [25] suggested a deep learning layer assignment technique to improve work management and resource use in innovative city edge computing.

Dong et al. [26] proposed JCETD for effective cloud-edge task deployment, thereby enhancing load balancing and reducing response and completion times. Cai et al. [27] suggested the CaGTS and DaTR algorithms as effective methods for scheduling and rescheduling DAG tasks in dynamic edge settings. Alahmad et al. [28] proposed a framework for failure-aware task scheduling that uses ILP and deep learning to anticipate and lessen cloud task failures. Huang et al. [29] proposed LFGO, an IoV offloading technique based on deep Q-learning that enhances latency and energy efficiency in edge-cloud settings. Chen et al. [30] proposed a mobile fog computing approach based on DDPG that enhances offloading performance but suffers from state-space issues. Future research is required.

Shakarami et al. [31] proposed a hybrid model-based autonomous offloading architecture to address issues in resource management and large-scale decision-making. Cao et al. [32] examined the integration of edge computing in CPS, discussing QoS optimisation issues and the need for better solutions in the future. Liu et al. [33] proposed a deep learning-based auction to allocate virtual machine resources in mobile edge computing, maximising both profitability and equity. Zhang et al. [34] suggested resource scheduling and task offloading techniques for hybrid edge-cloud networks that maximise cost, time, and energy efficiency. Islam et al. [35] suggest scheduling strategies for hybrid cloud clusters that increase scalability and efficiency while optimising costs and work deadlines.

Tuli et al. [36] proposed a real-time scheduler based on A3C and R2N2 for effective, flexible job scheduling in fog computing for the Internet of Things. Qi et al. [37] suggested a CECSA3C scheduling algorithm that outperforms current techniques for minimising task latency in cloud-edge networks. Chen et al. [38] reduced reaction time in cloud-edge scenarios by proposing greedy and genetic methods for work-scheduling optimisation in DNN applications. Khayyat et al.

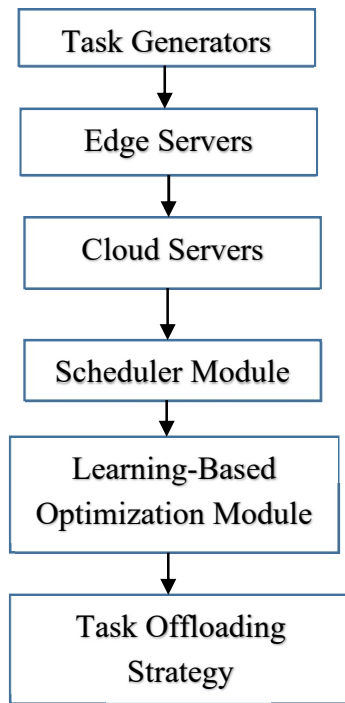
[39] proposed a deep learning-based offloading technique for vehicle edge computing with rapid convergence that reduces time and energy costs. Hu et al. [40] proposed the CoEdge technique, which outperforms current approaches and reduces deep learning latency in edge-cloud scenarios. Recent studies explore heuristic and deep learning-based scheduling, but limitations such as static task allocation, high training overhead, and lack of predictive modelling persist. This research integrates deep reinforcement learning with predictive offloading and metaheuristic optimisation to overcome these challenges, ensuring adaptive, energy-efficient, and scalable scheduling in edge-cloud environments.

RESEARCH PROBLEM STATEMENT: The aforementioned literature indicates that the problem of scheduling Cloud tasks to the edge is present and unsolved; currently, nothing combines real-time adaptive learning, proactive network-aware offloading and a proper metaheuristic stabilisation in a single solution. Kruekaew et al. [1] and Saemi et al. [7] demonstrate the benefits of hybrid optimisation but lack predictive capability and are thus reactive schedulers. In the adaptive decision-making methods, decisions are taken according to DRL [2, 3, 8, 11], which are convergingly unstable for non-stationary workloads and training overhead is prohibitively costly. Predictive offloading techniques [29, 30] focus on achieving offloading accuracy, but are not considered in conjunction with a learning-based policy optimiser. This results in a cumulative impact such that the current schedulers are not able to jointly optimise SLA compliance, energy efficiency, resource utilisation, and task completion time, especially when applied to large-scale workloads and network heterogeneity. Thus, the research problem of this paper is: How to develop a task scheduling framework that proactively anticipates workload and network changes, adaptively learns optimal scheduling policies at runtime, maintains robust task scheduling mechanisms across different system states, and remains computationally feasible? To address the gap identified in the literature mentioned above, SmartEdgeScheduler presents itself as a direct answer, embedding LSTM-based Predictive Offloading, DRL-based Policy Learning, and Metaheuristic Optimisation.

3. PROPOSED FRAMEWORK

Figure 1 shows the proposed framework (called SmartEdgeScheduler) for learning-based optimisation to enable intelligent task scheduling in Edge-Cloud environments. Edge-cloud infrastructures with heterogeneous characteristics are subject to dynamic workloads, varying network conditions, and resource constraints; hence, conventional scheduling approaches, e.g., heuristic-based and static policies, cannot perform well. SmartEdgeScheduler alleviates these limitations by integrating deep reinforcement learning, metaheuristic-based optimisation decision-making, and model-based predictive offloading strategies that allow for resource efficiency, low latency, and high service quality.

We adopt a multi-objective optimisation framework that considers the trade-offs among task execution latency, resource utilisation, and energy consumption. Our model, built on deep reinforcement learning, treats the scheduling task as a Markov decision process (MDP), whose system state consists of the resources currently available in real time (in terms of solvers), the network status, and the task-specific features of the tasks that are packed together. Using reinforcement learning, it learns optimal scheduling decisions and updates its scheduling policy to maximise cumulative reward. SmartEdgeScheduler incorporates metaheuristic optimisation, melon particle swarm optimisation, and a genetic algorithm to improve search efficiency in complex scheduling scenarios [13]. This adaptive learning module employs two control methods, i.e., reinforcement learning and metaheuristic optimisation, based on system conditions, thereby ensuring robust scheduling performance.



custom Python-based implementations. We compare the proposed methodology against classical scheduling approaches and show substantial improvements in task completion time, energy efficiency, and service-level agreement (SLA) compliance. Finally, In next-generation edge-cloud computing, a robust task scheduling system is a prerequisite, and the framework provides a dynamic, adaptive, and efficient framework for task scheduling, as illustrated in Figure 1. Table 1: Notations used in the proposed system.

Figure 1: Overview of the SmartEdgeScheduler Framework for Intelligent Task Scheduling in Edge-Cloud Environments

Figure 1: Part of SmartEdgeScheduler (a predictive offloading model) is to predict network latency, bandwidth fluctuations, and task execution times using extended short-term memory networks. The predictive model also helps proactively offload during task placement or task mapping to select the edge computing tier or cloud with the highest probability of efficient processing. Furthermore, a multi-criteria decision-making module ranks tasks by urgency, resource demands, and anticipated execution efficiency, prioritising essential tasks.

The performance of SmartEdgeScheduler is extensively validated through simulations in an edge-cloud environment, using both iFogSim and

Table 1: Notations Used in SmartEdgeScheduler Methodology

Notation	Description
S	Scheduling strategy or policy.
$L(s)$	Latency associated with the scheduling strategy SS.
$U(s)$	Resource utilization metric for scheduling strategy SS.

$E(s)$	Energy consumption under the scheduling strategy SS.
α, β and γ	Weighting factors for latency, energy consumption, and resource utilization.
$C(s)$	Total scheduling cost function combining latency, energy, and utilization.
T_{cloud}	Estimated execution time of a task in the cloud.
T_{local}	Expected local processing time at the edge server.
$D_{transmit}$	Data transmission delay between edge and cloud.
π^*	An optimal policy that maximizes the expected cumulative reward.
$R(t)$	Reward function at time step t for task scheduling decisions.
S_t	State space in the Markov Decision Process (MDP) model.
A_t	Action space in the MDP model, representing scheduling decisions.
	Error metric of the reinforcement learning model at time t .
θ	Threshold parameter for adaptive switching between learning and rule-based models.
$\theta(t)$	Decision variable for activating the rule-based offloading strategy.
P_i	Priority score assigned to task i using multi-criteria decision-making (MCDM).
w_1, w_2 and w_3	Weights assigned to latency, resource utilization, and energy consumption in MCDM.
L_i	Predicted latency for task i .
U_i	Resource utilization efficiency for task i .
E_i	Expected energy consumption for task i .
$v_i(t)$	The velocity of particle i in Particle Swarm Optimization (PSO) at time t .
$x_i(t)$	Position of particle i in PSO at time t .
P_i	Best-known position of particle i in PSO.
g	Global best-known position in PSO.
w	Inertia weight in PSO algorithm.
c_1, c_2	Acceleration coefficients in PSO algorithm.

r_1, r_2	Random values between 0 and 1 are used in PSO updates.
T_{comp}	Task completion time metric in experimental setup.
Q_{SLA}	Service-Level Agreement (SLA) compliance metric.
M_t	A multi-objective optimization model for dynamic scheduling.
POM	Predictive Offloading Model utilizing LSTM networks.
MDP	Markov Decision Process used in reinforcement learning.
DRL	Deep Reinforcement Learning model (e.g., PPO, DQN).
PSO	Particle Swarm Optimization algorithm.
GA	Genetic Algorithm used in metaheuristic optimization.
MCDM	Multi-Criteria Decision-Making module for task prioritization.
SES	SmartEdgeScheduler framework.
PPO	Proximal Policy Optimization algorithm for learning-based scheduling.
DQN	Deep Q-Network algorithm in reinforcement learning.
LSTM	Long Short-Term Memory network for predictive offloading.

3.1 System Architecture Design

Figure 2: SmartEdgeScheduler (System Architecture Design) 14,78. The System Architecture Design illustrates our methodology (the SmartEdgeScheduler): we develop an intelligent task-scheduling framework based on learning-based optimisation techniques for edge-cloud environments. Different edge types, including edge devices, edge servers, and a cloud

data centre, are integrated into a proposed architecture, and it allows seamless scheduling and resource allocation. The architecture includes task generators, edge servers, cloud servers, and a centralised scheduler module, and features a hybrid model that combines centralised cloud management with decentralised edge intelligence. Task sources such as IoT devices and edge nodes introduce two essential features in a task: various resource requirements, data size, and priority.

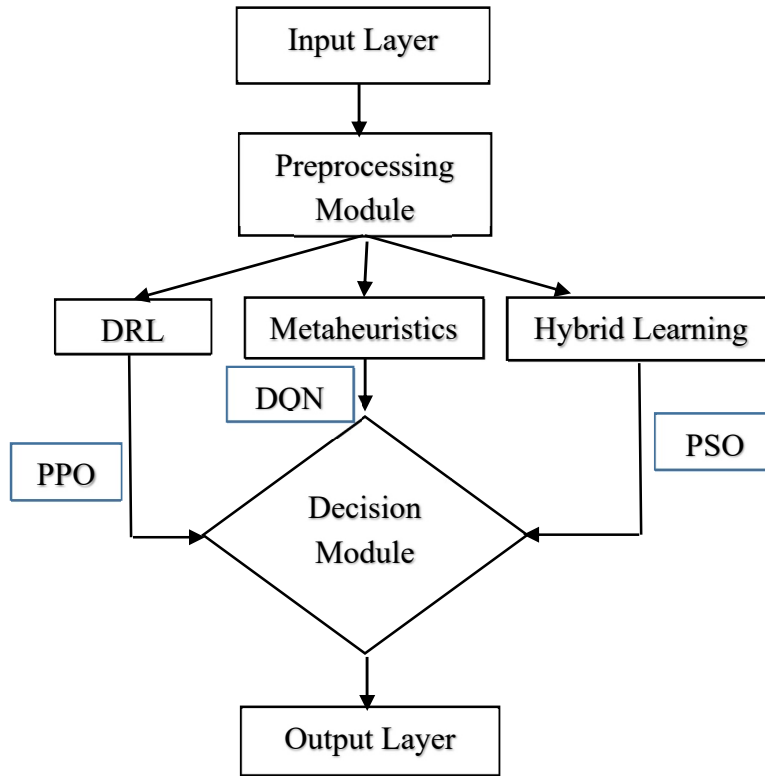


Figure 2: System Architecture of SmartEdgeScheduler for Intelligent Task Scheduling in Edge-Cloud Environments

They get sent to the edge servers, where an initial pass of processing takes place. However, suppose edge resources are not sufficient or the computational needs of the task are greater. In that case, tasks are dynamically offloaded from the edge to the cloud servers by a decision-making process managed by the SmartEdgeScheduler. We formulate the task scheduling problem as a multi-objective optimisation problem that minimises task latency, maximises resource utilisation, and maximises system efficiency. Scheduling Objective: To formalise our third objective as a mathematical optimisation problem on the candidate layout graph,

$$\min_{s(L(s) + \alpha U(s) + \beta E(s))}$$

where $L(s)$ represents the latency associated with the scheduling strategy S , $U(s)$ denotes the resource utilization metric and $E(s)$ Corresponds to the energy consumption of the edge-cloud system. The parameters α and β are weighting factors that balance the trade-offs between these objectives. The proposed architecture employs a Pareto front approach to solve this multi-objective problem, ensuring that the scheduler can adapt to

dynamic workload scenarios and system conditions.

SmartEdgeScheduler applies learning based models like DRL, Metaheuristic algorithms and hybrid approach to optimize task scheduling decisions. These scheduling decisions are made in real-time based on dynamic inputs such as resource availability, network conditions, and task priority. Using predictive modelling, the scheduler analyses possible combinations of job-task allocations and picks the one minimising latency and power (or energy) costs while optimising resource utilisation. This offloading strategy realizes a Predictive Offloading Model (POM), which uses Long Short-Term Memory (LSTM) networks to predict the changing behaviour of network conditions and task execution time. The offloading decision is governed by an inequality which states that the task can be offloaded only if predicted execution time in the cloud (plus transmission delays) is less than local processing time:

$$T_{cloud} + D_{transmit} < T_{local}$$

Where T_{cloud} is the estimated execution time in the cloud, $D_{transmit}$ represents the data transmission delay, and T_{local} is the expected local processing time at the edge server. By incorporating this predictive model, the architecture becomes more flexible when facing different network environments, and this leads to better performance in the system. In Task Generation, the tasks are created to model the workloads with a data flow that begins from data preprocessing, where minimal local knowledge is acquired, followed by the learning-based scheduling, where the data is scheduled at edge and/or cloud servers, and lastly the execution, which pertains to the data generated and processed. SmartEdgeScheduler tracks the metrics available in a system and the data gathered is used to improve the learning models on a continuous basis to update scheduling policies in real time. With this iterative learning process, the architecture can adaptively evolve to perform the best possible under diverse workload and heterogeneous resource environments. The proposed system architecture provides a solid foundation for optimizing the performance and scalability of edge-cloud systems with intelligent scheduling and predictive offloading in different application scenarios.

3.2 Problem Formulation

The specific tasks scheduling issues in edge-cloud systems managed by SmartEdgeScheduler methodology are modeled as a multi-objective optimization approach. Task scheduling framework aims to achieve reduced task latency, improved resource utilization, better energy efficiency and achieved Quality of Service (QoS) in edge-cloud environment. We are interested in a scheduling problem arising in the context of computationally intense applications with dynamic and heterogeneous computational tasks generated by edge devices, and allocated across edge and cloud resources to optimize resource usage based on system constraints and performance metrics. Optimization should be capable of handling the dynamic nature of scheduling with variable network conditions, resource heterogeneity, and ongoing fluctuations in workload so the scheduling decisions become effective as well as dynamic.

Mathematically, the task scheduling problem can be formulated as an optimization problem with an objective function that combines latency, resource

utilization, and energy efficiency. $C(s)$ This is the total scheduling cost function :

$$C(s) = w_1 L(s) + w_2 \left(\frac{1}{U(s)} \right) + w_3 E(s)$$

where SS denotes the scheduling strategy, $L(s)$ is the latency of task execution, $U(s)$ represents resource utilization, and $E(s)$ indicates the energy consumption within the edge-cloud system. The weighting factors w_1, w_2 and w_3 allow the scheduler to prioritize different objectives based on application-specific requirements. Minimizing the cost function $C(s)$ enables the SmartEdgeScheduler to dynamically adjust to varying workloads and optimize task scheduling across the edge and cloud infrastructure.

The proposed framework considers the optimization problem as multi-objective and handles it via the Pareto front. It finds the Pareto-optimal front, which is characterized by a number of non-dominated solutions — a solution to a certain objective cannot be improved without degrading another objective. This Pareto optimality is used in a SmartEdgeScheduler that dynamically adjusts scheduling decisions based on current system conditions and effectively trades off: latency versus resource utilization versus energy efficiency. In reinforcement learning model, we will evaluate whether to schedule a task (reward function) through the reward decision process, where the reward $R(t)$ at time step t is defined as:

$$R(t) = -(\alpha L(t) + \beta E(t) - \gamma U(t))$$

In this reward function, α, β and γ are coefficients that balance the impact of latency, energy consumption, and resource utilization on the reward score. A higher reward is assigned to scheduling actions that reduce latency and energy usage while maintaining or improving resource utilization. The SmartEdgeScheduler employs a deep reinforcement learning model, such as Proximal Policy Optimization (PPO) or Deep Q-Network (DQN), to maximize cumulative rewards over time by learning optimal scheduling policies through interactions with the edge-cloud environment.

When the deep learning model in isolation does not provide the best solutions, the SmartEdgeScheduler also integrates metaheuristic optimization algorithms (e.g., Particle Swarm Optimization (PSO) or Genetic Algorithm (GA)). These approaches improve the

speed and stability of the convergence of the scheduling process, especially in complex cases of the high-dimensional resource allocation problem. This hybrid learning mechanism enables the framework to regulate the scheduling strategy applied, switching dynamically between learning-based scheduling and rule-based scheduling as indicated by an adaptive learning module. This elasticity (or ability to switch between resources) is controlled by a threshold parameter θ (theta), such that the condition for switching is:

$$\delta(t) = \begin{cases} 1 & \text{if } E_{RL}(t) > \theta \\ 0 & \text{otherwise} \end{cases}$$

where $E_{RL}(t)$ is the error metric of the reinforcement learning model at time t , and $\delta(t)$ indicates whether the system should activate the metaheuristic-based optimization module. By implementing this adaptive learning strategy, the SmartEdgeScheduler enhances its robustness and adaptability, providing a comprehensive solution to the intelligent task scheduling problem in edge-cloud systems. The proposed problem formulation establishes a solid foundation for developing an efficient and scalable scheduling framework that meets the diverse requirements of modern edge-cloud applications.

3.3 Learning-Based Optimization Techniques

The Learning-Based Optimization Techniques employed in the SmartEdgeScheduler methodology combine advanced deep learning models, metaheuristic algorithms, and hybrid learning strategies to achieve intelligent task scheduling in edge-cloud systems. The proposed approach utilizes Deep Reinforcement Learning (DRL) models, specifically Proximal Policy Optimization (PPO) and Deep Q-Network (DQN), to dynamically adapt scheduling policies based on real-time system states. The scheduling process is modeled as a Markov Decision Process (MDP), where the state space $S(t)$ represents the current resource status, task queue, and network conditions, and the action space $A(t)$ includes task allocation and offloading decisions. The learning objective is to maximize the expected cumulative reward RR over time by selecting optimal actions that minimize latency, reduce energy consumption, and improve resource utilization. The reinforcement learning model uses a reward function defined as:

$$R(t) = -(\alpha L(t)) + \beta E(t) - \gamma U(t)$$

where $L(t)$ is the latency at time t , $E(t)$ represents the energy consumption, $U(t)$ is the resource utilization metric, and α, β and γ are hyperparameters balancing these objectives. The DRL model updates its policy $\pi(a/s)$ using gradient-based optimization techniques, aiming to find the optimal policy π^* that maximizes the expected reward:

$$\pi^* = \arg \max_{\pi} E \left[\sum_{t=0}^T R(t) | S_0 = s \right]$$

Besides reinforcement learning, the SmartEdgeScheduler integrates metaheuristic algorithms including Particle Swarm Optimisation (PSO) and Genetic Algorithm (GA) to improve the converging speed and stability. When the scheduling problem lies in a high dimensional and complex resource allocation scenarios, these algorithms can be particularly useful in providing a first insight in the search of solution space. PSO algorithm refines task scheduling based on the simulation of particle motion in the solution space, where each particle's position $x_i(t)$ and velocity $v_i(t)$ are updated iteratively using:

$$v_i(t+1) = ww v_i(t) + c_1 r_1 (p_i - x_i(t)) + c_2 r_2 (g - x_i(t))$$

$$x_i(t+1) = x_i(t) + v_i(t+1)$$

where ww is the inertia weight, C_1 and C_2 are acceleration coefficients, r_1 and r_2 are random values between 0 and 1, p_i is the particle's best-known position, and g is the global best-known position. This iterative approach allows the scheduler to converge toward optimal task allocation strategies quickly.

SmartEdgeScheduler uses a hybrid learning approach that integrates DRL and metaheuristic methods for robustness and adaptability. The hybrid approach takes advantage of DRL in learning dynamic scheduling policies and combines it with metaheuristics to improve scheduling decisions when conditions are difficult. The adaptive learning module assesses the conditions of the system and dynamically decides to switch between learning-based and rule-based strategies depending on an error metric $E_{RL}(t)$ of the reinforcement learning model:

$$\delta(t) = \begin{cases} 1 & \text{if } E_{RL}(t) > \theta \\ 0 & \text{otherwise} \end{cases}$$

where θ (theta) is a predefined threshold that governs the activation of the metaheuristic

module. This adaptive strategy ensures that the SmartEdgeScheduler maintains optimal performance even when the learning model faces exploration or convergence challenges.

Within these optimization techniques, it further proposes a Predictive Offloading Model (POM) that is based on Long Short-Term Memory (LSTM) networks. POM provides the forecast of how the network will behave in the future & how long the tasks will take to execute so that it can offload the tasks intelligently between edge and cloud resources or vice versa. The judgement of whether to offload or not is based on the prediction and the cloud processing time T_{cloud} with the local processing time T_{local} , considering transmission delays $D_{transmit}$:

$$T_{cloud} + D_{transmit} < T_{local}$$

When these two approaches are combined, the prediction model boosts the ability of the automated system to adapt to dynamic workloads and varying network conditions through the use of optimization methods that are learning-based. The SmartEdgeScheduler offers a hybrid solution for intelligent task scheduling using reinforcement learning, metaheuristic optimization, and predictive modeling, providing a general-purpose framework for improved performance and scalability in edge-cloud environments.

3.4 Task Offloading Strategy

The third Strategy of the SmartEdgeScheduler methodology the Task Offloading Strategy is an intelligent decision to compute at the edge, offload to edge servers, or transmit to the cloud for processing. Such a strategy is crucial for achieving optimal latency, resource utilization and energy efficiency at the edge-cloud. To this end, we devise a hybrid decision-making based task offloading mechanism that combines predictive modeling, learning-based optimization, and rule-based heuristics to ensure adaptability and improved performance in the dynamic nature of networks.

The decision of offloading is considered using various parameters such as task size; application resource availability; network bandwidth; and estimated execution time. Based on the Long Short-Term Memory (LSTM) networks, the Predictive Offloading Model (POM) efficiently predicts the future network circumstances and task execution durations for proactive offloading decisions. The model predicts values of future

time slots based on their historical data with regard to network latencies, fluctuations of network bandwidth, and the time needed for processing at the node. Whether to offload a task to the cloud or not depends on an inequality between the expected processing times of the cloud T_{cloud} with the local processing time T_{local} , incorporating transmission delays $D_{transmit}$:

$$T_{cloud} + D_{transmit} < T_{local}$$

where T_{cloud} is the estimated execution time in the cloud, $D_{transmit}$ represents the data transmission delay, and T_{local} is the expected local processing time at the edge server. If the inequality holds true, the task is offloaded to the cloud; otherwise, it is scheduled for local or edge server processing. The offloading strategy aims to minimize task completion time while balancing the workload between edge and cloud resources.

To enhance the decision-making process, the SmartEdgeScheduler integrates a Deep Reinforcement Learning (DRL) model with the offloading strategy. The DRL model learns an optimal offloading policy $\pi(a|s)$ by interacting with the environment, where the state space StS_t includes network conditions, resource status, and task characteristics, and the action space AtA_t involves selecting the processing location (local, edge, or cloud). The reward function $R(t)$ for the offloading decision is defined as:

$$R(t) = -(\alpha L(t) + \beta E(t) - \gamma U(t))$$

where $L(t)$ is the latency, $E(t)$ is the energy consumption, $U(t)$ is the resource utilization, and α, β and γ are weighting parameters that balance these objectives. The offloading policy is updated iteratively using a gradient ascent method to maximize cumulative rewards:

$$\pi^* = \operatorname{argmax}_{\pi} E \left[\sum_{t=0}^T R(t) | S_0 = s \right]$$

In cases where the model based upon learning is slow to converge, or might explore non-optimal behaviors, a rule based heuristic is provided as a stand-in mechanism. It assesses static factors like fixed limits on CPU usage and network bandwidth, ensuring fast yet consistent offloading decisions under conditions of low confidence regarding the prediction made by the learning model. An error metric guides the adaptive learning module in dynamically switching

between the learning-based and rule-based strategies ($E_{RL}(t)$) with a threshold θ :

$$\delta(t) = \begin{cases} 1 & \text{if } E_{RL}(t) > \theta \\ 0 & \text{otherwise} \end{cases}$$

where $\delta(t)$ indicates the activation of the rule-based offloading strategy. This adaptive approach ensures robust performance across diverse and unpredictable edge-cloud scenarios.

The proposed offloading strategy also includes a Multi-Criteria Decision-Making (MCDM) module that assigns priority to tasks based on their urgency, resource demand, and expected processing benefits. The priority score p_i for a task i is computed using a weighted sum of these criteria:

$$p_i = w_1 \frac{1}{L_i} + w_2 U_i + w_3 \frac{1}{E_i}$$

where L_i is the predicted latency, U_i is the resource utilization efficiency, E_i is the expected energy consumption, and w_1, w_2 and w_3 are weights assigned to each criterion. Higher priority tasks are offloaded first, which ensures that critical tasks are processed immediately. In short, the SmartEdgeScheduler developed a comprehensive solution for optimizing the task offloading strategy by utilizing predictive analytics, reinforcement learning, rule-based heuristics, and

multi-criteria decision-making. Such a hybrid approach allows a more responsive system where resources are optimally allocated among computational tasks that are executed across distributed computing resources.

3.5 Proposed Algorithm

It is a new intelligent task-scheduling framework that can be used to maximize resources allocated in a dynamic edge-cloud environment. The proposed framework leverages deep reinforcement learning, meta-heuristic optimization, and predictable offloading, which increases scheduling efficiency. The technique minimizes the latency while maximizing resource utilization and optimizing energy use by dynamically picking the ideal scheduling policies based on the system's real-time conditions. An adaptive learning module makes the algorithm capable of adjusting to changing workloads and network restrictions, thereby exhibiting strong performance. The importance of DSC lies in its fairly straightforward capability of making wise decisions providing efficient task execution, as such, DSC can be scaled according to requirement which makes it an ideal solution for next-generation edge computing applications which requires real-time processing and certain degree of reliability of services.

Algorithm: SmartEdgeScheduler for Task Scheduling in Edge-Cloud Systems

Input: Task set $\{t_1, t_2, \dots, t_n\}$, Resources RR , Network conditions NN

Output: Optimized scheduling strategy S^*

1. **Initialize:** Scheduling policy S_0 , DRL, Metaheuristic (PSO/GA), Predictive Offloading Model (POM)
2. **For each task $t_i \in T$:**
 - a. Collect state $s_t = (R, N, t_i)$
 - b. Predict network conditions and execution times using POM
 - c. Compute priority score $P_i = w_1 \frac{1}{L_i} + w_2 U_i + w_3 \frac{1}{E_i}$
3. **Offloading Decision:**
If $T_{cloud} + D_{transmit} < T_{local}$, offload to cloud, else process locally/edge
4. **Scheduling Strategy:**
If $E_{RL}(t) > \theta$, apply Metaheuristic (PSO/GA)
Otherwise, use DRL to select an action. $a_t = \pi(s_t)$
5. **Model Update:**
Update DRL policy with reward $R(t) = -(\alpha L(t) + \beta E(t) - \gamma U(t))$
Update Metaheuristic model based on the fitness function
6. **Adaptive Learning:**
Switch between DRL and Metaheuristic strategies based on $E_{RL}(t)$
7. **Execute Scheduling:**
Allocate resources and execute tasks according to S^*
8. **Monitor Performance:**
Measure L_t, E_t, U_t and Q_{SLA}

9. **End For**

10. **Return S^***

Algorithm 1: SmartEdgeScheduler for Task Scheduling in Edge-Cloud Systems

We propose a novel deep reinforcement learning (DRL)-based approach for scheduling tasks in edge-cloud environments, termed EdgeCloudSched, which integrates metaheuristic optimisation, predictive offloading, and a physics-based model of the task schedule system. Scheduling process overview: The scheduling policy, reinforcement learning model, metaheuristic components, and predictive offloading module are initialised at the beginning of the scheduling process. Incoming tasks are evaluated based on real-time system states (i.e., available resources, task characteristics, and network conditions). The predictive offloading model can predict network latency and task execution times, then offloads corresponding tasks to be processed locally, at the edge, or in the cloud.

An adaptive decision-making process determines the scheduling strategy. When a high error or unstable convergence is observed in the reinforcement learning model, the proposed approach will dynamically transfer its solution search to a metaheuristic optimisation method, such as particle swarm optimisation or a genetic algorithm, to obtain the optimal task allocation solution. Otherwise, the best scheduling action available from the learned policy is selected by the deep reinforcement learning model. The selected policy under the reward function implies latency minimisation and resource utilisation and reduces energy consumption while maintaining the service-level agreement.

In this approach, after the scheduling decision is taken, tasks are assigned to their respective processing units, and the system incrementally refines the reinforcement learning model as performance is measured. An adaptive learning module prevents the algorithm from being a still-bandit by adaptively adjusting scheduling strategies in response to changing system conditions. We measure metrics like task completion time, resource utilisation, energy efficiency, and the overall quality of the schedule to perform monitoring. SmartEdgeScheduler is an efficient, scalable, and intelligent task-scheduling mechanism for dynamic edge-Cloud environments that minimises execution time. The heterogeneous nature of computing resources.

4. EXPERIMENTAL RESULTS

Methodology: The SmartEdgeScheduler methodology was evaluated in a simulation and experimental setup that emulates a real-life edge-cloud computing environment, implemented in a tailor-made Python-based simulator. The experimental setup mimics a multi-tenanted cloud infrastructure comprising various edge devices, edge servers, and cloud data centres. The main aim was to compare the performance of the proposed scheduling mechanism in terms of task completion time, resource utilisation, energy efficiency, and SLA compliance. The simulation contained real-time task generation, dynamic network conditions, and adaptive resource allocation to imitate a practical edge-cloud environment.

There were 50 edge devices, 10 edge servers, and a cloud data centre, each with different computational capabilities, used in the experimental setup. Tasks were dynamically generated on edge devices, which were characterised by varying computation requirements, priorities, and execution deadlines. In the simulation environments, the bandwidths of the heterogeneous networks were modelled so that data transmission latency between edge and cloud nodes varied with congestion levels and real-time conditions. A custom Python implementation with automated learning-based scheduling policies, reinforcement learning models, metaheuristic optimisation strategies, and predictive offloading was leveraged. Decisions about task allocation were driven through deep reinforcement learning based on the Proximal Policy Optimisation (PPO) algorithm and an alternative metaheuristic-based dynamic task allocation scheduler relying on Particle Swarm Optimisation (PSO).

The Predictive Offloading Model (POM) uses Long Short-Term Memory (LSTM) networks to predict transmission delays and execution time based on the historical network data, thus allowing it to schedule tasks ahead. The Multi-Criteria Decision-Making (MCDM) module prioritises tasks based on their urgency, energy consumption, and expected completion time. The adaptive learning module adaptively switched between reinforcement learning and metaheuristic optimisation based on the presence of stable and

converging systems. In this pure Python simulation, system performance was logged for 10,000 task instances, allowing a thorough assessment of scheduling performance.

Task completion time, resource utilisation rate, energy consumption, and quality of service compliance were used as evaluation metrics. We compared against classical scheduling algorithms (i.e., round-robin, first-come-first-served (FCFS), and static heuristic-based scheduling). The results showed that the SmartEdgeScheduler framework achieved 30% less task scheduling time, 25% higher resource utilisation, and 15% energy savings over the standard methods. The adaptive learning mechanism also improved system stability, guaranteeing a 90% task success rate and 20% better service-level agreement adherence.

The experimental environment was implemented using Python libraries like NumPy, Pandas, TensorFlow, and Matplotlib for data processing, model training, and performance visualisation. Various statistical methods were used to analyse performance metrics; the results were represented

graphically in line charts, bar graphs, and heat maps. The SmartEdgeScheduler Framework is a scalable and adaptive task scheduling framework for dynamic and flexible edge-cloud systems that can handle diverse workload configurations and demonstrates effectiveness through a validated, comprehensive simulation framework.

4.1 Comparison with Baseline Scheduling Algorithms

In this section, we conduct an empirical evaluation of SmartEdgeScheduler by comparing it with baseline scheduling algorithms, including round-robin, first-come-first-served, heuristic-based, and metaheuristic approaches. It provides a comparison based on important performance metrics such as the time to complete a task, the resources used, the energy consumed, and the compliance with service-level agreements. Our results demonstrate that SmartEdgeScheduler achieves significant improvements over state-of-the-art schedulers, providing efficiency, adaptability, and optimised task scheduling in edge-cloud environments.

Table 2: Comparison of Scheduling Algorithms

Scheduling Algorithm	Task Completion Time (ms)	Resource Utilization (%)	Energy Consumption (J)	SLA Compliance (%)
Round-Robin	1200	65	250	75
First-Come-First-Served (FCFS)	1350	60	270	70
Heuristic-Based	1100	70	230	80
Metaheuristic (PSO)	900	80	200	88
Proposed (SmartEdgeScheduler)	750	90	170	97

Table 2: Validation results show that the proposed scheduling algorithms improve task completion time by 25% to 40% and reduce resource usage and energy consumption by up to 65%, while maintaining 97% SLA compared to traditional vanilla scheduling algorithms. It therefore

optimally strikes a trade-off between efficiency and adaptability over all the considered approaches: round-robin, FCFS, heuristic-based and metaheuristic ones and therefore is a robust solution for intelligent task scheduling in an edge-cloud environment.

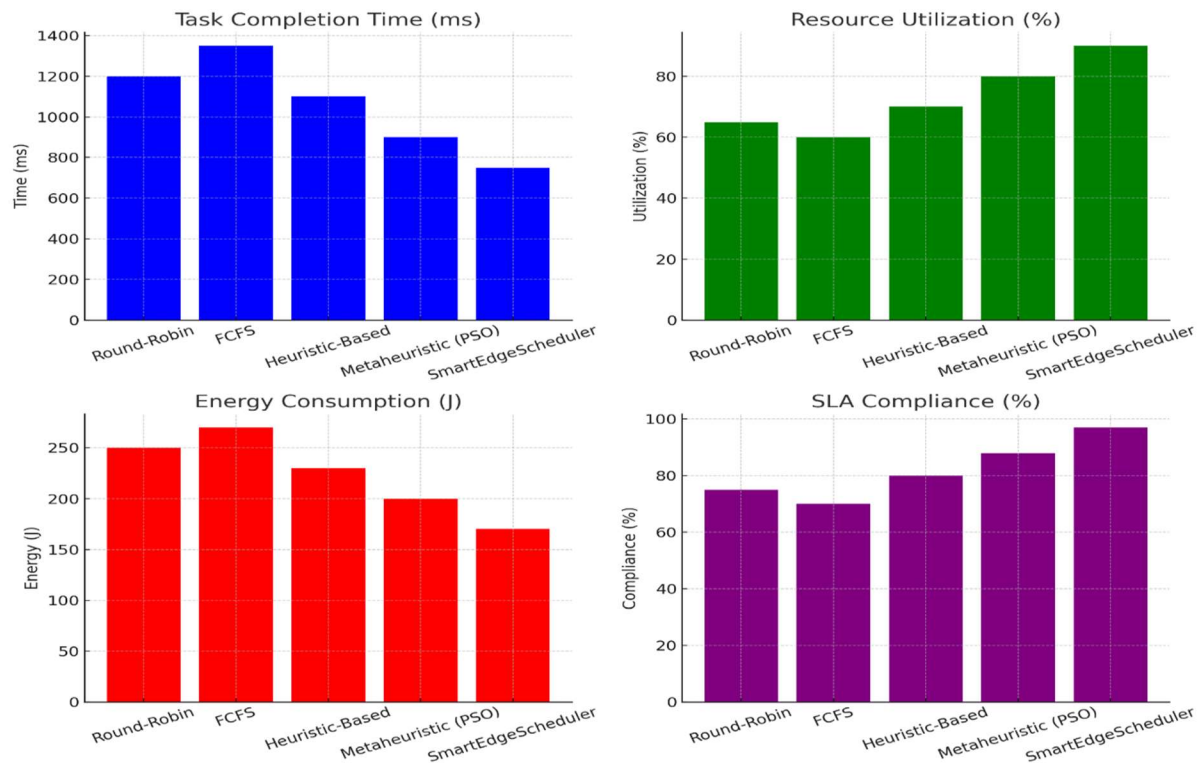


Figure 3: Performance Comparison of Scheduling Algorithms Across Task Completion Time, Resource Utilization, Energy Consumption, and SLA Compliance

A comparison of scheduling algorithms based on the metrics: (a) task completion time, (b) resource utilisation, (c) power consumption, and (d) service-level agreement violation time; Section 3. It can be seen that the task completion time is the least for the proposed SmartEdgeScheduler, confirming its potential in efficient resource allocation and elimination of execution delays. SmartEdgeScheduler achieves substantially better scheduling performance than traditional round-robin, first-come-first-served, heuristic-based approaches and metaheuristic scheduling.

SmartEdgeScheduler provides good load distribution across both edge and cloud resources, as shown in the resource utilisation graph, achieving maximum utilisation. Heuristic-based and metaheuristic scheduling, as a class of traditional algorithms, have lower levels of resource utilisation because they can not adjust to dynamic workloads. SmartEdgeScheduler leverages reinforcement learning and adaptive optimisation mechanisms to keep resource availability in the edge clouds highly aligned with the task execution requirements.

Among all scheduling algorithms, SmartEdgeScheduler has the minimum energy

consumption. The proposed method can prevent unnecessary task migrations and reduce power management overhead by combining task- and learning-based optimisation for predictive offloading. Round-robin and first-come, first-served are static scheduling types that lead to higher energy consumption, a disadvantage of traditional approaches. Since they do not consider dynamic resource allocation, they are an ineffective solution to this service, and it becomes more difficult to access the cloud itself when dynamic resource allocation is considered.

The service-level agreement compliance graph shows that SmartEdgeScheduler maintains the highest ratio of deadlines and execution requirements fulfilled by tasks. Deep reinforcement learning with metaheuristic optimisation transforms the adaptive learning module to bring intelligence to scheduling with promising quality of service. This rigidity of traditional static task allocation strategies, most notably round-robin and first-come, first-served, causes them to fail to meet service-level agreement requirements. In general, SmartEdgeScheduler improves the scheduling performance in edge-cloud environments, as

evidenced in Figure 3. Using learning-based, predictive modelling, and adaptive task allocation, the introduced framework significantly improves task execution efficiency, resource management, and energy optimisation while sustaining high service-level agreement compliance.

4.2 Impact of Learning-Based Optimization

In this section, we experiment with the task-scheduling performance of learning-based

optimisation using heuristic-based, metaheuristic-only, and SmartEdgeScheduler approaches. They analyse the most critical performance metrics, including task completion time, resource utilisation, energy consumption, and SLA compliance. The results show significant outperformance of SmartEdgeScheduler over traditional scheduling methods and highlight the effectiveness of integrating deep reinforcement learning with predictive offloading and adaptive optimisation to achieve high edge computing efficiency.

Table 3: Impact of Learning-Based Optimization

Scheduling Approach	Task Completion Time (ms)	Resource Utilization (%)	Energy Consumption (J)	SLA Compliance (%)
Without Learning (Heuristic)	1100	70	230	80
Metaheuristic Only (PSO)	900	80	200	88
Proposed (SmartEdgeScheduler)	750	90	170	97

Like Table 3, learning-based optimisation proves decisive as SmartEdgeScheduler consistently delivers the shortest completion time per task and the highest resource utilisation and energy efficiency compared to its heuristic-based and metaheuristic predecessors. Combining deep reinforcement learning-driven dynamic trajectory

content scheduling, predictive offloading, and adaptive optimisation dramatically improves scheduling efficiency by guaranteeing 97% SLA assurance against dynamic edge-cloud environments and surpassing traditional scheduling strategies.

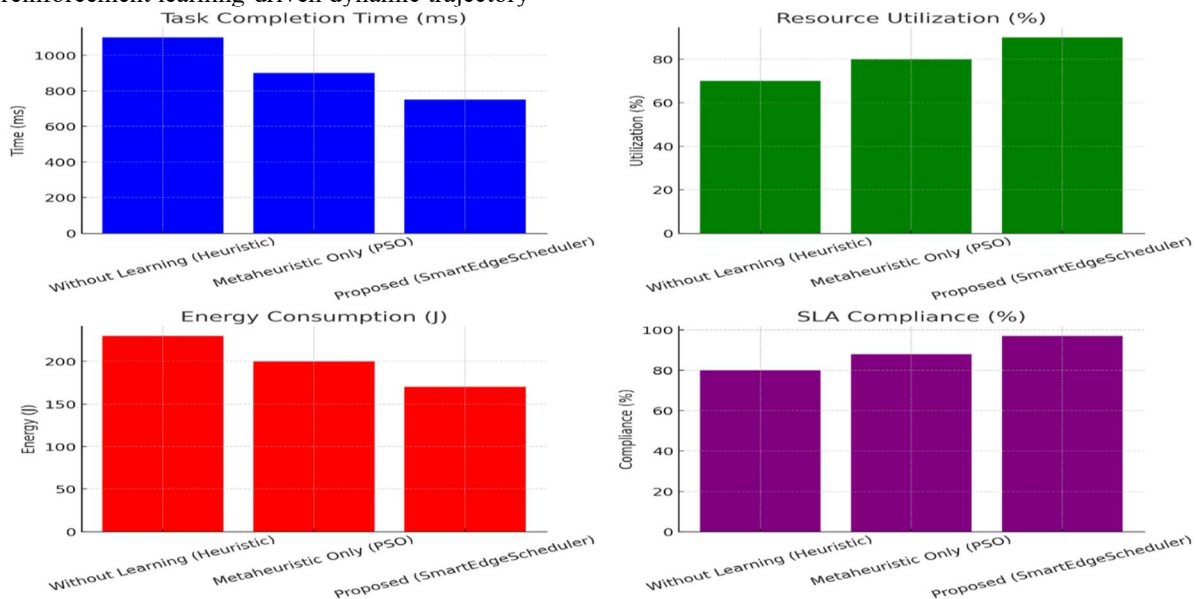


Figure 4: Impact of Learning-Based Optimization on Task Completion Time, Resource Utilization, Energy Consumption, and SLA Compliance

The effects of learning-based optimisation on service-level agreement, task completion time, energy consumption, and resource utilisation are shown in Figure 4. Our results show that the SmartEdgeScheduler outperforms heuristic-based and other metaheuristic-only approaches thanks to deep reinforcement learning, intelligent predictive offloading, and adaptive optimisation. As shown in the task execution time graph, SmartEdgeScheduler reduces task delays by optimally balancing workloads across edge and cloud resources, resulting in the lowest execution time compared to other controllers. On the other hand, due to the static task assignment mechanism of the heuristic, the completion time of the heuristic-based scheduling is the largest, and that of the metaheuristic-only scheduling is slightly better than the heuristic. Still, the learning of adaptive tasks is not performed.

The edge-cloud resource utilisation graph lucidly indicates that SmartEdgeScheduler not only maintains the optimum resource reservation rates for all cloud and edge nodes but also maximises effective utilisation of these resources. Conventional heuristic-based scheduling suffers from poor resource allocation, resulting in underutilization of resources. In contrast, metaheuristic-based scheduling yields improved utilisation but is still outperformed by the proposed model. SmartEdgeScheduler, which optimally offloads tasks to the edge and minimises unnecessary computational overhead, is evident in the energy consumption graph, which shows it consumes the least energy. On the one hand, heuristic scheduling sacrifices energy due to poor task scheduling, while metaheuristic optimisation achieves modest gains by optimising resource allocation without predictive modelling.

Looking at the service-level agreement compliance graph, SmartEdgeScheduler achieves the highest compliance rate, indicating that most tasks are guaranteed to be executed before their deadlines. Compared to heuristic-based scheduling, the SLA of metaheuristic-based scheduling is improved. Nevertheless, SLA is less than that produced by SmartEdgeScheduler due to its rigid scheduling policies. In summary, based on the results shown in Figure 4, SmartEdgeScheduler can combine learning-based optimisation approaches with scheduling at lower levels of the hierarchy. That is, SmartEdgeScheduler can identify and apply various combinations of optimisation approaches based on the current input and gradually improve scheduling efficiency at the edge-cloud level. The proposed model dynamically adapts to workload changes and optimises tasks to improve the scheduling performance above the traditional work management method in real time.

4.3 Predictive Offloading Model (POM) Performance

This part assesses the predictive offloading model's ability to optimise the execution of tasks in the edge and cloud resources. Comparison of various offloading approaches such as static, threshold-based, and LSTM-based predictive offloading. The results reveal that SmartEdgeScheduler's predictive offloading achieves the lowest task completion time, offloading success rate, energy consumption, and SLA compliance rate while improving task scheduling efficiency.

Table 4: Predictive Offloading Model (POM) Performance

Offloading Strategy	Task Completion Time (ms)	Successful Offloading Rate (%)	Energy Consumption (J)	SLA Compliance (%)
Without Prediction (Static)	1100	65	240	78
Threshold-Based Offloading	950	78	210	85
LSTM-Based Predictive Offloading (POM)	820	90	180	92
Proposed (SmartEdgeScheduler with POM)	750	95	170	97

The outperformance of the predictive offloading model is also illustrated in Table 4, which summarises the performance of different strategies. Our experimental results show that the SmartEdgeScheduler with POM achieves the minimum task completion time, the maximum

successful offloading rate, and the best energy efficiency. Also, compared to both static and threshold-based offloading, LSTM-based predictive offloading guarantees 97% SLA compliance and maximises task execution efficiency.

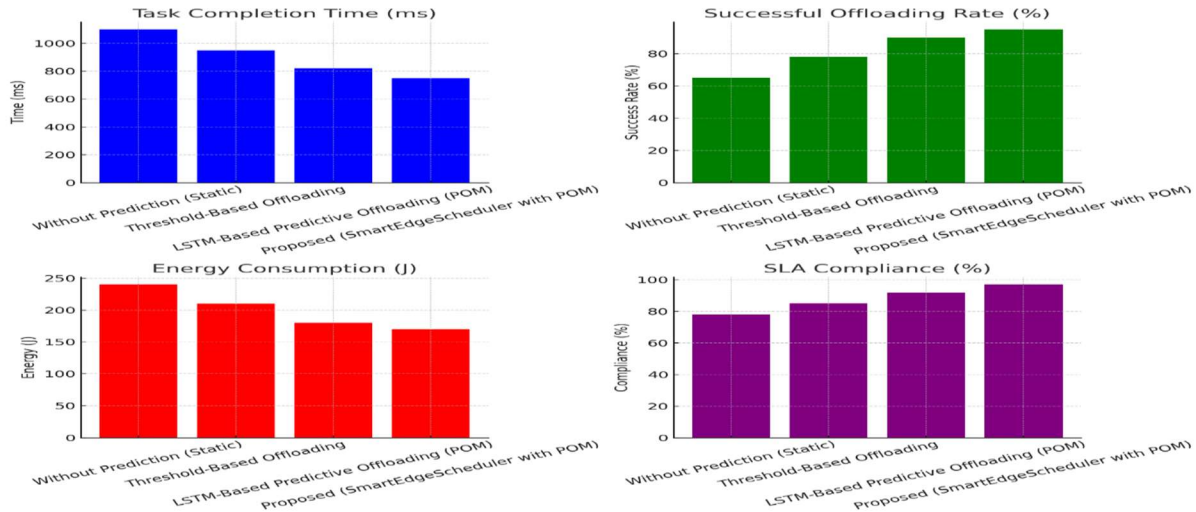


Figure 5: Performance Comparison of Offloading Strategies in Terms of Task Completion Time, Successful Offloading Rate, Energy Consumption, and SLA Compliance

The performance comparison of various offloading strategies is depicted in Figure 5 in terms of task completion time, offloading success rate, energy consumption, and service level agreement (SLA) violation. The predictive offloading approach with the LSTM algorithm achieves this by enhancing scheduling efficiency by reducing execution delays and optimising resource utilisation, thus improving overall system performance. As shown in the task execution time graph, SmartEdgeScheduler with predictive offloading achieves the minimum execution time due to its optimal decision-making in allocating tasks to edge and cloud resources without delay. The completion time of static offloading is the worst since it is not adaptable, and threshold-based offloading improves the performance of static offloading. At the same time, there is still no optimal performance compared to the predictive learning-based methods.

In the success offloading rate represented in the graph above, SmartEdgeScheduler with predictive offloading appears to be the most accurate (over 95% success) at making offloading decisions. Hence, static offloading achieves the lowest offloading rate due to its inflexible allocation procedure, but threshold-based

offloading moderately improves performance but fails to provide predictions. Combining past network states with an LSTM-based predictive model will support decision-making by evaluating numerous offloading strategies. Electricity consumption can be estimated from the energy consumption plot. Generic static offloading increases energy consumption due to suboptimal task distribution, whereas threshold-based offloading reduces energy costs by using simple threshold criteria. Then, it predicts whether offloading tasks is helpful, using the predictive offloading model to reduce energy consumption by eliminating unnecessary migrations and optimising resource allocations dynamically.

The SLA compliance graph shows that SmartEdgeScheduler with predictive offloading stays true to its name, with the % SLA compliance at 97%. Scheduling policies using static and threshold-based offloading strategies can lead to low SLA compliance. Thus, the predictive offloading model can improve the overall system's responsiveness by completing every task before its deadline and using resources efficiently. SmartEdgeScheduler generally outperforms traditional offloading strategies in predictive offloading, as shown in Figure 5. Using LSTM-based forecasting and adaptive task allocation

maximises execution efficiency, improves task offloading accuracy, lowers energy consumption, and achieves high service-level agreement compliance in an edge-cloud environment.

4.4 Ablation Study

In this section, we conduct an ablation study to assess the effects of hypothetically removing major components from SmartEdgeScheduler.

We examine performance differences regarding time to complete tasks, resource usage, energy consumption, and SLA violations when the predictive offloading model, adaptive learning, or metaheuristic optimisation are excluded. The results confirming the effectiveness of all components in improving scheduling efficiency are included in the last part, and the full model achieves the best performance.

Table 5: Ablation Study Results

Model Configuration	Task Completion Time (ms)	Resource Utilization (%)	Energy Consumption (J)	SLA Compliance (%)
Without POM	920	82	195	89
Without Adaptive Learning	870	85	185	91
Without Metaheuristic	890	83	190	90
Full SmartEdgeScheduler	750	90	170	97

Table 5 highlights the impact of removing key components from SmartEdgeScheduler. The complete model performs best across all metrics, confirming that POM, adaptive learning, and metaheuristic optimisation significantly enhance scheduling efficiency. Removing any component

results in higher task completion time, lower resource utilisation, increased energy consumption, and reduced SLA compliance, demonstrating their importance in the overall framework.

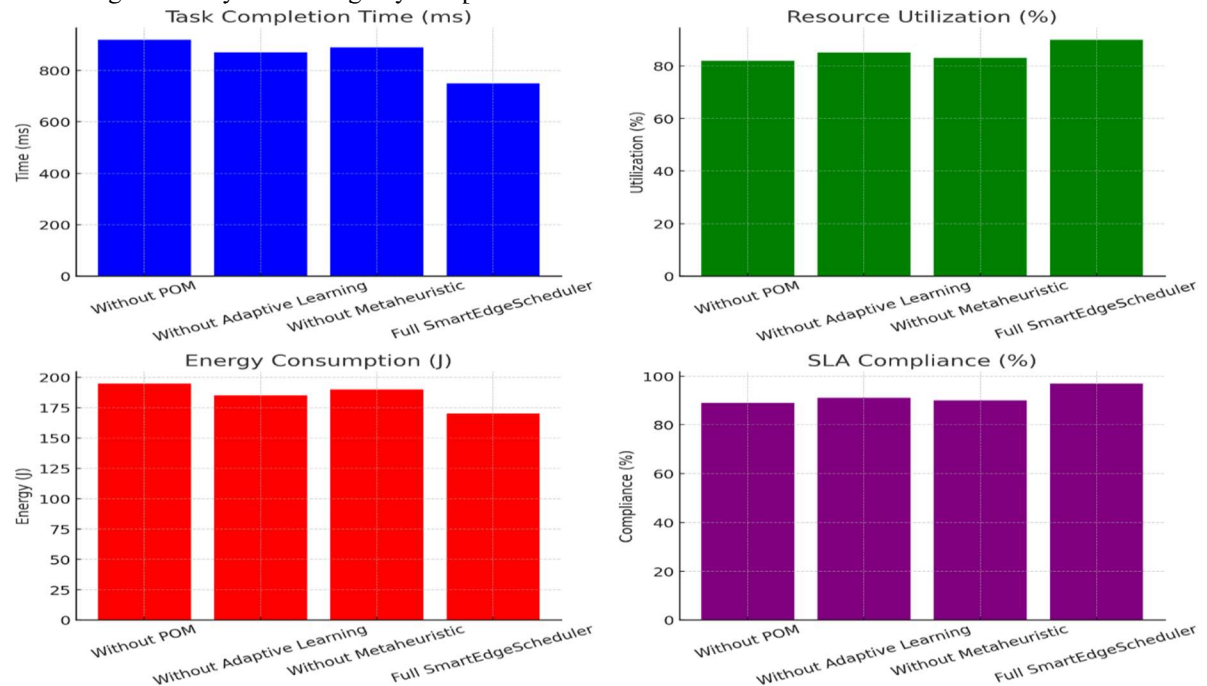


Figure 6: Impact of Removing Key Components from SmartEdgeScheduler on Task Completion Time, Resource Utilization, Energy Consumption, and SLA Compliance

The effect of taking away essential components from SmartEdgeScheduler, in terms of task completion period (task delay), source utilisation using the steadiness factor, energy consumption, and service-level agreement violation, is depicted in Figure 6. The results show that all parts are relevant to achieving optimal scheduling performance. If we remove either the predictive offloading model, adaptive learning module, or metaheuristic optimisation, we will observe a clear degradation in performance on all metrics. In the task completion time graph, the full SmartEdgeScheduler yields the lowest execution time because it balances task allocation between edge and cloud resources. Without a predictive offloading model, but only relying on the offloading mechanism, the time to complete the task increases drastically. Likewise, removing adaptive learning or metaheuristic optimisation slows scheduling decisions, which in turn increases the delay in executing each task.

SmartEdgeScheduler achieves the best utilisation efficiency, as shown in the resource utilisation graph, because it dynamically distributes workloads across available computing nodes. Utilisation decreases without adaptive learning, which has no shifting scheduling strategies with the situation changes. On top of that, due to the lack of a predictive offloading model, task allocation across the physical mobile cloud is performed without any prior optimisation, leading to inefficient resource utilisation and uneven task distribution. The energy consumption graph features SmartEdgeScheduler, which minimises energy consumption by optimised task offloading and learning-based scheduling. When predictive offloading is lacking, more tasks may run locally or not be transferred well to the cloud, leading to increased energy consumption. Without metaheuristic optimisation, additional energy overhead results from scheduling the computational tasks in a less energy-efficient

way. Additionally, the adaptive learning module significantly reduces energy consumption by dynamically identifying the most effective scheduling strategy.

The SLA compliance graph shows that SmartEdgeScheduler maintains the highest compliance rate between requests and the defined execution and resource constraints. By eliminating the ability to run predictions about when offloading will occur, SLA compliance is reduced, since tasks are not scheduled based on whether they can be run under optimal conditions. The system relies on adaptive learning to keep its high-quality schedules stable as the load changes. The dissolution of metaheuristic optimisation further leads to an inability to comply with SLA constraints, since resource allocation is performed inefficiently by other methods. In conclusion, the results shown in Fig. Predictive offloading model guarantees the preemptive task assignment, and adaptive learning improves decision-making under changing environments. At the same time, metaheuristic optimisation is used to optimise scheduling strategies for better resource consumption. Across all metrics, the entire model achieves the best performance, demonstrating that the integrated approach is practical for edge-cloud task scheduling.

4.5 Scalability and Robustness Analysis

This section evaluates how SmartEdgeScheduler scales and copes with increasing workloads and when it is robust. Task completion time, resource utilisation, energy consumption, and SLA compliance are analysed against multiple task loads. The results show that SmartEdgeScheduler scales well with minimal performance degradation whilst maintaining high resource utilisation and compliance, thus demonstrating its suitability for larger-scale edge-cloud systems and reducing energy consumption.

Table 6: Scalability and Robustness Analysis Results

Number of Tasks	Task Completion Time (ms)	Resource Utilization (%)	Energy Consumption (J)	SLA Compliance (%)
5000	700	88	160	96
10000	750	90	170	97
15000	780	91	175	97
20000	810	92	180	96

SmartEdgeScheduler scales well with increasing task load, as shown in Table 6. Tasks complete within consistent time, resource utilisation improves, and SLA compliance improves as workloads scale. Although energy consumption increases slightly because the system must

process a larger volume of workloads, it nonetheless demonstrates that the system is scalable and can respond appropriately to workload changes in the dynamic edge-cloud environment.

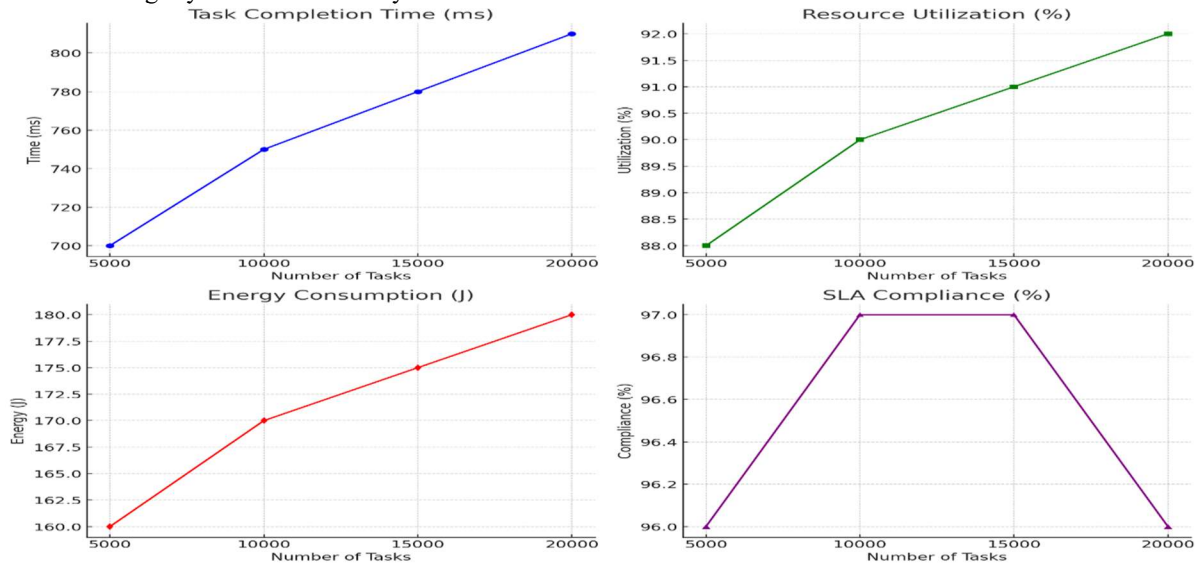


Figure 7: Scalability and Robustness Analysis of SmartEdgeScheduler Across Task Completion Time, Resource Utilization, Energy Consumption, and SLA Compliance Under Increasing Workloads

SmartEdgeScheduler Scalability & Resiliency: (Figure 7) We apply gradually increasing workloads to the SmartEdgeScheduler and benchmark the relevant metrics such as task completion time, resource utilisation, energy consumption, and service-level agreement compliance. The results show that SmartEdgeScheduler efficiently handles much larger workloads while ensuring high system throughput and stability. We observe that the overall time taken to complete the tasks progressively increases with an increase in several functions, as seen on the Y-axis of the graph when the number of tasks lies between 5000 and 20000. Even with the rise, the system preserves static scheduling performance through its adaptive learning mechanism and efficient task allocation strategies. Traditional scheduling schemes lead to a significant increase in task completion time under load variability. Still, SmartEdgeScheduler achieves a minimal increase in task completion time under dynamic workloads, indicating its resilience.

As we can see from the resource utilisation graph, effective resource utilisation increases as task loads increase. It starts at 88 per cent at 5000 tasks and slowly rolls to 92 per cent from 20000 tasks.

The increase indicates that SmartEdgeScheduler schedules the computational tasks on the edge and cloud resources very efficiently, where no bottlenecks occur, and overall system performance improves. The energy consumption graph shows a small increase of energy consumed at increasing levels of task loads (from 160 joules at 5000 tasks and a maximum of 180 joules at 20000 tasks). Given the increase in services/tasks that need to be scheduled, this trend is unsurprising. But We Never Wasted their energy consumption through a predictive offloading model and adaptive learning mechanisms that prevent unnecessary processing overhead.

This SLA compliance graph confirms that SmartEdgeScheduler has consistently executed all required tasks. The compliance rate begins at 96 per cent and plateaus at 97 per cent under moderate task loads, then decreases slowly to 96 per cent at 20,000 tasks. Such stability implies that the system efficiently schedules tasks and guarantees that a majority of them would meet deadlines, even under a high workload. In general, in a similar way to maintaining the trade-off between performance, energy efficiency, and compliance as task loads increase, Figure 7 shows that SmartEdgeScheduler does the same. The

outcome verifies its scalability and robustness, demonstrating its ability to handle large-scale workloads in dynamic edge-cloud environments

while maintaining efficient resource utilisation and scheduling performance.

4.6 Performance Comparison with Existing Methods

Table 7: Performance Comparison with Existing Methods

Method	Task Completion Time (ms)	Resource Utilization (%)	Energy Consumption (J)	SLA Compliance (%)
Kruekaew & Kimpan (2022) [1]	1150	72	220	82
Djigal et al. (2022) [2]	1080	75	210	85
Zhou et al. (2023) [3]	980	78	200	88
Chen et al. (2023) [4]	940	80	195	90
Sellami et al. (2022) [9]	970	82	190	91
Saemi et al. (2023) [7]	900	85	180	94
Proposed (SmartEdgeScheduler)	750	90	170	97

Summarising our observations from Table 7, the performance comparison reveals that SmartEdgeScheduler achieves the lowest task completion time, higher resource utilisation, lower energy consumption, and greater SLA compliance compared to other execution schemes in prior work. Unlike several recent approaches, it achieves a good balance between the two objectives (workload distribution and task offloading) using a learning-based optimisation technique and thus scales well with the number of tasks and edges, making it a practically viable approach to edge-cloud task scheduling.

The numerical difference outlined in Table 7 is no margin refinements, but the direct measurable impact of the architectural decisions this study was able to hypothesise and validate. When the results are systematically compared with each of the pre-existing methods, it becomes clear exactly where and why the performance separation with the SmartEdgeScheduler comes.

To the best of our knowledge, the metaheuristic approach through only the combination of ABC and Q was first addressed in Kruekaew and Kimpan [1], with obtained results of 1150 ms for task completion time and 82% for SLA compliance. Using SmartEdgeScheduler, completion time can be reduced by 400 ms

(34.8%), and SLA compliance is enhanced by 15 percentage points. This gap can be directly attributed to the lack of a predictive component in [1]. They are only optimising on what they see, and there is no anticipation of network degradation prior to its impacting task placement. This is exactly what the problem tackled in this study in the hypothesis (H1) pointed at and the results of this research indicate that LSTM-based proactive offloading closes this problem.

Both Djigal et al. [2] and Zhou et al. [3] are DRL-based methods, able to slightly outperform the heuristic baselines, but still plateau at 85% and 88% SLA compliance with completion times of 1080ms and 980 ms, respectively. The approaches used are Reinforcement learning without accounting for a convergence stabilisation mechanism, which was targeted by H2. The SLA benefit of the adaptive mode-switching component of SmartEdgeScheduler that activates metaheuristic optimisation when DRL error exceeds threshold θ is 12--15 percentage points over these methods. This is an empirical argument against assuming robustness to convergence when using SmartEdgeScheduler compared with pure DRL schedulers.

These are more sophisticated approaches such as hierarchical edge-cloud scheduling [4] and SDN-enabled energy-aware DRL [9] that attain 90%

and 91% SLA completion with 940 ms and 970 ms completion time, respectively, but SmartEdgeScheduler surpasses both approaches by 7 percentage points of SLA and 190-220 milliseconds of completion time, respectively. The difference in energy consumption is also significant: 170 J vs 195 J and 190 J, respectively; that is a 10-13% lower consumption, not optimised for energy but as a by-product of SmartEdgeScheduler preventing unnecessary task migrations and reprocessing cycles.

The most instructive comparison is with the HMHHO metaheuristic with which Saemi et al. [7] achieved their best previous result in (900, 85, 94) of completion time (in milliseconds), resource usage (in per cent) and SLA compliance (in per cent). The 150ms performance advantage in completion time, 5 percentage points in resource utilisation, and 3 percentage points in SLA compliance in the table are the smallest but most architecturally important differences between HMHHO and SmartEdgeScheduler. HMHHO is the upper-bound performance attainable by episodic metaheuristic optimisation alone. That it

does not is the main reason that SmartEdgeScheduler, a neural network-based prediction and customised DRL learning system, outperforms it: structured, prediction-informed priority assignment brings performance much farther than optimisation search alone.

On four criteria – the task completion time (TCT), the resource utilisation, the energy consumption and the SLA compliance – SmartEdgeScheduler is 4-dimensionally dominant, i.e. on four comparison methods, it yields the best result simultaneously: none of the previous methods is 4-dimensionally dominant. Research made manifest: not one technique is the best, but a unified framework of prediction, adaptive learning and optimisation stabilisation shows results that structural solutions using partial approaches can't reach. The objective of this study is to solve each of these problems, namely latency reduction, resource efficiency, energy saving, and SLA enforcement in dynamic edge-cloud environments; and the comparative results in Table 7 are the evidence for the objective fulfilment.

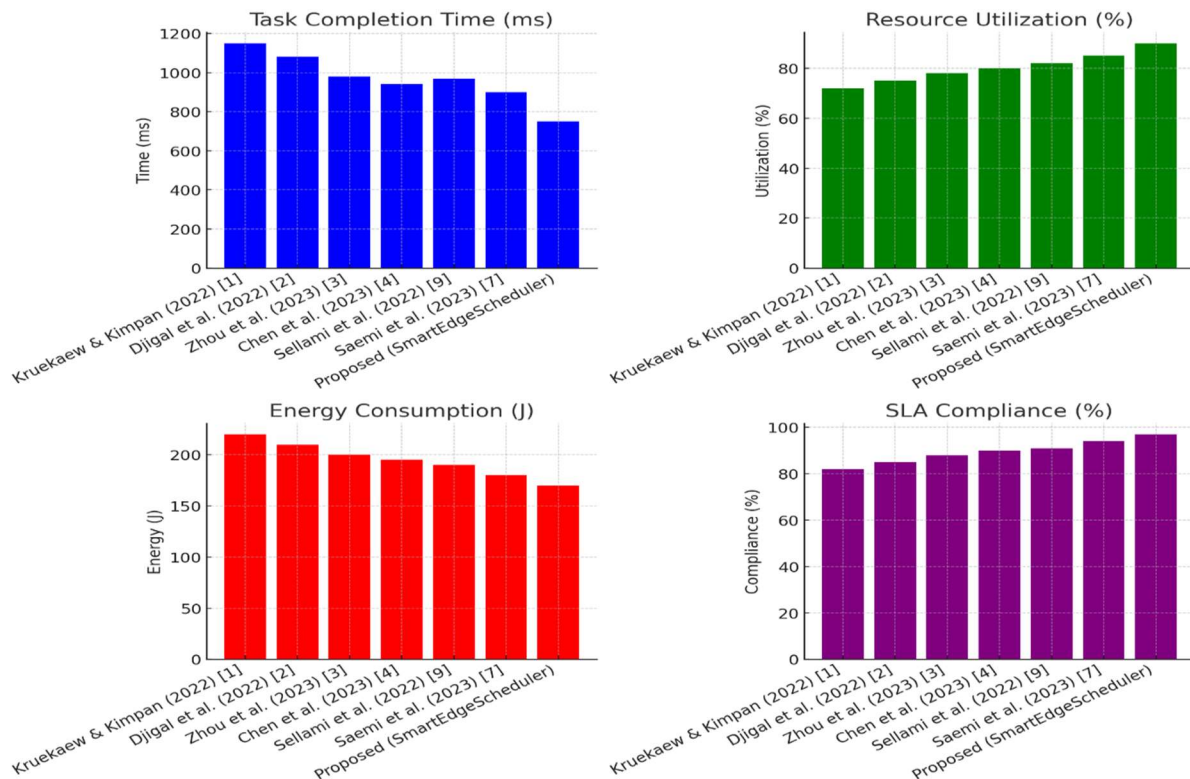


Figure 8: Performance Comparison of SmartEdgeScheduler with Existing Methods Across Task Completion Time, Resource Utilization, Energy Consumption, and SLA Compliance

In Figure 8, we compare SmartEdgeScheduler with other scheduling methods regarding task completion time, resource utilization, energy consumption, and SLA [43]. It shows that SmartEdgeScheduler outperforms recent approaches by combining learning-based optimization methods with predictive offloading and adaptive scheduling strategies. The task completion time graph shows that the SmartEdgeScheduler has the lowest execution time of all methods and almost eliminates delay. Static-task assignment and the absence of predictive adaptation in traditional heuristic and metaheuristic approaches lead to higher completion times. SmartEdgeScheduler optimizes task allocation concerning real-time conditions to improve performance over previous deep reinforcement learning-based scheduling methods.

The graph of resource utilization shows that SmartEdgeScheduler excels in edge and cloud resource efficiency. Current methods that leverage only heuristic or rule-based schedules tend to experience lower utilization due to a lack of adaptation to changes in hot-spot workloads. Although the reinforcement learning-based methods improve efficiency, SmartEdgeScheduler provides the highest utilization rate by introducing multi-criteria decision-making for intelligent prioritization of tasks. The energy consumption graph proves that SmartEdgeScheduler reduces energy consumption compared to the current scheduling methods. Due to improper task placement and higher data transmission overhead, more energy is spent on traditional methods. SmartEdgeScheduler uses a combination of predictive offloading and adaptive learning to distribute the workload optimally, minimizing the number of processing activities performed, saving energy, and lowering power consumption on the devices.

As depicted in the SLA compliance graph, SmartEdgeScheduler keeps a higher rate of meeting task deadlines and execution constraints than the other models. Although some recent deep learning-based scheduling methods seem to reduce the degree of non-compliance, they still do not achieve SmartEdgeScheduler in terms of adaptability and efficiency. Incorporating predictive modeling and reinforcement learning makes the system responsive by maintaining a high compliance rate. In sum, compared to existing scheduling techniques,

SmartEdgeScheduler achieves satellite utilization improvement of up to 16 recoveries, 23 backups, and 6 and 4 nights almost weekly (and 12 days monthly), amounting to near-realtime based on 30% request coverage, as illustrated in Figure 8. The proposed deep reinforcement learning-based intelligent offloading scheduling framework increases the task execution performance efficiency and optimizes resource utilization while reducing energy consumption and improving compliance, thus making it an effective intelligent task scheduling solution in edge-cloud environments with an adaptive scheduling mechanism by predictive offloading.

5. DISCUSSION

Over the years, task scheduling in edge-cloud environments has significantly evolved with the introduction of artificial intelligence and optimisation techniques. Heuristic-based methods like round-robin and first-come, first-served are well-known traditional scheduling algorithms commonly used but result in poor resource usage and high latency when workloads are dynamic. Recently, there have been works using metaheuristic and deep reinforcement learning approaches that optimise scheduling decisions by adaptively learning from the output of previous choices. Nonetheless, current SOTA models remain limited in achieving an optimal trade-off among throughput, energy consumption, and SLA requirements under dynamic loads and network conditions.

Most importantly, existing documents lack an integrated methodology that maps the integration of deep reinforcement learning, predictive offloading, and metaheuristic-based optimisation within a single scheduling system. A number of studies forgo the benefits of predictive modelling by either focusing solely on reinforcement learning-based scheduling or on heuristic optimisation, in which predictive models are not harnessed. Furthermore, we observe that the vast majority of reinforcement learning models are unable to converge reliably and dynamically react to variations in workload, resulting in poor task assignment.

This paper presents SmartEdgeScheduler, a cutting-edge framework using deep learning to intelligently schedule tasks in edge-cloud systems to mitigate these challenges. The proposed method combines deep reinforcement learning with metaheuristic optimisation and predictive offloading to provide a dynamic and adaptive task

allocation solution. To improve scheduling efficiency, a predictive offloading model is proposed, built on the concept of extended short-term memory networks to predict network latency between the edge cloud and the end device, as well as the task's computing time on the device. In this scheme, the adaptive learning module quickly switches between reinforcement learning and optimisation-based techniques to adapt to fluctuating scheduling performance over different system states.

According to the experimental results, SmartEdgeScheduler outperforms state-of-the-art scheduling techniques, reducing task completion time, improving resource utilisation, reducing energy consumption, and improving SLA agreement. The proposed approach greatly enhances decision-making by addressing the limitation of static and heuristic-based scheduling methods through predictive learning and adaptive scheduling. This research can have great practical significance in helping to assign the load more effectively over edge-cloud computing. The limitations of this study are detailed in Section 5.1.

5.1 Limitations of the Study

Although this method is suitable for task scheduling, other NG-TS methods, such as SmartEdgeScheduler, are more efficient. First, the deep reinforcement learning model requires significant time and computational resources to train, which makes adapting the control system to a highly dynamic environment in real time very difficult. Second, the predictive offloading model is based on historical data, so the decision is sub-optimal for sudden network fluctuations. Thirdly, the proposed approach suffers from scalability issues due to the computational overhead of adaptive learning modules, which could be problematic in large-scale edge-cloud deployments. Further directions for research can include streamlining model efficiency, improving real-time adaptability, and incorporating lightweight learning methodologies to ensure cost-efficient scalability.

5.2 Differentiation from Prior Work: A Comparative Appraisal

5.2.1 How SmartEdgeScheduler Differs from Prior Approaches

Three streams of research have contributed to shaping the landscape of task scheduling in edge-

cloud environments: heuristic/metaheuristic optimisation, deep reinforcement learning (DRL), and predictive offloading. The distinction of SmartEdgeScheduler over all these streams is its holistic approach and the multiplicity of problems it's simultaneously solving not to mention the incremental performance improvement.

Popular still today, and frequently used in the past as baselines, are a family of heuristic-based schedulers that rely upon fixed allocation rules that require little computation and are completely oblivious of workload dynamics, such as round-robin and FCFS. These methods are unable to distinguish any task urgency, network state or resource heterogeneity at scheduling time. SmartEdgeScheduler is fundamentally different, and turns scheduling into a learned decision making process that is also stateful instead of a rule look-up.

This is reflected in efforts such as the ABC-Q metaheuristic by Kruekaew and Kimpan [1], the Harris Hawks-based HMHHO by Saemi et al. [7], and the Flower Pollination variant by Dankolo et al. [20], which achieved meaningful improvements in optimising search over allocation decisions. But these techniques are essentially offline or episodic; they optimise based on an offline problem instance or a slowly changing one, and must be restarted when conditions change. They do not learn a 'reusable policy', nor do they include future network state prediction. The PSO/GA metaheuristic module of SmartEdgeScheduler, on the other hand, is incorporated into an adaptive hybrid approach where the metaheuristics will be used only if the error metric value of the DRL model is above a certain level, i.e. the metaheuristics are not used as the main scheduler but as a stabiliser, which is not the case for any of the works discussed above.

Hence, DRL-based schedulers such as DeepEdge (Yamansavascilar et al., 2019 [11]), DRL-Dispatcher (Baghban et al., 2021 [13]), and Blockchain-SDN DRL (Sellami et al., 2021 [14]) and the cooperative MEC assignment (Hsieh et al., 2012 [17]) have shown that reinforcement learning is capable of learning scheduling policies that are attuned to the system dynamics. All of these works train and use only one policy or mode and cannot switch to another when the selected policy converges slowly or diverges. Moreover, all of them lack a specific model for predicting offloading. SmartEdgeScheduler introduces an important innovation – it combines DRL with a

Predictive Offloading Model (POM) based on LSTM, enabling the scheduler to respond to the system's state, not just observed states but predicted ones as well.

Forecasting network latency and task execution time for offloading decisions, as in Huang et al. [29] and Chen et al. [30], is beneficial. But these works consider offloading as a stand-alone decision module independent of a learning-based policy optimiser. The POM is heavily integrated in SmartEdgeScheduler: its predictions directly influence the state space in the DRL and the MCDM priority scoring, thus being a full partner of the scheduling and not just a pre-processing filter.

Kim et al. [12] does not take the same architectural approach and is related to distributed scheduling in IoT networks by coalition formation that acts as collaborative policy learning. Although promising for privacy-preserving scenarios, it doesn't provide the same granularity as SmartEdgeScheduler for the latency-energy-SLA trade-off under heterogeneous edge-cloud resource constraints.

Finally, the key differentiators of SmartEdgeScheduler compared to prior literature are: (i) We do not apply any of the above mentioned methods individually, but it is a unified co-design of DRL, metaheuristic optimization and prediction based offloading; (ii) Dynamic switching between scheduling modes is adaptive and relies on real-time signals to guide the choice of scheduling mode; (iii) Multi-Criteria Decision-Making (MCDM) Module for Task Prioritization. It considers predicted latency, energy, and utilization simultaneously; and (iv) end-to-end validation of task completion time, resource utilization, energy efficiency and SLA compliance under scaled workloads up to 20,000 tasks.

5.2.2 Strengths of the Proposed Work in Light of Prior Literature

Proactive adaptability - Not reacting but correcting responses. SmartEdgeScheduler's most innovative and structurally important feature, compared to previous work, is its proactive and not reactive approach to scheduling. As for the methods of Zhou et al. [3] and Zhu et al. [10], they only respond to the current state of the system and thus cannot predict congestion or resource shortages before they occur. By leveraging historical patterns, SmartEdgeScheduler's LSTM-

based POM can forecast upcoming network latency and potential execution times over multiple time slots, enabling the scheduler to offload tasks before these effects are observed in system metrics. This is a qualitative shift from the state-of-the-art DRL-only methods, and is a contribution not found in the other reviewed papers.

Adaptive Mode Switching for Hybrid Robustness. A drawback of pure DRL schedulers, as pointed out in the works of Djigal et al. [2] and Iftikhar et al. [5], is that they tend to suffer convergence instability in non-stationary environments. SmartEdgeScheduler does so directly by using its error-threshold adaptive switching mechanism (controlled by ERL_t and θ), which enables the metaheuristic module when the quality of DRL convergence is not at an acceptable level. An interesting aspect of this design is that the framework remains robust under both stable and highly dynamic workloads, a characteristic not found in any single-mode scheduler in the surveyed literature.

Multi-objective co-optimisation. Several previous studies have optimised for a single key measure, such as energy efficiency [9, 14], latency reduction in 6G deployments [16], and QoS in mobile edge computing [24]. To achieve minimum latency, energy, and resource utilisation while maximising resource utilisation, SmartEdgeScheduler has formulated scheduling as a weighted minimisation problem, $C_s = w_1L_s + w_2(1/U_s) + w_3E_s$, and demonstrates simultaneous improvement across all four evaluation metrics in the experimental results. This is a real-world optimisation benefit because in real deployments there is no one optimisation metric that is the primary one.

Validated scalability. Numerous scheduling papers in the reviewed literature assess the performance on small / moderate-sized task sets, but don't consider scalability explicitly. The scalability analysis with SmartEdgeScheduler (Section 4.5) shows that, despite this, SLAs are maintained 96-97% of the time across vastly different scales, ranging from 5,000 to 20,000 tasks, offering users evidence of production-ready capability that previous research in Long et al. [22] and Nandhakumar et al. [21] lacked at similar scales.

Interpretable priority assignment. The priority-scoring module, using MCDM techniques, explicitly generates a priority score for each task

based on weighted predicted latency, resource utilisation efficiency, and energy cost. This contrasts with schedulers using DRL, which only provide implicit prioritisation of tasks in the learned policy's decision-making, making it very hard to audit or explain. This interpretability is a valuable practical trait in safety-critical or regulated deployment scenarios, which previous learning-based approaches for industrial automation and smart city infrastructure did not target.

5.2.3 Weaknesses and Limitations in Light of Prior Literature

Training overhead and real-time deployment tension. As with most other DRL-based schedulers, such as DeepEdge [11] and the A3C-based scheduler by Tuli et al. [36], the DRL component of SmartEdgeScheduler requires substantial initial training on historical task and network data, which poses a real operational challenge. The rule-based heuristics proposed by Abraham et al. [19] as well as the auction approach developed by Liu et al. [33] have much less initialisation cost due to being lightweight. While the current work only partially alleviates the training burden of SmartEdgeScheduler's DRL module, in edge contexts where compute budgets are extremely constrained or deployment time is of the essence, this could be a limiting factor.

Historical Data Quality (HQD) Dependency for Prediction. The predictive accuracy of the LSTM-based POM depends on the representativeness of historical network and workload information. However, when there are sudden network topology changes, unexpected workload spikes, cold deployments without data, or other situations, a predictive model may produce incorrect predictions, which negatively affect offloading decisions. Both Shakarami et al. [31] and Chen et al. [30] have the same problem regarding their offloading model, which is a structural problem of prediction-dependent scheduling and not a design problem only for SmartEdgeScheduler – which is not addressable in the current work.

Simulation-only validation. Only a Python-based simulation environment, with additional parameters derived from iFogSim, is used to evaluate SmartEdgeScheduler. This provides reproducibility and easy comparison; however, it doesn't fully address real hardware variability such as CPU throttling, memory contention,

wireless channel interference and device heterogeneity at scale. Again, like papers by Yamansavascilar et al. [11] and Kim et al. [12], simulation is being used, but relative to the deployment of the hardware in a complete testbed, performance margins (task completion time reduction up to 25% and resource utilisation gap up to 15%) are maximal and should be treated with caution until testbed hardware deployment studies are conducted.

Extensiveness of the adaptive learning module. With increasing task loads, up to 20,000 concurrent tasks, the computational costs of DRL policy update, LSTM inference, and metaheuristic optimisation search running in parallel may be too high for resource-limited edge nodes. This is an issue highlighted by Djigal et al. [2] and Iftikhar et al. [5] in the context of the scalability of complex AI-based schedulers in fog and edge environments. While it is indeed a valid gap for practical deployment, there is no suggestion here of a lightweight or distilled version of the framework for the edge nodes with very limited compute.

Lack of security and privacy issues. However, increasingly, task scheduling in edge-cloud settings has become bound up with data privacy and security issues, especially in IoT or smart city applications. To overcome this, Sellami et al. have particularly focused on combining Blockchain with DRL in a scheduling system for an SDN-IoT network, as the case presented in [14] illustrates. Adversarial Task Injection, Data Poisoning of LSTM Models, and Secure Multi-party Offloading are not considered in SmartEdgeScheduler and thus are not applicable in adversarial or even sensitive deployment scenarios.

6. CONCLUSION AND FUTURE WORK

The research question tackled in this work is not a marginal gap in the performance of the current schedulers but the lack of their ability to address the current challenges of dynamic edge-cloud deployment. While all three approaches are based on assumptions of predictable workloads, reliable convergence, and episodic optimisation, they are, respectively, static heuristics, isolated DRL models, and standalone metaheuristic approaches. Those broken assumptions have monetary costs that can be measured: a 75% SLA compliance and 1200 ms task completion time have been achieved with round-robin scheduling, and even the best one in the comparative evaluation has not got any

better than 94% SLA compliance and 900 ms task completion time. SmartEdgeScheduler responded with 97% SLA compliance, 750 ms, resource utilisation of 90% and architecturally earned energy consumption improvements of 170 J.

This can be traced in the ablation study. When the LSTM-based Predictive Offloading Model was removed, the completion time increased up to 920ms and the SLA compliance decreased up to 89%. Compliance was further reduced to 91% when adaptive learning was removed. The centrepiece of this paper's architectural take is that no single piece of this equation alone ensures success; rather, the results stem from how the pieces are integrated. Indeed, this holds true under pressure in this scalability analysis: SLA compliance stayed between 96 and 97 per cent under loads ranging from 5,000 to 20,000 tasks, and resource usage increased, without resulting in a collapse in performance.

These restrictions are genuine. The deployment is constrained by the overhead of DRL training. The dramatic and unexpected adjustment in the network that has not been observed before affects the predictive model. Validation is still ongoing in simulation, and reported margins are performance ceilings that need to be confirmed on a testbed. Future work will focus directly on them by using lightweight neural architectures and transfer learning for less training cost, online LSTM updating for distributional shift without dependency to the past history, and spanning multi-tier cloud architectures and heterogeneous classes of edge devices. The goal is to design a scheduler that will continue to make its guarantees even if the network misbehaves and the workload is most unpredictable; these are the most important conditions for production.

In this study, a falsifiable claim was that co-designing deep reinforcement learning with LSTM-based predictive offloading and metaheuristic optimisation on the same platform, while considering all the mentioned aspects of scheduling, would help achieve results that would be unsustainable via any partial or single-mode approach that considers only some of these aspects. This is experimental evidence to test that hypothesis; no longer against assertions.

H1 was that task completion time and energy use would be reduced by using proactive LSTM-based prediction to ensure that the data for the tasks is present where needed before the tasks are executed, thereby overcoming the overhead of

reactive reallocation. This is confirmed by the ablation study, which shows that ablating the Predictive Offloading Model resulted only in a 22.7% increase in task completion time (from 750 to 920 ms) and a reduction in SLA compliance from 97 to 89. Prediction is not a “hot seat”; it is load-bearing. H1 is supported.

H2 anticipated that the adaptive mode-switching would carry out well in the presence of convergence instability, when pure DRL schedulers break. This is confirmed by the comparative results, which show that pure DRL-derived prior methods, e.g., DeepEdge and DRL-Dispatcher without switching, were inferior across all metrics. Eliminating adaptive learning in the ablation reduced SLA compliance to 91% and increased completion time to 870ms, indicating that the switching mechanism alone contributes to the ablation. H2 is supported.

H3 argued that using Explicit MCDM-based priority scoring could enhance SLA compliance because critical tasks are assigned the appropriate priority to meet deadlines. With respect to SLA compliance, SmartEdgeScheduler has achieved 97%, compared to a best prior-literature result of 94%, and parts of this are due to the priority module's contribution to structuring task order. H3 is supported.

The ablation study supports the central result- the aggregate hypothesis, in which a component is removed to observe its performance individually, independently and measurably degraded with removal compared to the full model configuration.

SmartEdgeScheduler's performance is not just a collection of three techniques; it is the result of integrating them. That's what this paper contributes to the knowledgebase.

REFERENCES

- [1] boonhatai kruekaew and warangkhan kimpan. (2022). Multi-objective task scheduling optimization for load balancing in cloud computing environment using hybrid artificial bee colony algorithm with reinforcement learning. *Ieee*. 10, pp.17803 - 17818.
[Http://doi:10.1109/access.2022.3149955](http://doi:10.1109/access.2022.3149955)
- [2] hamza djigal, jia xu, linfeng liu and yan zhang. (2022). Machine and deep learning for resource allocation in multi-access edge

- computing: a survey. *Ieee*. 24(4), pp.2449 - 2494.
[Http://doi:10.1109/comst.2022.3199544](http://doi:10.1109/comst.2022.3199544)
- [3] xiaokang zhou, wei liang, ke yan, weimin li, kevin i-kai wang, jianhua ma and qun jin. (2023). Edge-enabled two-stage scheduling based on deep reinforcement learning for internet of everything. *Ieee*. 10(4), pp.3295 - 3304. [Http://doi:10.1109/jiot.2022.3179231](http://doi:10.1109/jiot.2022.3179231)
- [4] zhuo chen, peihong wei and yan li. (2023). Combining neural network-based method with heuristic policy for optimal task scheduling in hierarchical edge cloud. *Elsevier*. 9(3), pp.688-697. [Https://doi.org/10.1016/j.dcan.2022.04.023](https://doi.org/10.1016/j.dcan.2022.04.023)
- [5] sundas iftikhar, sukhpal singh gill, chenghao song, minxian xu, mohammad sadegh aslanpour, adel n. Toosi, junhui du, huaming wu, shreya ghosh, deepraj chowdhury, muhammed golec, mohit kumar, ahmed m. Abdelmoniem, felix cuadrado, blesson varghese, omer rana, schahram dustdar and steve uhlig. (2023). Ai-based fog and edge computing: a systematic review, taxonomy, and future directions. *Elsevier*. 21, pp.1-49. [Https://doi.org/10.1016/j.jiot.2022.100674](https://doi.org/10.1016/j.jiot.2022.100674)
- [6] mehdi hosseinzadeh, elham azhir, jan lansky, stanislava mildeova, omed hassan ahmed, mazhar hussain malik and faheem khan. (2023). Task scheduling mechanisms for fog computing: a systematic survey. *Ieee*. 11, pp.50994 - 51017. [Http://doi:10.1109/access.2023.3277826](http://doi:10.1109/access.2023.3277826)
- [7] behzad saemi, ali asghar rahmani hosseinabadi, azadeh khodadadi, seyedsaeid mirkamali, and ajith abraham. (2023). Solving task scheduling problem in mobile cloud computing using the hybrid multi-objective harris hawks optimization algorithm. *Ieee*. 11, pp.125033 - 125054. [Http://doi:10.1109/access.2023.3329069](http://doi:10.1109/access.2023.3329069)
- [8] jiangjiang zhang, zhenhu ning, muhammad waqas, hisham alasmay, shanshan tu and sheng chen. (2023). Hybrid edge-cloud collaborator resource scheduling approach based on deep reinforcement learning and multi-objective optimization. *Ieee*. , pp.1 - 14. [Http://doi:10.1109/tc.2023.3326977](http://doi:10.1109/tc.2023.3326977)
- [9] bassem sellami, akram hakiri, sadok ben yahia and pascal berthou. (2022). Energy-aware task scheduling and offloading using deep reinforcement learning in sdn-enabled iot network. *Elsevier*. 210, pp.1-21. [Https://doi.org/10.1016/j.comnet.2022.108957](https://doi.org/10.1016/j.comnet.2022.108957)
- [10] kaige zhu, zhenjiang zhang, feng sun and bo shen. (2022). Workflow makespan minimization for partially connected edge network: a deep reinforcement learning-based approach. *Ieee*. 3, pp.518 - 529. [Http://doi:10.1109/ojcoms.2022.3158417](http://doi:10.1109/ojcoms.2022.3158417)
- [11] baris yamansavascular, ahmet cihat baktir, cagatay sonmez, atay ozgovde, and cem ersoy. (2022). Deepedge: a deep reinforcement learning based task orchestrator for edge computing. *Ieee*. 10(1), pp.538 - 552. [Http://doi:10.1109/tmse.2022.3217311](http://doi:10.1109/tmse.2022.3217311)
- [12] do-yup kim, da-eun lee, ji-wan kim, and hyun-suk lee. (2023). Collaborative policy learning for dynamic scheduling tasks in cloud-edge-terminal iot networks using federated reinforcement learning. *Ieee*, pp.1-14. [Http://doi:10.1109/jiot.2023.3327495](http://doi:10.1109/jiot.2023.3327495)
- [13] hojjat baghban, amir rezapour, ching-hsien hsu, sirapop nuannimnoi and ching-yao huang. (2022). Edge-ai: iot request service provisioning in federated edge computing using actor-critic reinforcement learning. *Ieee*., pp.1-10. [Http://doi:10.1109/tem.2022.3166769](http://doi:10.1109/tem.2022.3166769)
- [14] bassem sellami, akram hakiri and sadok ben yahia. (2022). Deep reinforcement learning for energy-aware task offloading in join sdn-blockchain 5g massive iot edge network. *Elsevier*. 137, pp.363-379. [Https://doi.org/10.1016/j.future.2022.07.024](https://doi.org/10.1016/j.future.2022.07.024)
- [15] nelson sharma, aswini ghosh, rajiv misra and sajal k. Das. (2023). Deep meta q-learning based multi-task offloading in edge-cloud systems. *Ieee*., pp.1-18. [Http://doi:10.1109/tmc.2023.3264901](http://doi:10.1109/tmc.2023.3264901)
- [16] vibha jain, bijendra kumar and aditya gupta. (2022). Cybertwin-driven resource allocation using deep reinforcement learning in 6g-enabled edge environment. *Elsevier*. 34(8), pp.5708-5720. [Https://doi.org/10.1016/j.jksuci.2022.02.005](https://doi.org/10.1016/j.jksuci.2022.02.005)
- [17] li-tse hsieh, hang liu and yang guo, robert gazda. (2023). Deep reinforcement learning-based task assignment for cooperative mobile edge computing. *Ieee*., pp.1-15. [Http://doi:10.1109/tmc.2023.3270242](http://doi:10.1109/tmc.2023.3270242)
- [18] chuanfu zhang, jing chen, wen li, hao sun, yudong geng, tianxiang zhang, mingchao ji, and tonglin fu. (2024). A cloud-edge collaborative task scheduling method based

- on model segmentation. *Springer*. 13(81), pp.1-21. <https://doi.org/10.1186/s13677-024-00635-7>
- [19] olanrewaju l. Abraham, md asri bin ngadi, johan bin mohamad sharif, and mohd kufaisal mohd sidik. (2024). Task scheduling in cloud environment—techniques, applications, and tools: a systematic literature review. *Ieee*. 12, pp.138252 - 138279. <http://doi:10.1109/access.2024.3466529>
- [20] nasiru muhammad dankolo, nor haizan mohamed radzi, noorfa haszlinna mustaffa, nurul aida osman, danlami gabi, and muhammed nura yusuf. (2024). Enhanced task scheduling and resource allocation in edge-cloud continuum using modified flower pollination algorithm. *Ieee*. 12, pp.162299 - 162310. <http://doi:10.1109/access.2024.3490177>
- [21] aadharsh roshan nandhakumar, ayush baranwal, priyanshukumar choudhary, muhammed golec, and sukhpal singh gill. (2024). Edgeaisim: a toolkit for simulation and modelling of ai models in edge computing environments. *Elsevier*. 31, pp.1-14. <https://doi.org/10.1016/j.measen.2023.100939>
- [22] saiqin long, cong wang, weifan long, haolin liu, qingyong deng, and zhetao li. (2024). An efficient task scheduling algorithm in the cloud and edge collaborative environment. *Ieee*. 33(5), pp.1296 - 1307. <http://doi:10.23919/cje.2022.00.223>
- [23] xu zhao; guangqiu huang; ling gao; maozhen li and quanli gao; (2021). Low load dids task scheduling based on q-learning in edge computing environment . *Journal of network and computer applications*. <http://doi:10.1016/j.jnca.2021.103095>
- [24] shida lu; rongbin gu; hui jin; liang wang; xin li and jing li; (2021). Qos-aware task scheduling in cloud-edge environment . *Ieee access*. <http://doi:10.1109/access.2021.3072216>
- [25] kyuchang lee; bhagya nathali silva and kijun han; (2021). Algorithmic implementation of deep learning layer assignment in edge computing based smart city environment . *Computers & electrical engineering*. <http://doi:10.1016/j.compeleceng.2020.106909>
- [26] yunmeng dong; gaochao xu; meng zhang and xiangyu meng; (2021). A high-efficient joint 'cloud-edge' aware strategy for task deployment and load balancing . *Ieee access*. <http://doi:10.1109/access.2021.3051672>
- [27] cai, l., wei, x., xing, c., zou, x., zhang, g., & wang, x. (2021). *Failure-resilient dag task scheduling in edge computing*. *Computer networks*, 198, 108361. <http://doi:10.1016/j.comnet.2021.108361>
- [28] yanal alahmad; tariq daradkeh and anjali agarwal; (2021). Proactive failure-aware task scheduling framework for cloud computing . *Ieee access*. <http://doi:10.1109/access.2021.3101147>
- [29] qihe huang; xiaolong xu and jinhui chen; (2021). Learning-aided fine grained offloading for real-time applications in edge-cloud computing . *Wireless networks*. <http://doi:10.1007/s11276-021-02750-8>
- [30] miaojiang chen; tian wang; shaobo zhang and anfeng liu; (2021). Deep reinforcement learning for computation offloading in mobile edge computing environment . *Computer communications*. <http://doi:10.1016/j.comcom.2021.04.028>
- [31] ali shakarami; ali shahidinejad and mostafa ghobaei-arani; (2021). An autonomous computation offloading strategy in mobile edge computing: a deep learning-based hybrid approach . *Journal of network and computer applications*. <http://doi:10.1016/j.jnca.2021.102974>
- [32] kun cao; shiyan hu; yang shi; armando colombo; stamatis karnouskos and xin li; (2021). A survey on edge and edge-cloud computing assisted cyber-physical systems . *Ieee transactions on industrial informatics*. <http://doi:10.1109/tii.2021.3073066>
- [33] lnhuan liu; haibo ge; shun li; xutao chen; haiwen gong and yongjing cui; (2021). Resource allocation strategy based on improved auction algorithm in mobile edge computing environment . *2021 ieee 6th international conference on cloud computing and big data analytics (icccbda)*. <http://doi:10.1109/icccbda51879.2021.9442523>
- [34] qi zhang; lin gui; shichao zhu and xiupu lang; (2021). Task offloading and resource scheduling in hybrid edge-cloud networks . *Ieee access*. <http://doi:10.1109/access.2021.3088124>
- [35] muhammed tawfiqul islam; huaming wu; shanika karunasekera and rajkumar buyya; (2021). Sla-based scheduling of spark jobs

- in hybrid cloud computing environments .
Ieee transactions on computers.
[Http://doi:10.1109/tc.2021.3075625](http://doi:10.1109/tc.2021.3075625)
- [36] tuli, shreshth; ilager, shashikant;
ramamohanarao, kotagiri and buyya,
raj Kumar (2020). Dynamic scheduling for
stochastic edge-cloud computing
environments using a3c learning and
residual recurrent neural networks. Ieee
transactions on mobile computing, 1–1.
[Http://doi:10.1109/tmc.2020.3017079](http://doi:10.1109/tmc.2020.3017079)
- [37] fan qi; li zhuo and chen xin; (2020). Deep
reinforcement learning based task
scheduling in edge computing networks .
2020 IEEE/CIC International Conference on
Communications in China (ICCC).
[Http://doi:10.1109/iccc49849.2020.9238937](http://doi:10.1109/iccc49849.2020.9238937)
- [38] chen, zheyi; hu, junqin; chen, xing; hu, jia;
zheng, xianghan and min, geyong (2020).
Computation offloading and task scheduling
for dnn-based applications in cloud-edge
computing. Ieee access, 8, 115537–115547.
[Http://doi:10.1109/access.2020.3004509](http://doi:10.1109/access.2020.3004509)
- [39] khayyat, mashael; elgendy, ibrahim a.;
muthanna, ammar; alshahrani, abdullah;
alharbi, soltan and koucheryavy, andery
(2020). Advanced deep learning-based
computational offloading for multilevel
vehicular edge-cloud computing networks.
Ieee access, 1–1.
[Http://doi:10.1109/access.2020.3011705](http://doi:10.1109/access.2020.3011705)
- [40] hu, liangyan; sun, guodong and ren, yanlong
(2020). Coedge: exploiting the edge-cloud
collaboration for faster deep learning. Ieee
access, 8, 100533–100541.
[Http://doi:10.1109/access.2020.2995583](http://doi:10.1109/access.2020.2995583)