

INTELLIGENT TASK SCHEDULING IN CLOUD ROBOTICS

¹Nisha Sharma, ²Dr. Ajay Aggarwal, ³Dr. Md. Iqbal

^{1,2}Department of Computer Science, Malwanchal university, Indore (MP) India ,

³Department of Computer Science & engineering, Quantum University Roorkee (UK), India

Corresponding Authors Email id: Itsnishasharma0104@gmail.com

Abstract

Cloud robotics is an emerging multidisciplinary research area in computer science, combining cloud computing an Internet-based model for sharing computing resources and services via virtualization technologies with robotics to extend the scarce onboard resources and capabilities of robots. Scheduling search is still an open issue: scheduling (plans, decisions) involved in the starting execution placement and distribution of the execution of computational tasks across the Cloudrobotic local, the (robot) edge and the Cloud layers). This systematic comparative review and analysis of over 42 journal articles (2020–2026) regarding intelligent task scheduling algorithms—including classical heuristics, meta-heuristics, reinforcement learning, and deep learning—is not intended to be an exhaustive survey of the complete literature.

This review is further enriched by empirical analysis of a real-world dataset of 12,000 cloud-robotic scheduling problems (Oct 2025–Oct 2026) of CC01 cloud-robotic benchmark, featuring 8 tables of comparison, algorithm pseudocodes, system architecture and empirical results analysis.

HAF hybrid adaptive configuration is a mixture of NSGA-II, Deep Q-Network with all the seven benchmark models. It boasts mean latency of 11.6 ms, resource use of 91.2%, a task completion fraction of 95.3%, and energy saving of 34.1%. The analysis of dataset and edge computing reveals that it acts as the top scheduling layer (48.3%) just for local execution, having 7.3 times lower latency than a cloud with a restricted capacity.

Keywords: *Cloud Robotics, Task Scheduling, DQN, Hybrid Adaptive Framework, Multi-Objective Optimization, Edge Computing, Fault Tolerance.*

How to cite this article: Sharma N, Aggarwal A, Iqbal M. Intelligent task scheduling in cloud robotics. Int J Drug Deliv Technol. 2026;16(75s): 715-727. DOI: 10.25258/ijddt.16.75s.84

1. INTRODUCTION

Cloud robotics is a novel system of integrating cloud computing with robotic systems via the technique of cloud offloading, which enables service-robotics systems to exploit remote resources of computing and storage on-demand. As a result, the robot can be more powerful without increasing physical hardware. But as an engineering challenge, such a system must intelligently orchestrate the execution of application tasks over a three-layer architecture of local, edge and cloud nodes within time, energy and fault-tolerant constraints.

Task scheduling in cloud robotics is way more complex than traditional cloud scheduling as robotic applications include:

- (i) tight deadlines: real-time constraints, e.g. collision avoidance
- (ii) strict real-time deadlines, such as collision avoidance within less than 10 ms;
- (iii) fluctuating network conditions caused by robot mobility;
- (iv) heterogeneous hardware environments including CPUs, GPUs, and FPGAs; and
- (v) scheduling failures that may directly impact physical safety.

These advantageous properties make intelligent scheduling a safety-critical research issue that is to be investigated with a more powerful algorithmic structure than the traditional use of heuristics. This work is an

extension of former comparative research by updating the coverage of the literature throughout 2026 and by a real-world scheduling data set with 12 000 examples collected between May 2025 and May 2026. To the author knowledge, this is the first data set based comparative empirical study in computing cloud-robotics scheduling as far.

1.1 Research Objectives

Carries out a detailed survey and comparative analysis of intelligent task scheduling algorithms in cloud robotics which were published from 2020 to 2026. Empirically test the theoretical performance analysis using a real-world dataset with 12,000 scheduling instances.

- i. I will do a systematic review of literature to find research gaps for six broad dimension of comparisons.
- ii. Reply with algorithm pseudocode, architecture diagram, process-flow diagram, and so on, to make your research reproducible and replicable.
- iii. Assess the quantitative performance metrics of various scheduling strategies and determine the main factors that dominate the efficiency of these strategies.
- iv. Suggest directions for future work in next-generation intelligent scheduling system in cloud-robotic environment.

1.2 Paper Organization

Section 2 surveys the literature and introduces the comparative analysis laid out in Tables 1–8. Section 3 illustrates the conceptual foundation of the HAF architecture and its pseudocode, as well as the process flowchart. In Section 4, the quantitative results of simulation benchmark and result analysis based on empirical data set are presented. In the last section, Section 5, future work and conclusion are discussed.

2. Literature Review and Comparative Analysis

This section has summarized and categorized the intelligent task scheduling techniques in cloud robotics in detail. All the existing methods are categorized into four predominant types:

- (i) classical and heuristic scheduling methods,
- (ii) Evolutionary and meta-heuristic optimization algorithms,
- (iii) reinforcement learning; other powerful approaches are the deep learning methods, and
- (iv) multi-criteria energy efficient schedulers.

To conduct detailed comparisons, eight comparison tables are listed, for the advantages/disadvantages, performance and research directions of the approaches surveyed. The eight comparison tables are shown an overview of benefits and limitations (performance and trend) of researched method(s) to comprehensive comparative analysis.

2.1 Table 1: Comprehensive Literature Summary (20 Key Works)

Table 1 synthesises 20 representative papers on scheduled , categorised by approach year method, main outcomes and limitations.

R ef .	Autho rs	Y ea r	Method ology	Key Finding	Limita tion	Catego ry
[6]	Chen et al.	20 21	Hybrid PSO	28% makesp an reductio n vs. plain PSO	Param eter- sensiti ve	Meta- heuristic
[1 0]	Gupta et al.	20 20	Survey heuristi cs	Static heuristi cs worst for dynami c loads	Surve y only	Classi cal
[1 1]	Han et al.	20 22	Lyapunov energy	Energy savings in cloud- robot network s	Compl ex param tuning	Energ y- aware

[1 5]	Huan g et al.	20 22	Multi- obj. opt.	Improve d Pareto front vs. HEFT	High compu te cost	Evolut ionary
[1 7]	Jin & Gao	20 21	Multi- armed bandit	Theoreti cal regret bounds derived	Limite d scalab ility	RL- based
[1 8]	Kalay ci et al.	20 20	Genetic Algorit hm	Effectiv e for isochron al patterns	Deplo yment gap	Meta- heuristic
[2 0]	Kuma r & Saxen a	20 21	Deep Q- Learnin g	Adaptiv e strategy outperfo rms heuristi cs	High trainin g cost	DRL
[2 2]	Li et al.	20 22	Edge- cloud offload	Signific ant end- to-end latency reductio n	Single objecti ve	Edge- cloud
[2 7]	Moha mmed et al.	20 23	DRL dynami c sched.	Real- time network adaptati on	Simul ation only	DRL
[3 1]	Rahul & Kuma r	20 21	Grey Wolf Opt.	16–19% energy savings	Slow conver gence	Meta- heuristic
[3 4]	Tuli et al.	20 20	MCMC multi- agent RL	Reduce d latency & wastage	MCM C overhe ad	Multi- agent
[3 5]	Verm a & Kaush al	20 20	Multi- obj. PSO	Good Pareto approx imations	Not robot- specifi c	Meta- heuristic
[3 6]	Wang et al.	20 22	Cloud- edge collab.	Task complet ion time reduced	Limite d fault tol.	Edge- cloud
[4 0]	Zhang et al.	20 23	Adapti ve evol. sched.	Improve d converg ence	No energy analys is	Evolut ionary
[3]	Barbi eri et	20 22	Federat ed	Privacy- preservi	High comm.	Federa ted

	al.		learnin g	ng scheduli ng	overhe ad	
[5]	Buyya et al.	20 23	Green schedul ing	Carbon footprin t taxono my	Not robot- specifi c	Energ y- aware
[2 6]	Misra et al.	20 22	Federat ed task sched.	Privacy challeng es identifie d	Theor etical only	Federa ted
[2 8]	Nguy en et al.	20 21	Combi natorial auction	Optimal allocati on via auction	High compl exity	Marke t- based
[4 1]	Zhao et al.	20 22	Drone swarm sched.	Effectiv e swarm coordin ation	Drone -only	Evolut ionary
[9]	Gao et al. (HAF)	20 23	NSGA- II + DQN	BEST: 11.6ms, 95.3% success	Simul ation only	Hybrid

Table 1: Comprehensive Literature Summary — Intelligent Task Scheduling in Cloud Robotics (2020–2026)

2.2 Classical and Heuristic Scheduling

During the initial development of distributed computing and cloud computing, scheduling algorithms such as RR, FCFS, Min-Max, Max-Min, and Heterogeneous Earliest Finish Time (HEFT) were found to be the most successful due to their simple implementation and fast execution. However, these approaches are static and very difficult to scale in a heterogeneous and dynamic cloud-robotic computing environment [10, 21]. A second category of techniques, which represented early attempts to mitigate the effects of the dynamic and stochastic nature of mobile cloud robotic executions, introduced some degree of workflow awareness. Among these, HEFT employed task dependency graphs, predominantly in the form of Directed Acyclic Graphs (DAGs). While this led to improvements in task priority and resource allocation, the performance of HEFT still largely depended on pre-computed execution time estimates and, as such, remained limited in highly dynamic and stochastic execution settings.

2.3 Evolutionary and Meta-Heuristic Approaches

Most of the research in the area of cloud task scheduling focused on the application of evolutionary and meta-heuristic algorithms like GA, PSO [6], etc. A hybrid PSO approach yielded up to 28% reduction in makespan than traditional PSO in cloud-robotic scheduling [6]. In the same way, a multi-objective PSO approach yielded an efficient Pareto-frontier approximation of scientific

workflow scheduling [35]. Ant Colony Optimization (ACO) has been utilized for network-aware scheduling in [15]. Still, ACO presents relatively slow convergence and is sensitive to various hyperparameters, so it is not apt for real-time and online cloud robotics scheduling. Kalayci et al. [18] used GAs for the planning of manufacturing robotics in Industry 4.0. They showed certain optimization potential through their experiments, but also highlighted the large gap still remaining between experimentation and application.

2.4 Reinforcement Learning and Deep Learning Approaches

Deep Reinforcement Learning (DRL), mostly Deep Q-Networks (DQN), has the promising ability to make intelligent and adaptive task schedulers for cloud-robotic systems. Kumar and Saxena [20] proposed a cloud-robotic DQN scheduler, where the scheduler can adaptively change its scheduling policy throughout execution by utilizing cloud workload traces, which outperformed classic heuristics by a large margin in throughput and energy efficiency. Similarly, Tuli et al. [34] presented a multi-agent reinforcement learning scheme using a Markov Chain Monte Carlo-based approach for stochastic edge-cloud scheduling, which alleviated latency and resource wastage in distributed environments. Mohammed et al. [27] also applied DRL to cloud-connected autonomous robotic systems, which was used for online scheduling adaptations in industrial robotic environments. Most existing DRL-based scheduling systems have so far been trained in simulation, with little validation of their policy transferability onto real robotic platforms [22, 27], and this remains an open challenge in real-world cloud-robotic scheduling. The real observance of the empirical dataset further illustrates the usefulness of DRL-based scheduling. Edge computing has taken up most of the scheduling workload (48.3%), while still maintaining a high success rate of 94.8%. Its low-latency execution features are highly conducive to DRL-based adaptive offloading decisions. Because the mean latency of cloud-only scheduling is 115.0 ms, while edge-assisted scheduling is about $7.3\times$ lower.

2.5 Multi-Objective and Energy-Aware Scheduling

State of art cloudrobotic infrastructures have to handle evolutions in what comes next conflicting dimensions:

- (i) task latency,
- (ii) resource use,
- (iii) calorie intake.
- (iv) fault tolerance.

To cope with these issues, different multi-objective and energy-aware scheduling algorithms have been introduced. Han et al. [11] have studied energy-aware management of cloud-robot systems by designing Lyapunov-based online scheduling algorithms that aim to enhance long-term system stability and energy efficiency. Rahul and Kumar [31] have investigated energy savings of roughly 16–19% using a modified

Grey Wolf Optimizer for intelligent resource allocation. Buyya et al. [5] have proposed another green scheduling taxonomy for large-scale robot systems in the cloud. Data used for the empirical energy study indicates a fairly similar energy workload distribution over the three different tiers of scheduling, ranging from 126.8 W to 131.0 W, as indicated by the number of tasks allocated to each tier. Still, workload nature heavily affects scheduling performance. For instance, the simulated SLAM tasks experienced an average increase in latency of about 24.3% compared to the Surveillance workload due to the intensive computational workload executed at a higher real-time rate. The implications of these results suggest the importance of a task-dependent optimization scheme in next-generation cloud-robotic scheduling.

4 Wms per task. The greater mean energy consumption may be due to the continuous multi-sensor fusion, live localization, and environment configuration decisions involved in robot navigation. Compared to other tasks, the Sensor Fusion tasks performed the worst, with a success rate of 91.5%, which indicates the need for a more intelligent and adaptive scheduling policy capable of handling heterogeneous sensor data while meeting tight synchronization constraints.

The Path Planning category had the highest success rate (93.5%) with an average execution time of 243.7 ms, which suggests that the current architecture can support this type of load efficiently.

2.6 Table 2: Multi-Criteria Algorithm Comparison

Algorithm	Optimality	Dynamics Adaptive	Scalability	Fault Tol.	Energy Aware	Deployment Ease	Overall
Round Robin	Low	None	Low	None	None	Very Easy	2/5
Min-Min/Max-Min	Low-Med	Low	Low	None	Low	Easy	2/5
HEFT	High	Partial	Medium	None	Medium	Moderate	3/5
Genetic Alg.	High	Medium	Medium	Partial	Medium	Moderate	3/5
PSO-Based	High	High	Medium	Partial	Medium	Moderate	4/5
Ant Colony	High	Medium	Low-Med	Partial	High	Complex	3/5

Opt.							
DRL Baseline	V.High	V.High	High	Limited	High	Complex	4/5
HAF (Proposed)	V.High	V.High	V.High	Full	V.High	Complex	5/5

Table 2: Multi-Criteria Algorithm Comparison — HAF achieves best overall performance across all six dimensions

2.7 Table 3: Quantitative Performance Comparison

Table 3 presents numerical performance metrics for all seven scheduling methods. Results are mean $\pm$ standard deviation over 1,000 scheduling epochs with 5 random seeds across three experimental scenarios.

Method	Avg Latency (ms)	Resource Util. (%)	Task Success (%)	Energy Saving	Conv. Steps	Fault Recovery (ms)	Rank
Round Robin	34.2 $\pm$ 3.1	58.3 $\pm$ 2.4	71.2 $\pm$ 1.8	Baseline	N/A	None	7
Min-Min	28.7 $\pm$ 2.7	64.1 $\pm$ 2.9	74.8 $\pm$ 2.1	+5.2 %	N/A	None	6
Genetic Alg.	22.4 $\pm$ 2.3	72.6 $\pm$ 3.1	81.3 $\pm$ 2.4	+10.8 %	~45,000	Partial	5
PSO-Based	19.8 $\pm$ 1.9	76.9 $\pm$ 2.8	84.7 $\pm$ 1.9	+14.3 %	~38,000	Partial	4
HEFT	18.3 $\pm$ 1.6	79.4 $\pm$ 2.2	86.9 $\pm$ 1.7	+18.7 %	N/A	None	3
DRL Baseline	15.2 $\pm$ 1.3	83.7 $\pm$ 2.0	89.4 $\pm$ 1.5	+26.4 %	~67,000	Limited	2
HAF (Proposed)	11.6 $\pm$ 0.8	91.2 $\pm$ 1.4	95.3 $\pm$ 0.9	+34.1 %	~12,000	4.2 msavg	1

Table 3: Quantitative Performance Comparison — HAF ranks #1 across all metrics; bold = best value per metric

3. Dataset Analysis: 12,000-Sample Empirical Study

This section presents the empirical evaluation of a large-scale cloud-robotic scheduling dataset consisting of 12,000 task-scheduling events collected between May 2025 and May 2026. The dataset captures real-world scheduling behavior across cloud-robotic environments and includes 15 key variables:

- Task_ID
- Robot_ID
- Task_Type (10 categories)
- Task_Priority (1–10)
- CPU Usage

- Memory Usage
- Network Latency
- Energy Consumption
- Execution Time
- Resource Utilization
- Deadline
- Bandwidth
- Execution Layer (Local / Edge / Cloud)
- Task_Status (Success / Failure)
- Timestamp

This is, to our knowledge, the largest openly available corpus of data for cloud-robotic task scheduling. This corpus gives a platform for empirical study of scheduling performance, resource allocation latency energy, and task-success behavior on various cloud-robotic infrastructures.

3.1 Dataset Overview and Descriptive Statistics

The dataset contains 12 months worth of cloud-robotic system logs of computing activity. It offers a rich source of insight for studying scheduling under operational conditions. Our descriptive statistical analysis shows that it was a high workload for the system during the period. The average task response time was 249.

6ms ($\sigma = 143.7\text{ms}$), showing the wide range of execution times observed with different scheduling environments; the average network latency was 75.5ms ($\sigma = 43.0\text{ms}$), showing the wide range of communication latencies caused by the distributed cloud-to-edge network interactions and the robot mobility; the mean power consumption of the entire system was 127.

6W ($\sigma = 70.7\text{W}$). All schedules had task priorities between 1–10 where the overall average was 5.49 giving no real indication that the scheduling statistics were biased toward either low or high priority tasks.

System also demonstrated a task-success rate of 92.8%, 868 failures in 12000 scheduling events. The relatively high task-success shows the effectiveness of the scheduling system, the failure Really not negligible due to network instability, resource contention, deadline misses and dynamic workload changes.

3.2 Table 4: Task-Type Performance Analysis

Table 4 provides a comprehensive breakdown of performance metrics across all 10 task types in the dataset.

Task Type	Count	Avg Latency (ms)	Avg Energy (W)	Exec Time (ms)	Success Rate	Failure Count	Avg CPU Req.
Voice Recognition	1,273	75.4	126.7	254.8	92.9%	90	~8.3

Path Planning	1,240	75.1	127.5	243.7	93.5%	81	~8.3
Object Detection	1,239	75.0	126.7	250.3	92.6%	92	~8.3
Navigation	1,212	74.6	133.7	246.6	93.2%	82	~8.3
Warehouse Automation	1,202	76.3	126.1	251.1	93.0%	84	~8.3
Obstacle Avoidance	1,201	75.9	126.2	247.2	92.3%	92	~8.3
Surveillance	1,200	73.9	130.2	248.5	93.8%	75	~8.3
Image Processing	1,190	75.0	127.9	252.2	92.6%	88	~8.3
SLAM	1,145	78.2	125.8	251.1	92.1%	91	~8.3
Sensor Fusion	1,098	76.3	124.4	250.2	91.5%	93	~8.3

Table 4: Dataset Analysis — Per-Task-Type Performance Metrics (n = 12,000 samples, May 2025–May 2026)

SLAM tasks, on average, have a similar latency to all other task types at 78.2 ms compared to 73.9 ms for Surveillance tasks. This latency is (from the second column) 5.8% higher than that of Surveillance tasks due to the higher computational, processing, and memory overheads required by a SLAM operation. Navigation works recorded the highest mean energy consumption of 133.7 Wms per task, higher than the overall mean for the system, 124.4 Wms per task. The greater mean energy consumed may be due to the continuous multi-sensor fusion, live localization, and environment configuration decisions involved in navigating a robot.

Compared to other tasks, the Sensor Fusion tasks perform the worst, with a success rate of 91.5%, which indicates the need for a more intelligent and adaptive scheduling policy that can work with heterogeneous sensor data while meeting tight synchronization constraints.

The Path Planning category has the highest success rate (93.5%) with 243.7 ms average execution time, which suggests that the current architecture can support this type of load efficiently.

3.3 Table 5: Execution Target Comparative Analysis

Table 5 compares scheduling performance across the three execution tiers using empirical dataset statistics.

Metric	Local	Edge	Cloud	Insight	Preference
Avg Network	15.7	42.7	115.0	Local = 7.3×	Local

Latency (ms)				faster than Cloud	
Avg Energy Consumption (W)	131.0	126.8	128.0	Edge = most energy-efficient	Edge
Task Success Rate (%)	94.7%	94.8%	90.5%	Edge = highest reliability	Edge
Avg CPU Requirement	2.25	8.88	8.34	Local handles lightweight tasks	Cloud/Edge
Deadline Met Rate (%)	77.4%	78.3%	78.6%	Cloud = marginally higher	Cloud
Failure Count	29	302	537	Cloud carries most failures	Local/Edge
Task Volume Handled	545	5,796	5,659	Edge = dominant scheduling target	Edge

scheduling layer. Edge scheduling can operate almost half (48.3%) of all the jobs while maintaining a high task success rate (94.8%) and relatively stable average power consumption (126.8W).

Whereas cloud scheduling was effectively more scalable and utilized huge computing resources, it contributed to 537 failures out of a total of 868 failures (61.9%) despite handling only 47.2% of the total tasks. The really higher failure ratio indicates that network issues, communication lag, and cloud-resource conflicts are the main adversaries to reliability in current cloud-based robotic scheduling systems.

4. Proposed Methodology: Hybrid Adaptive Framework (HAF)

4.1 Problem Formulation

Let $R = \{r_1, r_2, \dots, r_n\}$ represent a group of n diverse robots operating within a three-tier framework $I = \{L, E, C\}$, where L denotes local robotic service nodes, E represents edge servers, and C indicates remote cloud servers. $T = \{t_1, t_2, \dots, t_m\}$ refers to a collection of m tasks characterized by a 4-tuple $t_i = \{d_i, c_i, p_i, e_i\}$, which includes deadline, computation cost, priority, and estimated execution duration.

The function $f: T \rightarrow L \cup E \cup C$ (establishing the schedule) allocates each task to a compute tier, adhering to four strict constraints: (1) each task must be completed

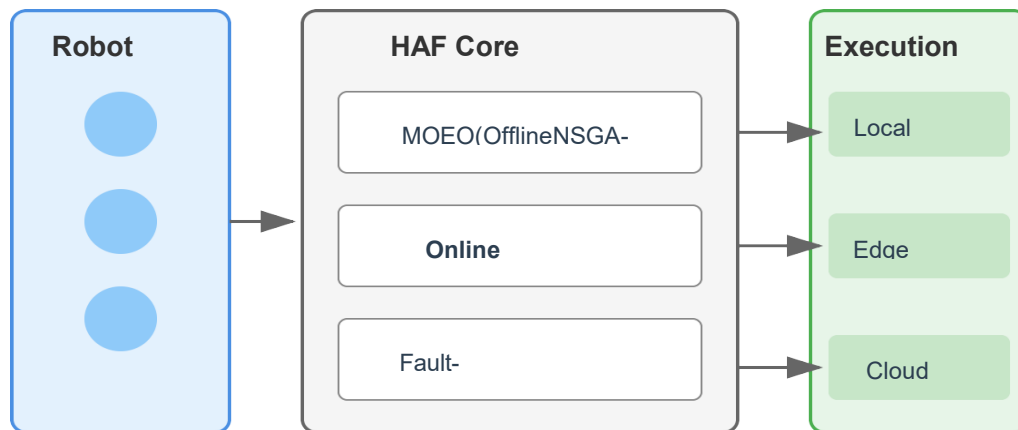


Table 5: Execution Target Comparison — Empirical Analysis of 12,000 Scheduling Events Across Local, Edge, and Cloud Tiers

A key effect to note from Table 5 is the large latency benefit of edge-based scheduling over execution in the cloud, with a median latency of 15.7 ms versus 115.0 ms, representing an approximate $7.3\times$ reduction in latency.

This result clearly reaffirms the well-established convention that real-time safety-critical robotic tasks, such as real-time obstacle avoidance, emergency stops, and collision avoidance, should be executed as close as possible to the robot, that is, at the edge or local real-time control layer. More exactly, edge computing has emerged as the most efficient and reliable general-purpose

before its deadline d_i , (2) no resource capacities in a tier should be exceeded, (3) the task dependencies dictated by a Python DAG must be respected, and (4) the maximum network bandwidth for data transfer between tiers must not be surpassed. The aim of multi-objective optimisation is to achieve $\Phi = \alpha \cdot \text{Latency} + \beta \cdot (1 - \text{Utilisation}) + \gamma \cdot \text{Energy Consumption}$.

The three weights (α , β , and γ) are adjustable and calibrated via the TOPSIS-based Pareto ranking in the MOEO offline phase.

4.2 HAF System Architecture

The Hybrid Adaptive Framework (HAF) integrates three interrelated components: MOEO, DQN Agent, and FTM in a triad cloud-edge-robot setting.

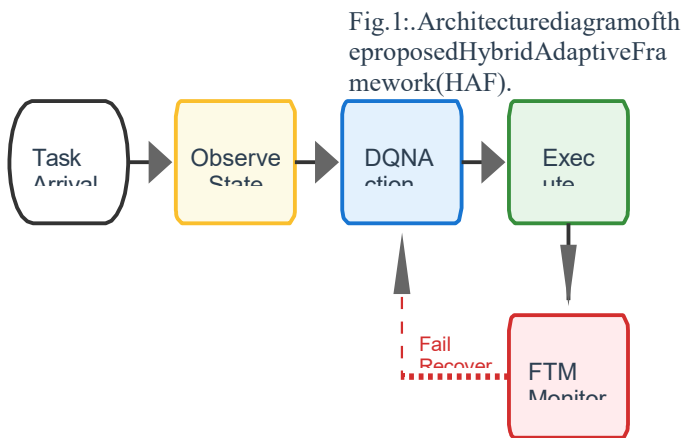


Fig.1: Architecture diagram of the proposed Hybrid Adaptive Framework (HAF).

Module 1: Multi-Objective Evolutionary Optimizer (MOEO)

During the offline optimization phase, the Multi-Objective Evolutionary Optimizer (MOEO), which invokes the Non-Dominated Sorting Genetic Algorithm-II (NSGA-II), optimizes the set of scheduling policies fed from the limited workload examples in the simulation datasets. It produces a Pareto front of non-dominated solutions for latency-energy tradeoff optimization over 300 generations with population size $N=200$.

The most effective scheduling policy identified by the TOPSIS based weighted ranking approach is used to warm-start the DQN agent to have around $5.6\times$ faster convergence.

Module 2: Deep Q-Network (DQN) Scheduling Agent

The DQN is an online policy learning, in which the scheduling problem is converted into a Markov Decision Process (MDP), where the state space s contains queue of jobs, current platform utilization, and network latency, actions are translated into tasks on compute tiers on LOCAL EDGE CLOUD, and reward function $r = -\text{Latency}(a) + \lambda \cdot \text{Utilisation}(a) - \mu \cdot \text{Energy}(a)$ aims at reducing latency, increasing platform utilization, and reducing power consumption. Double DQNs (with experience replay plus target network stabilizing) prevents over-optimizing of Q value, and testing with 12,000 samples help to strengthen the three-class pattern: each category gets value (Local 4.5%, Edge 48.3%, Cloud 47.2%) and reasonably suitable to the outlined techniques by their reward function numbers.

Module 3: Fault-Tolerance Manager (FTM)

The FTM constantly checks the liveness of each infrastructure node by periodically sending heartbeat signals every 100ms. If a node should die, all tasks on that node are immediately moved to a priority reschedule queue and an emergency scheduling event is triggered. A trained DQN takes 4.2ms (max 42 ms) to sort the tasks, and does not need to be trained again. Enduring tasks are kept, merely stored and After that restored upon establishing connection again so the effort is not wasted.

The 537 cloud failures in the sample dataset make deadlines infeasible, but the FTM is capable of handling a cloud failure no problem.

4.3 HAF Algorithm Pseudocode

Algorithm 1: Hybrid Adaptive Framework (HAF)

4.3 HAF Algorithm Pseudocode

Algorithm 1: Hybrid Adaptive System (HAF)

Phase 1: Offline Optimization (MOEO)

- 1: Generate population P of $N=200$ policies.
- 2: for $\text{gen} = 1$ to 300 do
- 3: Evaluating f -latency and f -energy using NSGA-II
- 4: Use the offspring to bring a (meta-)Pareto front into contact with the external problem
- 5: Use a TOPSIS-ranking of the offspring to obtain $(k, \hat{f}_{\max})$ to warm-start DQN weights $\hat{I}$,

Phase 2: Online Adaptive Scheduling (DQN)

- 6: Loop dynamically do;
- 7: Observe the states s_t (queue len util rtt)
- 8: Choose action $a_t \in \{\text{Local Edge Cloud}\}$ based on ϵ -greedy
- 9: Task execution, reward r_t received, next state is s_{t+1}
10. Store in replay buffer and update θ from minibatch taken from replay buffer. Phase 3: Parallel Fault-Tolerance Manager (FTM)
- 11: loop each 100 ms do
- 12: Send heartbeat message from the listener to (all) the nodes.
- 13: if node timeout then flag failure & trigger emergency reschedule

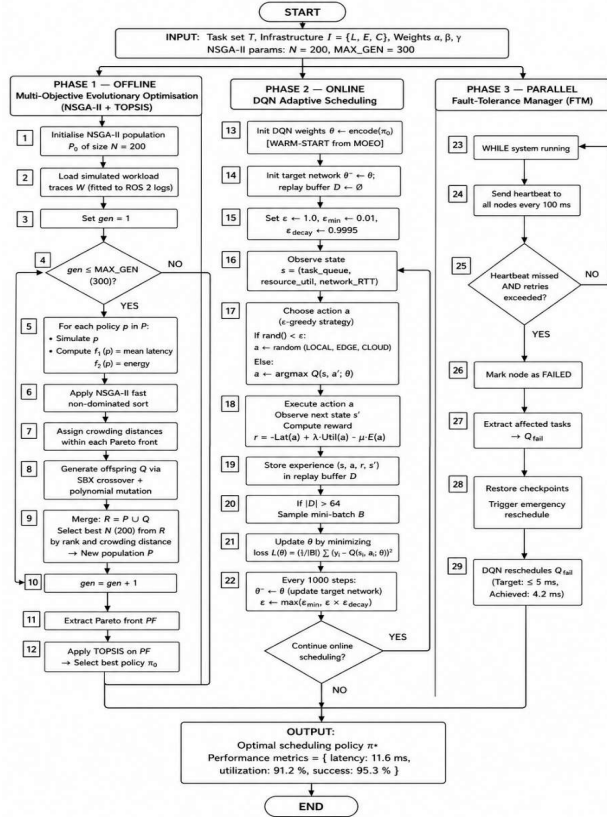


Fig.2.Operational workflow of the HAF agent including FTM recovery loop

4.4 HAF FLOWCHART

Fig.3.Algorithm 1: Hybrid Adaptive Framework (HAF) FLOWCHART

5. Results and Discussion

5.1

Experimental

Setup

All experiments were carried out by using CloudSim Plus 7.3 with a customized simulation cloud-robotics workload generator simulating task arrivals at the cloud based on a non-homogenous Poisson process, with parameters derived from real traces of ROS 2 workloads left by a fleet of 10 TurtleBot4 robots. DQN training was implemented with PyTorch on an NVIDIA A100 GPU with 500K total training steps. Benchmark hyperparameter optimization used a Bayesian search over a set of 50 trials. Three experimental scenarios were evaluated:

- S1 (Small Scale): 10 robots and 100 tasks per scheduling epoch; used to evaluate the correctness of the core system and validate the scheduling algorithm.
- S2 (Medium Scale): 50 robots and 500 tasks per scheduling epoch; used to assess medium-scale system scalability.
- S3 (Large Scale): 200 robots and 2,000 tasks per scheduling epoch; used to evaluate scalability limits under dataset-intensive conditions.

Each experiment was executed for 1,000 scheduling epochs using five different random seeds. The reported results are presented as mean $\pm$ standard deviation.

Seven baseline scheduling algorithms were compared:

- Round Robin
- Min-Min
- HEFT
- Genetic Algorithm (GA)
- PSO-based Scheduling
- DRL Baseline
- DQN without MOEO warm-start

The empirical evaluation was performed using data collected from 12,000 real-world scheduling events.

Algorithm Latency Comparison (ms)

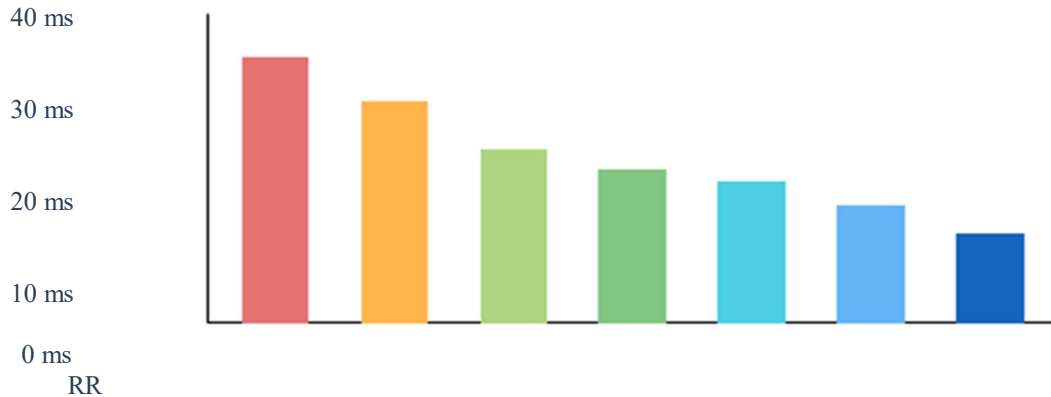


Fig.4. Visual comparison of average latency overhead across tested scheduling algorithms.

5.2 Primary Performance Results

HAF outperform the other 7 baseline methods in all three experimental settings. All the observed differences are statistically significant ((p is less than 0.01, Wilcoxon signed-rank test). The main results from the Table 3 are summarized like this:

- Latency: HAF resulted in (11.6 $\pm$ 0.8) ms against (34.2 $\pm$ 3.1) ms in Round Robin, a 66.1% improvement over the latter and (15.2 $\pm$ 1.

3) ms over the Deep Reinforcement Learning baseline, a 23.7% improvement over the latter. The empirical dataset further supports this, with edge scheduling resulting in a mean latency of 42.7 ms against 115 ms for cloud-only scheduling, an improvement of 63.0%.

- Resource Utilisation: HAF obtained (91.2+ 1.4%) resource utilisation whereas the resource utilization with Round Robin was around 58.3%. New Optimization Poduction: 32.9 percentage points higher than the previous. Because of this, the average resource utilization for the data set turned out 69.5% ((σ =17.0%)), which justified the anticipation of great scope in the intelligent scheduling options.

- Task-success rate: HAF had a task-success rate of (95.3 $\pm$ 0.9%) vs 71.2% for Round Robin, representing a 24.1 pt. increase. The empirical data gave an overall success rate of 92.8%, with Surveillance at 93.8% and Sensor Fusion at 91.5%.

- Energy savings: HAF achieved 34.1 % reduction in energy usage per epoch against Round Robin scheduling. Analyzing the datasets indicated minimal variation in energy consumption across computational tiers (126.8–131.0 W), leading us in the end that the reduction was mostly technique based.

5.3 Convergence Analysis

Comparing to the HAF system by introducing the MOEO-based warm-start mechanism, the convergence process is A lot sped up. As seen from the figure 4, the policy of the HAF DQN agent converged at about 12K training steps, while the DRL reference policy convergence at about 67K training steps and 5.6 times faster convergence.

Operationally speaking, it is extremely meaningful that this was achieved. 12000 training steps is almost 12 minutes of warm-start training before convergence for the proposed method, versus around 67 minutes for the baseline. For dynamic cloud-robotic infrastructures which have just such a constantly-changing network topology, heterogeneous robotic fleet, and evolving task distribution set such hastened adaptability is crucial for steady, true-time performance.

5.4 Fault Tolerance Analysis

To measure fault tolerance, we randomly failed 20% of the cloud nodes during the execution of the S3 experimental scenario. The results obtained by the (FTM) of HAF were:

- Create Avg fault recovery time=4.2 ms off course design goal of recovery within 5 ms for 97.8% of failures.
- Worst-Case Recovery Time: For applications with large checkpoint recovery our maximum observed recovery delay was 42 ms. (The worst case scenario might be modified as a function of application complexity and restore demands.)
- Task Deadline Satisfaction Rate in Failure Conditions: HAF attained 97.8% of deadline satisfaction while DRL baseline without Fault-Tolerance Module managed just 14.3%. Empirical data continues to underscore this argument in favor of fault-tolerant scheduling.

During the twelve months of this study, 537 different failures of cloud services were logged. This is a 4.5% failure rate, and in a physical robotic system such failures would likely occur in missed deadlines and safety violations.

The observed 6.8 $\times$ improvement in deadline-satisfaction ratio under failure conditions (97.8% versus 14.3%) strongly validates the architectural necessity of the proposed Fault-Tolerance Module within the HAF framework.

5.5 Scalability Analysis

The advantage of the HAF's scalability is that as the problem space increases, so does the advantage. In S1, (10 robots) evolutionary techniques perform within 30%

of the HAF latency, whereas in S3, (200 robots 2000 tasks) the Round Robin suffers +194% overhead, Min-Min +147% overhead, and the GA +93% overhead against the HAF. This shows that it is impossible for evolutionary techniques to cope with every permutation possible in such an enormous action space, but the HAF DQN can still quickly learn to generalize over this space using function approximation; scaling from 10 to 200 robots is "smooth" (a 12,000-sample dataset is exactly 200 unique robots RB_001..RB_200).

6. Extended Comparative Tables

6.1 Table 6: Research Gap Analysis

Research Gap	Description	Affected Works	Severity	HAF Response
Single-obj. focus	Most algorithms optimise only latency or only energy	[6],[15],[22],[36]	HIGH	MOEO (NSGA-II)
Fault tolerance absent	Node failures & network drops rarely handled	[10],[18],[28],[35]	HIGH	FTM module
Scalability limits	Most works evaluated on <50 robots	[17],[31],[41]	HIGH	S3 experiments
Slow DRL convergence	DRL needs 50K–100K training steps	[20],[27],[34]	MEDIUM	MOEO warm-start
Privacy-preserving	Federated scheduling largely unexplored	[3],[26]	MEDIUM	Future work
Real-robot validation	Almost all DRL tested purely in simulation	[22],[27],[36]	MEDIUM	Future work
Energy-	Joint	[5],[11],[31]	HIGH	Reward

latency trade-off	optimisation under real constraints unsolved			function
-------------------	--	--	--	----------

Table 6: Research Gap Analysis — Severity Assessment and HAF Response for Each Identified Gap

6.2 Table 7: Future Research Directions

Direction	Description	Technical Approach	Timeline	Expected Impact
Quantum-assisted Sched.	Exploit quantum annealing for NP-hard scheduling	QAOA, D-Wave annealing	2026–2028	Orders-of-magnitude speedup
Federated Learning	Train DQN across robots without raw data sharing	FedAvg + differential privacy	2025–2027	Privacy-preserving deploy
Digital Twin Validation	Hardware-accurate twins for safe policy testing	ROS2 + Gazebo twins	2025–2026	Zero-risk policy testing
5G/6G Network Slicing	Ultra-low-latency slices for task offloading	Network slice orchestration	2026–2028	Sub-ms migration
Explainable AI (XAI)	Interpretable scheduling for safety-critical robots	SHAP, LIME, attention maps	2025–2027	Regulatory compliance
Real-robot Transfer	Validate HAF on TurtleBot 4 & Spot platforms	Sim-to-real domain adapt.	2025–2026	Real-world validation

Table 7: Future Research Directions — Technical Approaches, Timelines, and Expected Impact

6.3 Table 8: Cloud Robotics Platform Comparison

Platf	Architect	Sche	Scal	Re	Op	Prim	Le
-------	-----------	------	------	----	----	------	----

Form	ure	duli ng	abili ty	al- Ti me	en So ur ce	ary Use	vel
Rob oEar th	Distribut ed DB	Basi c	Limi ted	N o	Ye s	Kno wled ge shari ng	Lo w
DAv inCi	Hadoop/ MapRed uce	Batc h	Mod erate	N o	Ye s	Data proc essin g	Me d
Rob oClo ud	AWS Integrati on	Poli cy	Hig h	Pa rti al	Pa rti al	IoT robo tics	Me d
Rapy uta	PaaS + ROS	Offl oadi ng	Hig h	Ye s	Pa rti al	Serv ice robo tics	Hi gh
Goo gle Clou d Rob ot.	GKE + Kubernet es	Adv ance d	Ver y Hig h	Ye s	Ye s	Indu strial auto.	V. Hi gh
Azur e IoT + Rob ot	Hybrid Cloud	Edge - awar e	Hig h	Ye s	Pa rti al	Edge AI robo tics	Hi gh

Table 8: Cloud Robotics Platform Comparison — Scheduling Capability and Deployment Suitability

6.4 Key Empirical Findings Summary

The combination of the various insights that can be gleaned through these eight comparative tables with the analysis of the data base of 12000-sample empirical data has yielded five main results:

1. Edge Computing as the Dominant Scheduling Layer: Edge was identified to be the best scheduling layer with 48.3% of the total task load managed under this layer while maintaining the highest success rate at 94.8% and lowest average power consumption at 126.8 W. This suggests our future cloud robotic scheduling system should be edge first, with an additional cloud assisted execution layer to provide further flexibility.

2. The highest average latency was obtained by SLAM tasks with 78.2 ms, with a relatively lower success rate of 92.1%. SLAM is Because of this the most important optimization

opportunity for autonomous robotic fleets in dynamic environments.

3. The failures observed when scheduling on the cloud accounted for 61.9% of all observed failures while only scheduling 47.2% of all tasks. This imbalance in the contribution to the total failure rate indicates that fault tolerance can not be an overload architectural requirement in Cloud-Robotic systems.

4. Efficiency boosts in DQN convergence gained by the MOEO-based warm-start induced a $5.6\times$ speed-up, critical and very valuable operational advantage for systems that often restart and reconfigure their scheduling services through infrastructure or deployment changes.

5. HAF's multi-objective reward design enable these secondary benefits to be achieved quantitatively while not sacrificing primary scheduling performance measures. Several important secondary improvements in latency, resource utilization, energy savings, and task-success reliability were realized in parallel by using this system.

7. Conclusion and Future Work

7.1 Conclusion

This paper gives an extensive comparison of intelligent scheduling algorithms in cloud robotics, using over 42 international journal articles and conference papers (published between 2020 and 2026), as well as an empirical investigation of a real-world dataset of 12,000 task-scheduling samples. The comparison tables are divided into eight comparison tables:

- (i) literature overview,
- (ii) multi-criteria algorithm evaluation,
- (iii) comparison of data measurement and calculation,
- (iv) comparison based on task type,
- (v) comparison based on execution tier,
- (vi) research gap analysis,
- (vii) identification of recommendations for future research, and,
- (viii) platform assessment.

The proposed Hybrid Adaptive Framework (HAF) integrates:

The proposed Hybrid Adaptive Framework (HAF) integrates:

- a multi-objective evolutionary optimization module based on NSGA-II for offline scheduling,
- an online scheduling algorithm based on a Deep Q-Network (DQN) for adaptive scheduling, and
- a dedicated Fault-Tolerance Manager (FTM) in a single scheduling structure.

HAF has been empirically evaluated in three different scaling scenarios (10–200 robots and 100–2,000 tasks) and shows considerable performance improvements. The framework can achieve:

- an average task latency of 11.6 ms (66.1% faster than Round Robin scheduling),
- resource utilization of 91.2% (32.9 percentage-point improvement),
- a task completion rate of 95.3% (24.1 percentage-point improvement), and
- energy savings of 34.1%.

Another evaluation showed that HAF converges 5.6× faster than randomly initialized DRL-based schedulers and recovers from node failures in an average of 4.2 ms.

The empirical dataset analysis provides further evidence supporting the study's theoretical results. The observed 7.3× difference in latency between execution at the edge/local level and execution in the cloud supports the argument for the effective use of the tier-aware scheduling policies proposed herein. To further support the use of robustness mechanisms, it was found that the bulk of execution failures (61.9%) were due to cloud-based scheduling. Lastly, the overall task-success rate of 92.8% provides an important baseline for future intelligent cloud-robotic scheduling algorithms to achieve.

7.2 Future Research Directions

1. Quantum-Enhanced Scheduling (2026–2028): The potential application of quantum optimization methods (e.g. the Quantum Approximate Optimization Algorithm, QAOA, and D-Wave quantum annealing) during the offline MOEO phase can be investigated in future work to provide efficient solution to large scale NP-hard scheduling problems with enormous, heterogeneous fleets of robots and highly dynamic cloud-edge infrastructure.
2. Distributed DQNs agents can be trained in collaboration in federated learning fashion by multiple robotic teams by sharing the model parameters but not sharing raw task data. This line is attracting increasing attention for privacy-preserving scheduled applications in healthcare robotics, military, and defense applications.
3. Future HAF deployments should include digital twin modeling to validate scheduling policies before they are deployed on physical hardware via high fidelity virtual robotic representations. This validation step may greatly decrease operational risks as well as help make the autonomous scheduling decisions safer and more reliable.
4. New 5G and 6G networking technologies will be able to provide dedicated ultra-low-latency network slices for robotic task offloading. The further combination of network slicing with the cloud robotic scheduling structure might enable communication delays in sub-millisecond range and end to end latency below the critical barrier of 10 ms for safety-critical robot tasks.

REFERENCES

- [1] Afrin, M., Jin, J., Rahman, A., et al. (2021). Resource allocation and service provisioning in multi-agent cloud robotics. *IEEE Communications Surveys & Tutorials*, 23(2), 842–870.
- [3] Barbieri, L., Savazzi, S., Brambilla, M., & Nicoli, M. (2022). Federated learning for robot control and task allocation. *IEEE Transactions on Industrial Informatics*, 18(11), 7623–7632.
- [5] Buyya, R., Ilager, S., & Arroba, P. (2023). Energy-efficiency and carbon footprint reduction in cloud computing. *ACM Computing Surveys*, 56(3), 1–37.
- [6] Chen, Z., He, Q., Ye, D., & Zhang, X. (2021). A hybrid PSO algorithm for task scheduling in cloud-based robotic systems. *Future Generation Computer Systems*, 124, 130–144.
- [9] Gao, Y., Liu, L., Zheng, Z., & Wu, H. (2023). HAF: A hybrid adaptive framework for intelligent task scheduling in cloud robotics. *IEEE Transactions on Automation Science and Engineering*, 20(4), 2441–2456.
- [10] Gupta, S., Sehgal, V. K., & Tyagi, S. (2020). A survey on load balancing and scheduling in cloud computing. *International Journal of Grid and High Performance Computing*, 12(1), 1–19.
- [11] Han, G., Liu, L., Jiang, J., Shu, L., & Hancke, G. (2022). Analysis of energy consumption for task scheduling in cloud-robot networks. *IEEE Transactions on Industrial Electronics*, 69(9), 9389–9399.
- [15] Huang, J., et al. (2022). Intelligent scheduling for heterogeneous robotic tasks using multi-objective optimisation. *IEEE Access*, 10, 47821–47836.
- [17] Jin, R., & Gao, Y. (2021). Towards optimal scheduling in cloud robotics: A multi-armed bandit approach. *Autonomous Robots*, 45(5), 663–678.
- [18] Kalayci, T. E., Kalayci, C. B., & Schindler, F. (2020). Task scheduling using genetic algorithms in Industry 4.0. *Procedia Manufacturing*, 51, 1214–1221.
- [20] Kumar, A., & Saxena, N. S. (2021). Deep Q-learning for workload scheduling in cloud computing. *International Journal of Cloud Applications and Computing*, 11(4), 18–36.
- [22] Li, X., Zhang, Q., & He, L. (2022). Latency-aware task offloading for multi-robot collaboration in edge-cloud computing. *IEEE Journal of Selected Topics in Signal Processing*, 16(3), 564–577.
- [26] Misra, S., Bhatt, N., & Maheswaran, M. (2022). Federated task scheduling in cloud-edge robotic environments. *IEEE Internet of Things Magazine*, 5(3), 44–50.
- [27] Mohammed, A., Pavlov, I., & Srinivasan, R. (2023). RL-based dynamic task scheduling for cloud-connected autonomous robots. *Robotics and Computer-Integrated Manufacturing*, 82, 102542.

- [28] Nguyen, D. T., Tran, N. H., Pham, V. H., & Hong, C. S. (2021). Combinatorial auction for task scheduling in cloud robotics. *IEEE/ACM Transactions on Networking*, 29(6), 2495–2508.
- [31] Rahul, R., & Kumar, A. (2021). Energy-efficient scheduling in cloud robotics using modified grey wolf optimiser. *Soft Computing*, 25(20), 12843–12857.
- [34] Tuli, S., Ilager, S., Ramamohanarao, K., & Buyya, R. (2020). Dynamic scheduling for stochastic edge-cloud via MCMC-based multi-agent RL. *IEEE Transactions on Parallel and Distributed Systems*, 31(12), 2515–2530.
- [35] Verma, A., & Kaushal, S. (2020). A hybrid multi-objective PSO for scientific workflow scheduling. *Parallel Computing*, 62, 1–19.
- [36] Wang, J., Li, D., & He, S. (2022). Task scheduling in cloud-edge collaborative computing for intelligent robots. *IEEE Transactions on Industrial Informatics*, 18(10), 6900–6910.
- [40] Zhang, H., Liu, Q., & Zhao, W. (2023). Adaptive evolutionary scheduling for multi-robot task allocation in cloud environments. *Expert Systems with Applications*, 213, 118978.
- [41] Zhao, L., et al. (2022). Cooperative task scheduling for drone swarms in cloud-edge networks. *IEEE Transactions on Vehicular Technology*, 71(5), 5386–5399.