



Federated learning and deep reinforcement learning synergy: opportunities for multi-cloud serverless deployment

I Gusti Ngurah Wikranta Arsa^{*1}, Arief Setyanto², Andi Sunyoto², Alva Hendi Muhammad²

Institut Teknologi dan Bisnis STIKOM Bali, Indonesia¹
Universitas Amikom Yogyakarta, Indonesia²

Article Info

Keywords:

Multi-cloud, Serverless, Federated Learning, Deep Reinforcement Learning, PRISMA, Optimization

Article history:

Received: January 19, 2026

Accepted: April 27, 2026

Published: August 01, 2026

Cite:

I. G. N. W. A. Arsa, Arief Setyanto, Andi Sunyoto, and Alva Hendi Muhammad, "Federated Learning and Deep Reinforcement Learning Synergy: Opportunities for Multi-Cloud Serverless Deployment", *KINETIK*, vol. 11, no. 3, Aug. 2026.
<https://doi.org/10.22219/kinetik.v11i3.2694>

*Corresponding author.

I Gusti Ngurah Wikranta Arsa

E-mail address:

arsa@stikom-bali.ac.id

Abstract

The advancement of distributed computing has enabled the use of multi-cloud and serverless computing, which are beneficial due to their flexibility, scalability, and cost efficiency. There are, of course, pertinent challenges associated with these computing paradigms, such as resource heterogeneity, cold-start latency, vendor lock-in, and privacy. Recent trends in Federated Learning (FL) and Deep Reinforcement Learning (DRL) offer promise in tackling these challenges. FL systems enable decentralised, privacy-preserving model training across heterogeneous systems, while DRL systems enable adaptive models for real-time decision-making to optimise system resources and improve performance. This Systematic Literature Review (SLR) covers the years 2020 to early 2026 and examines the intersection of FL and DRL in multi-cloud serverless computing, following the PRISMA methodology. A primary analysis of 50 quality studies was undertaken to answer four privacy-related resource management questions. The results demonstrated FL increases privacy and scalability utilizing decentralised training. Consolidating the Federated DRL and Multi-Agent stacks increases the system by obtaining a better trade-off and optimization among latency, energy, and operational efficiency. However, a few gaps still exist, such as the absence of a more holistic framework, elusiveness in cross-system integration and collaboration, and a lack of concrete real-world applications. More work is needed to build a cohesive Federated Learning framework to improve sustainability and security in the multi-cloud, serverless systems of the future. Beyond a systematic literature review, this study provides a comparative synthesis of FL–DRL-based approaches and proposes a conceptual orchestration framework as a foundation for intelligent resource allocation in serverless multicloud environments.

1. Introduction

Distributed computing has evolved rapidly with the adoption of multicloud and serverless architectures, fundamentally changing how organizations operate by offering greater flexibility, scalability, and cost efficiency. These paradigms enable dynamic resource deployment and abstract infrastructure management, making them well suited for modern data-intensive applications. At the same time, Federated Learning (FL) has emerged as a promising distributed learning paradigm capable of achieving predictive performance comparable to centralized training while reducing communication overhead and preserving data privacy, since raw data remain local to participating nodes [1]. Consequently, FL is increasingly considered relevant for heterogeneous and privacy-sensitive distributed environments.

However, the transition toward serverless multicloud systems introduces significant orchestration challenges, including resource heterogeneity, cold-start latency, vendor lock-in, and privacy concerns [2][3]. While Deep Reinforcement Learning (DRL) has been applied to dynamic scheduling, real-time load balancing, and auto-scaling in serverless platforms [4], and FL has addressed decentralized and privacy-preserving learning, existing research remains fragmented. Most studies focus on isolated optimization objectives, rely heavily on simulation-based evaluations, or lack standardized orchestration frameworks suitable for heterogeneous multicloud environments. Furthermore, unresolved issues such as statistical and system heterogeneity, data imbalance, and cost-performance trade-offs continue to limit the practical applicability of current solutions [5][6][7][8].

The integration of FL and DRL therefore represents a promising approach for intelligent resource orchestration in serverless multicloud systems, where privacy preservation, scalability, and efficient resource allocation are critical. FL protects sensitive data within local environments while minimizing communication overhead and supporting regulatory compliance such as GDPR [9], whereas DRL—particularly multi-agent methods—effectively optimizes performance metrics including latency, throughput, and energy consumption under dynamic conditions [10][11].

Although prior work has explored hybrid federated architectures and cryptographic enhancements to improve privacy and efficiency [12], a unified orchestration perspective is still lacking. Accordingly, this paper goes beyond a conventional systematic literature review by providing a comparative synthesis of FL–DRL-based approaches and formulating a conceptual orchestration framework as a foundation for intelligent, scalable, and interoperable serverless multicloud systems.

2. Research Method

The trend of Federated Learning (FL), Deep Reinforcement Learning (DRL), multicloud, and serverless computing based on research outputs over the past five years (2020 to 2025) in Scopus can be observed as follows: Figure 1

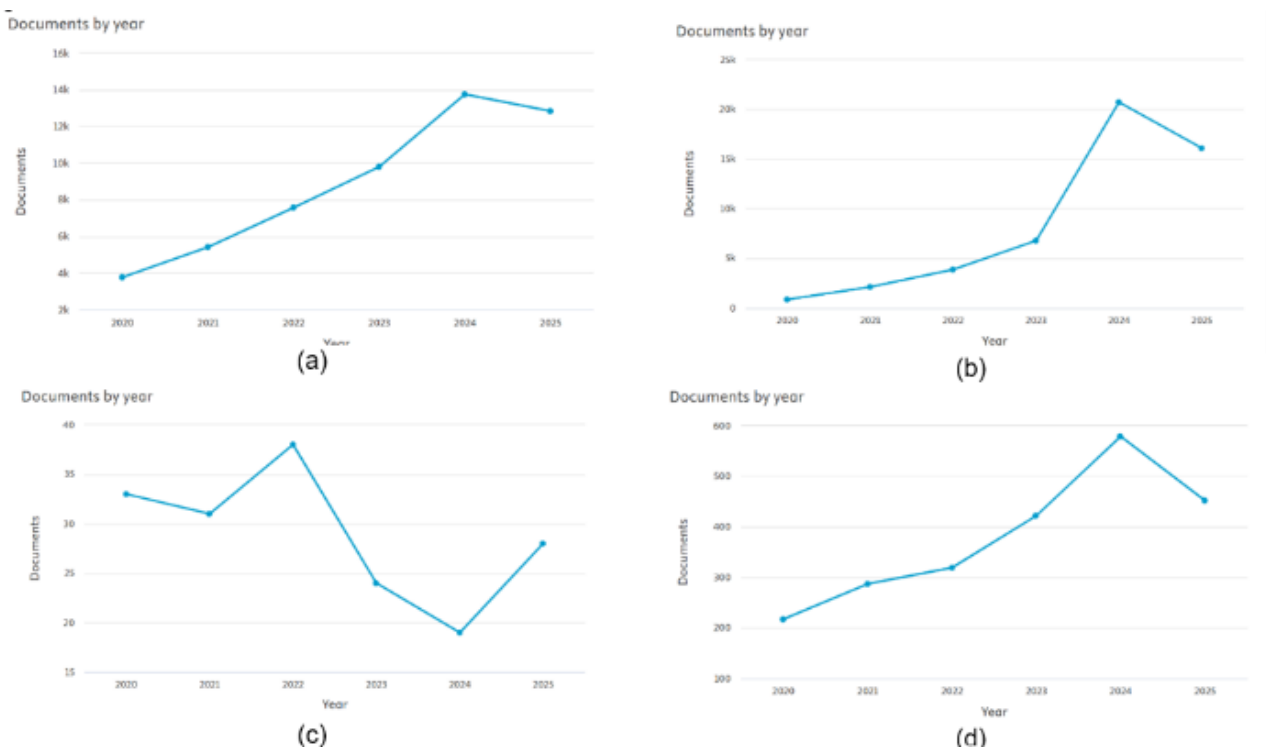


Figure 1. Trends in Research Topics

According to the keywords "Deep Reinforcement Learning", the highest number of publications, 13,772, was indexed in 2024. The keyword "Federated Learning" yielded 50,422 publications, the highest number for 2024, with 20,720. In contrast, the area of research on "Multicloud" yielded only 173 publications, with the highest number in 2022 (38). The highest number of publications focusing on "Serverless" environments was 2,276, with 2024 the highest year (579 publications). The research output for these topics shows that there are currently active trends, particularly the combination of Federated Learning with Deep Reinforcement Learning in Multicloud Serverless environments. The publication trends indicate a rapidly growing research interest in Federated Learning and Deep Reinforcement Learning, while studies explicitly addressing multicloud serverless environments remain relatively limited. This imbalance highlights the relevance and novelty of investigating FL–DRL integration in multicloud serverless contexts.

The results of these trends explain the continued relevance within these research fields. The purpose of this research is a Systematic Literature Review (SLR) Methodology and more specifically, the PRISMA framework. The procedure has four specific steps: Identification, Screening, Eligibility and Inclusion. It is depicted in Figure 2. Searching for literature: The first step was to search several well-known databases such as ScienceDirect, Springer, Scopus, Wiley and IEEE Explore to investigate the latest publications on this subject. We used three keyword strategies to collect studies that are relevant to Federated Learning (FL), Deep Reinforcement Learning (DRL), and the combination of them in multicloud serverless environments:

"Federated Learning" AND "Serverless" AND "Multi-cloud" (yielding 55 records).

"Deep Reinforcement Learning" AND "Resource Management" AND "Multi-cloud" (yielding 201 records).

"FL AND DRL AND Serverless AND Multi-cloud" (yielding 46 records).

There were a total of 239 records after the initial data collection. Duplicate records needed to be removed during the screening process. Additionally, other irrelevant studies to the main topics of focus also had to be removed. Studies

that were outside the desired publication date of 2020 to the present year (2026), and that were also unverified peer reviews, also had to be removed, which excluded 73 records for the remaining studies, and they were then moved to the next stage in the process. During the eligibility stage, the remaining studies' abstracts and titles were examined. . Flow of the PRISMA can be seen in Figure 2.

Out of the previous 239 records, 84 of the studies were deemed to meet the eligibility criteria to be then moved to the next step, while the remaining 155 were irrelevant, in that they did not have any comparisons to focus on topics like Federated Learning (FL), Distributed Reinforcement Learning (DRL), and Multicloud Serverless Resource Management. An inclusion step, which is a full review of the remaining studies, analyses all the potentially eligible studies in detail, and after some unrelated studies were removed, a total of 50 studies were included in the final data set. With these final studies, a foundational basis was provided to answer the primary objectives of the research, such as the identification of methods, or to provide an overview of how FL and DRL have all been in studying the use of a serverless multicloud environment, to further provide detail on any underlying issues.

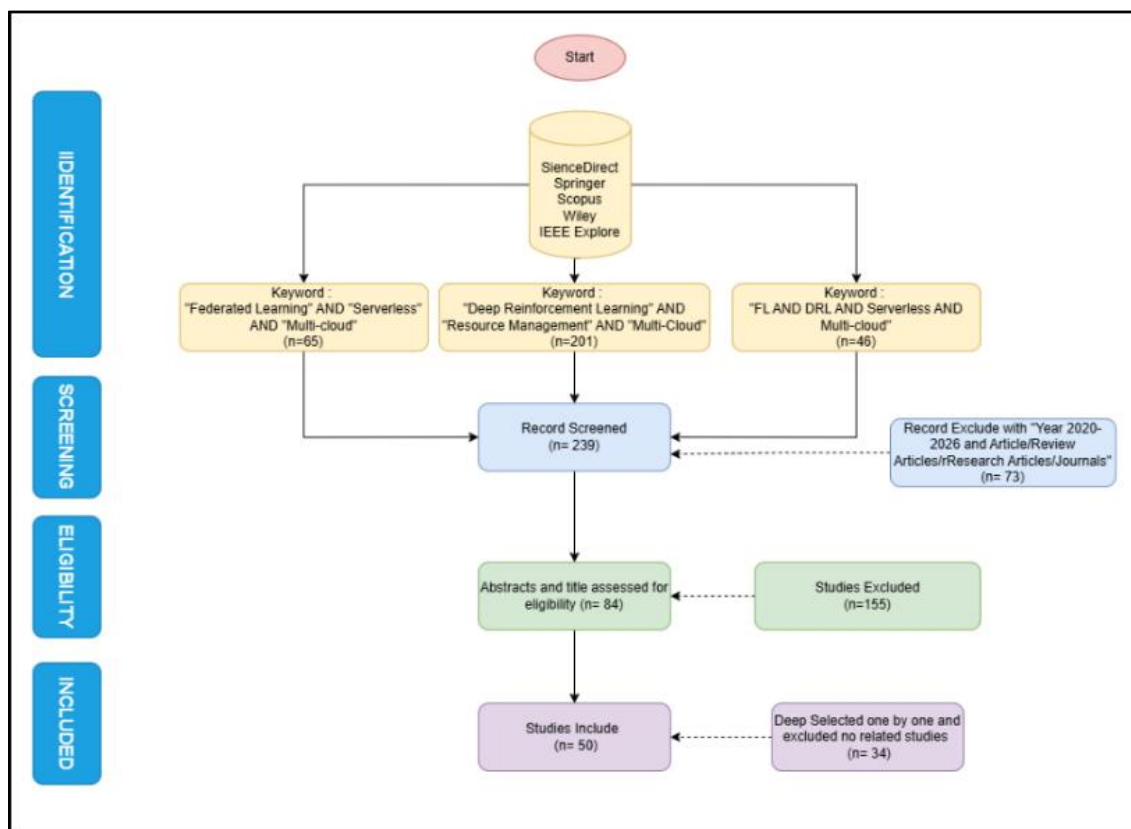


Figure 2. PRISMA-based Flow Diagram Illustrating the SLR Process

This study, after systematically recognizing, screening, and including pertinent studies released between 2020 and early 2026, involves the curation of datasets consisting of quality articles, journals, research papers, and review papers, all of which have been peer-reviewed. These documents form the basis of the study of the combination of Federated Learning and Deep Reinforcement Learning in multi-cloud serverless environments. This study is guided by a single overarching research question: How can the integration of Federated Learning (FL) and Deep Reinforcement Learning (DRL) be effectively designed and optimized to address privacy preservation, scalability, dynamic resource allocation, and performance optimization in multicloud serverless environments while improving latency, energy efficiency, and cost–performance trade-offs, and what research gaps remain for developing standardized and interoperable architectures?

In addition to study selection, this review applies a systematic comparative synthesis approach, extracting performance metrics (latency, energy, cost, SLA), architectural patterns, and learning strategies from selected papers. The synthesis results are then abstracted into high-level architectural components and orchestration functions, which serve as the basis for formulating a conceptual framework.

3. Results and Discussion

This section presents a comparative synthesis of FL and DRL approaches to identify recurring performance patterns, architectural characteristics, and orchestration requirements across the reviewed studies. In Table 1 presents the mapping of 50 papers curated according to the research topic. There are 15 papers referenced for the FL topic, 16 for the DRL topic, 10 for multi-cloud, and 9 for serverless learning. Each of these topics will later be used to address the research questions identified earlier.

Table 1. Topic from Resource Paper

Topic	Resource Paper
Federated Learning	[13][14][15][16][17][18][19][20][21][22][23][24][25][26][27]
Deep Reinforcement Learning	[28][29][30][31][32][33][34][35][36][37][38][39][40][41][42][43]
Multi-cloud	[44][45][46][47][8][48][49][50][51][52]
Serverless	[53][54][55][56][57][58][59][60][61]
Total	50

3.1 Federated Learning Privacy and Scalability

The following analysis relies on the literature previously mapped in RQ1. FL is a promising approach to overcoming the challenges of data privacy and scalability in distributed systems, particularly those utilising cloud computing resources. In the literature, FL is a paradigm for data privacy and scalability challenges in distributed systems, cloud computing, and related areas [19][25]. A centralised controller typically computes a global model. However, there are other ways to integrate FL into a blockchain system for the distributed computation of global model updates across network slice tenants. These tenants collaborate to minimise a global loss function by optimising the weighted average of their respective local loss functions [17]. While the literature is silent on the term “multicloud serverless,” the application of FL in cloud computing and IoT offers a glimpse into its architecture and the opportunity to develop new architectures. FL is particularly relevant to modern ecosystems that include IoT devices and edge computing, as the paradigm facilitates collaborative model training while keeping the training data localised [13]–[16][18][19][25]. Each client’s data resides locally, as model updates are sent, thereby mitigating the risk of information leakage while also supporting compliance with privacy regulations [13]–[16][18][19].

Paillier Homomorphic Encryption and Blinding Method, and other privacy measures emergent in the architecture, can be complemented with other privacy-enhancing core operationalization in FL, and other detective privacy in FL prevention of Inference Attacks [19]. FL also postulates and augments the ability to incur less communication costs; thus, the cost of communication FL provides a potential cost of communication. Telecommunications, IoT, and pharmaceuticals [21]. FL provides advantages and also minimizes the disadvantages of centralized architecture, such as single-point failures and data-at-scale leaks. FL also consolidates data from several devices to sustain the quality of the global model [14][16][22][23][27].

Regarding the architecture, FL scales by distributing computational load across devices attached to the network and the main central processing unit [14][25][26]. By this method, there is less communication overhead—FL transfers model updates rather than sending larger raw-data files. FL efficiently distributes model updates, which benefits IoT networks [14][24]. FL is integrated with edge, fog, and cloud architectures and can support large-scale privacy-preserving training [14][19]. The FDSAC, along with other adaptive techniques, dynamically repositioned FL nodes and other techniques to the edge, helping accelerate convergence and better utilize disparate, highly heterogeneous resources [19][20][26].

To summarise, Federated Learning (FL) addresses privacy concerns by retaining raw data within local environments and exchanging only model updates, optionally enhanced with cryptographic techniques, which distinguishes it from centralized learning approaches that require extensive data movement and global coordination. This privacy-by-design paradigm not only reduces the risk of data leakage but also improves scalability by minimizing communication bottlenecks and enabling parallel training across heterogeneous devices and cloud resources. However, compared to centralized and fully coordinated orchestration models, FL still faces challenges related to data heterogeneity and training efficiency, particularly in dynamic multicloud environments. Existing studies further indicate that while blockchain, encryption, and differential privacy can strengthen security during model exchange, their effectiveness depends heavily on coordinated aggregation, model selection, and deployment mechanisms. In multicloud settings, the process of global model aggregation followed by selective dissemination to local nodes highlights the critical need for an orchestration framework that can manage collaboration, ensure consistency, and scale learning across distributed environments, as illustrated in Figure 3.

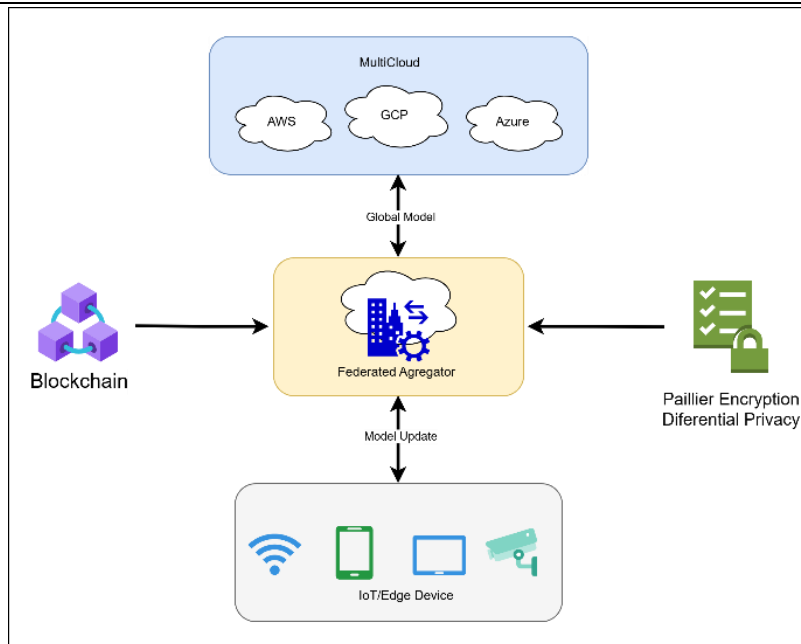


Figure 3. The Relationship Between Privacy, Local Models, and Global Models in Federated Learning

3.2 Deep Reinforcement Learning: dynamic resource allocation and performance

Deep Reinforcement Learning (DRL) is the cutting edge of dynamic resource allocation and performance profiling in distributed systems. Within the context of the second research question (RQ2), the confluence of how the technologies within DRL can make 'rational' and adaptable sets of decisions within environments characterised by uncertainty by employing deep learning and reinforcement learning.

Dynamic resource allocation and performance optimisation in serverless multicloud environments are among the many challenges in the current computing age. The integration of deep reinforcement learning (DRL) and deep neural networks has shown promise for adaptive decision-making in changing environments [30][33]. Unlike previous optimisation techniques, DRL is the first of its kind to determine solutions by adaptive interaction with the environment. This feature makes it an excellent fit for heterogeneous structures. Some of the many positive applications of DRL policy are in resource allocation, which is energy and cost-effective, workload allocation, energy and cost-effective dynamic resource allocation, and efficient resource scheduling [30][43]. In data centres, Adaptive Learning with an Actor-Critic Model has been employed. Q-learning techniques and models for data centres are discussed, along with their use to adapt to desired auto-scheduling and Heuristic Scheduling models and to mitigate the cold-start problem. TF-DDRL has been proposed for scheduling optimisation in the edge and cloud integration layers of the IoT. The solutions that integrate multiple agents for dynamic task allocation in Fog Networks, such as DMARL, are other recent examples [30][42]. An offloading system that incorporates a UAV to support a digital twin-based industrial IoT also incorporates DRL [38].

Resource management based on Deep Reinforcement Learning (DRL) aims to achieve three goals: enhancing system efficiency and functionality, reducing operational costs, and maximising the operational system's (the energy system's) efficiency. Performance-optimal energy efficiency and energy-optimal task performance (the workload) are achieved through DRL-based federated workload scheduling. Cross-domain resource allocation and federated approaches enable distributed learning and supported ecosystems [28][29]. Such approaches promote scalability and the sustainability of the distributed ecosystem. However, in the multicloud environments, systems exhibit DRL challenges in computational complexity and convergence speed as well as in (data) heterogeneity. Recent studies have integrated DRL and federated learning (FL) with privacy-preserving mechanisms to address those challenges. In addition, more advanced (hybrid) approaches for improving operational effectiveness and system reliability have integrated DRL with other machine learning types (e.g., heuristic) and predictive modelling [32][39]. Innovative federated Deep Double Q-Network (FD3QN) is a good example. It vertically integrates federated learning into D3QN for cross-domain resource scheduling, achieving lower latency than conventional DRL and 15% higher energy efficiency [29]. FUD3QN is another model that optimises resource allocation in hybrid multi-server MEC architectures [31][39].

Additionally, Adaptive Federated DRL (AFDRL) is an optimisation technique for predictive edge content caching that mitigates non-IID data and resource issues through flexible client assignments and a DDQN-based architecture [40]. Moreover, dynamic monitoring systems achieve 99.66% accuracy with low false positive rates [37]. Examples of DRL federated learning are growing rapidly, particularly by integrating self-supervised learning (FedSSL) and lightweight

models, such as those optimised with LightGBM, to improve performance in resource-constrained environments [41]. Regarding these approaches, providing DRL and adaptive federated learning enables intelligent optimisation for seamless interoperability and advanced cloud-edge collaboration across multiple systems [31][34][35].

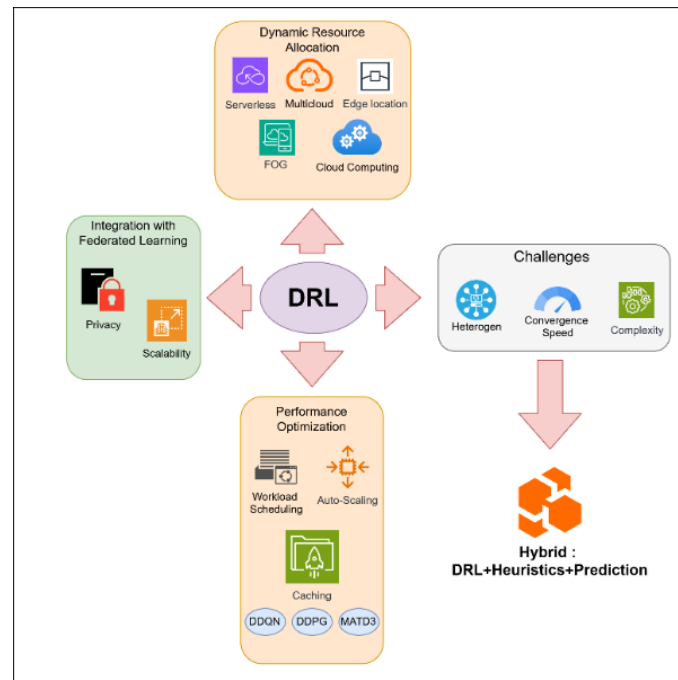


Figure 4. Conceptual Framework for DRL Implementation

Figure 4 illustrates the role of Deep Reinforcement Learning (DRL) as a core mechanism for adaptive and responsive resource management across heterogeneous computing environments, including serverless, multicloud, edge, fog, and cloud systems. Compared to static or heuristic-based resource allocation approaches, DRL enables dynamic resource placement, autoscaling, and workload scheduling that improve elasticity and utilization under fluctuating workloads. However, standalone DRL solutions face limitations related to convergence speed, computational complexity, and sensitivity to heterogeneous infrastructures. The integration of Federated Learning (FL) addresses these limitations by enabling privacy-preserving and scalable distributed training without transferring raw data, thereby supporting regulatory compliance and cross-domain collaboration. Furthermore, advanced DRL algorithms such as DDQN, DDPG, and MATD3 enhance decision quality for scheduling, caching, and scaling, while recent studies suggest that combining DRL with lightweight models, attention mechanisms, and heuristic or predictive approaches further improves adaptability in non-IID and resource-constrained environments. These findings indicate that FL–DRL integration is not merely an algorithmic enhancement, but a necessary foundation for intelligent orchestration, where coordinated learning, decision-making, and deployment strategies are required to achieve efficient, scalable, and privacy-aware resource allocation in serverless multicloud systems.

3.3 Integrating Federated Learning and Deep Reinforcement Learning

Combining FL with DRL for RQ3 (called Federated Reinforcement Learning) is a new direction that seeks to leverage the best of both worlds to solve problems in distributed systems, with a focus on IoT and edge computing. The joint method is intended to maximize several performance measures while efficiently guaranteeing data privacy [24].

Table 2 summarizes conceptual approaches that integrate Federated Learning (FL) and Deep Reinforcement Learning (DRL) to address challenges in distributed and heterogeneous environments such as Industrial IoT (IIoT) and edge computing. These approaches demonstrate that FL–DRL integration improves cost efficiency, learning performance, and privacy preservation, while enabling multi-objective optimization across task allocation, energy consumption, response latency, resource migration, and operational cost. In collaborative edge networks, federated deep Q-learning–based caching further mitigates uncertainty in user demand and content popularity by minimizing long-term system costs. Overall, these studies highlight the effectiveness of FL–DRL integration in enhancing scalability, efficiency, and adaptability in complex distributed systems.

Table 2. Concept FL and DRL Integration

Author	Conceptual Integration Approach	Description
S. Demil and M. R. Abdmeziem [14]	DRL-assisted FL	DRL supports FL in IIoT environments, enabling more efficient learning, cost reduction, and privacy preservation.
Sekhar, et. al [30]	Federated Multi-Objective DRL	This approach simultaneously optimizes task scheduling, energy efficiency, response latency, and resource migration costs.
Zhou, et.al [36]	Federated Deep Q-Learning Caching	A federated deep Q-learning caching scheme minimizes long-term system costs in collaborative edge networks under demand and content popularity uncertainties.
Xiao, et.al and Zhang , et.al [29], [39]	Synergistic Integration for Performance using Algorithms like FD3QN	Combining FL and DRL with innovative algorithms such as MATD3 and FD3QN significantly enhances performance and supports safe UAV (Unmanned aerial vehicles) trajectory planning.
G. Nagabushnam and K. Hoon [42]	Distributed Multi-agent DRL Frameworks	A distributed multi-agent DRL approach optimizes deadlines, energy consumption, makespan, and scheduling latency in heterogeneous fog networks.
A. S. M. S. Sagar, A. Haider, and H. S. Kim , Sekhar, et. al [24], [30]	FRL for Resource Allocation and Task Scheduling	Federated reinforcement learning intelligently optimizes resource allocation and task scheduling within federated and edge-cloud frameworks.

Another advancement is a well-developed integration strategy that leverages algorithm developments such as MATD3 and FD3QN, which have been shown to improve performance and support safe inter-UAV planning for offloading decisions. A distributed multi-agent DRL framework can be developed to optimize deadlines, energy consumption, makespan, and scheduling latency in heterogeneous fog networks. Furthermore, Federated Reinforcement Learning (FRL) emerge as a robust paradigm for intelligently optimizing resource allocation and task scheduling in federated and edge-cloud architectures. According to the paper, this integration approach emphasizes combining FL and DRL to achieve multi-objective optimization while addressing the challenges of privacy, heterogeneity, and cost-performance trade-offs in modern computing environments.

Table 3. The Impact of FL and DRL Integration

Author	Proposed Method	Impact
S. Demil and M. R. Abdmeziem [14]	MADRL	Achieves higher global accuracy (94.81%) compared to a random offloading approach (93.51%) in the presence of malicious nodes.
Shinu M. Rajagopa, et.al. [19]	FedSDM	Outperforms Fog and Cloud in terms of energy consumption, network usage, cost, execution time, and latency (i.e.) 0.3%, 2%, 15%, 11%, and 3% when compared with Fog and 1.6%, 31%, 41%, 24% and 85% against Cloud respectively
Xiao, et.al [29]	FD3QN	Exhibits a 25% decrease in average service latency and a 15% improvement in energy efficiency compared to traditional DRL approaches and maintains a high task completion rate of over 98% under dynamic network conditions.
M. C. Sekhar, et.al [30]	FedTaskRL	Achieved 28 kWh Use of Energy, 145 ms Average Response Time, 97.5% of SLA fulfillment, and a lower migration cost of \$4.8.
Zhou, et.al [36]	FD3PG	Compared with FADE, MADRL, and DDPG, FD3PG achieves a significant decrease in average delivery delay, approximately 10%, 11%, and 35% on the Synthetic dataset and 12%, 14%, and 48% on the MovieLens Latest Small dataset.

The impact of Federated Learning (FL) and Deep Reinforcement Learning (DRL) integration on latency, energy efficiency, and cost–performance trade-offs, as summarized in Table 3, demonstrates clear advantages over standalone DRL approaches. Compared to conventional DRL methods such as DQN and DDPG, federated DRL models—including FD3QN highlighting the effectiveness of collaborative and decentralized decision-making in dynamic environments [19], [29], [30]. Beyond latency optimization, FL–DRL integration also improves energy efficiency by enabling energy-aware

resource management, where FL reduces local device battery consumption and FD3QN further lowers overall energy usage by up to 15% [19], [29], [30]. In addition, federated deep Q-learning–based caching and multi-objective DRL frameworks minimize long-term system costs by jointly optimizing resource migration, workload placement, and operational efficiency, thereby improving cost–performance trade-offs in heterogeneous settings [14], [30], [36]. Collectively, these findings indicate that FL–DRL integration makes a significant and non-trivial contribution to intelligent distributed systems by simultaneously optimizing latency, energy, and cost while addressing privacy and heterogeneity challenges, which in turn necessitates a robust orchestration strategy—particularly in serverless multicloud environments characterized by diverse service providers and operational constraints

3.4 Standardized frameworks for FL-DRL integration in a multicloud serverless environment

These findings indicate that FL–DRL integration is not a trivial extension of existing resource management techniques but represents a necessary architectural shift toward intelligent orchestration in serverless multicloud systems. The development of standardized frameworks for integrating FL and DRL in multicloud serverless environments presents several research gaps and opportunities for future work, current challenges, and research gaps aligning with research question RQ4.

Table 4 outlines the challenges and opportunities in researching the integration of FL and DRL in multicloud serverless environments. Most RL-based approaches to serverless scheduling still rely on simulations that fail to represent the complexity of real-world distributed systems, necessitating the need for more realistic simulators. In addition, current Function-as-a-Service (FaaS) frameworks are designed for cloud environments. They are therefore ill-equipped to handle the edge-cloud continuum, which is highly diverse, characterized by significant network latency, and resource-constrained. Other gaps include a lack of protocol standardization, the absence of a unified framework for multi-objective optimization, and minimal consideration of multi-user and horizontal offloading.

Table 4. Challenges and Opportunities Framework

Author	Challenge	Opportunity
[59], [52]	Real-world complexity not captured in simulation environments.	Develop more realistic simulators or conduct experiments in real multicloud settings.
[55]	FaaS frameworks designed mainly for cloud, not edge-cloud continuum.	Design new architectures to handle heterogeneity, high latency, and limited edge resources.
[47]	Lack of standardization for multi-cloud collaboration.	Establish standardized protocols and frameworks for seamless FL-DRL integration across providers.
[49]	Absence of unified frameworks for multi-objective optimization.	Create frameworks that optimize latency, energy consumption, and host utilization simultaneously.
[58]	Limited consideration of multi-user and horizontal offloading.	Develop scheduling algorithms that support multi-user classes and horizontal offloading.
[49]	Neglect of trust, security, and SLA compliance.	Integrate trust, security, and SLA compliance mechanisms into resource allocation decisions.
[47], [53]	Data privacy and model efficiency not fully optimized.	Enhance FL-DRL integration to improve privacy and computational efficiency.
[50]	Performance, scalability, and cost-effectiveness remain low.	Apply AI-driven approaches (e.g., DRL) for predictive container orchestration to reduce costs.
[52]	High variance in heterogeneous environments.	Develop adaptive strategies to minimize energy variance and SLA violations.
[58]	DRL solutions require computationally intensive training per node.	Design lightweight and efficient DRL training methods for distributed environments.
[50]	Lack of optimal container placement and allocation strategies.	Identify algorithms for optimal initial container allocation in dynamic environments.
[58]	Centralized schedulers in FaaS frameworks.	Develop decentralized schedulers for multicloud-to-edge continuum scenarios.
[8]	Interoperability issues between legacy interfaces and cloud apps.	Create compatible interfaces for seamless integration of on-premises and cloud applications.

Security, trustworthiness, and SLA compliance are also often overlooked, while FL-DRL integration still needs improvement to ensure data privacy and model efficiency. Other challenges include improving performance, scalability, and cost-effectiveness; managing performance variations across heterogeneous environments; and the high computational intensity of DRL training. The factors of container placement strategy, limitations of centralised schedulers, and the interoperability between current interfaces and cloud applications necessitate the creation of novel adaptive, distributed, and resource-efficient structures and algorithms. Overall, this gap presents research opportunities to develop innovative solutions that optimize performance, privacy, and efficiency in a dynamic multicloud ecosystem.

Table 5. Category vs Reserch Focus

Primary Category	Research Focus Area	Author
Frameworks & Architecture	- Cloud-Edge Collaborative DRL Frameworks	[45]
	- Multi-agent Architectures	[53], [54]
Optimization & Efficiency	- DRL for Cost and Latency Optimization	[60]
	- Optimization of Sustainability	[56]
	- Robustness and Efficiency	[57]
	- Performance Metrics	[51]
	- Function Chains and Workflow Scheduling	[51], [59]
Scheduling & Deployment	- Multi-Agent DRL (Edge or Cloud)	[61]
	- Decentralized and Privacy-Preserving Orchestration	[50]
Security & Privacy	- Security and Data Integrity	[50]
	- Security Requirements	[39], [58]
	- Practical Training Environments	[58]
Practicality & Evaluation	- SLA-based Reasoning Models	[39], [54]
	- System Parameters (VM clusters and dynamic serverless workloads)	[46]
	- Federated, Hybrid, and Sky Serverless Computing	[8], [56]
Emerging Paradigms	- the serverless edge-cloud continuum	[8], [58]

Table 5 clearly shows several research opportunities for integrating FL and DRL into multicloud serverless architectures, enabling the development of collaborative cloud-edge architectures that prioritize scalability, lightweightness, and interoperability. The primary challenge is optimizing function placement on serverless platforms, which presents opportunities for adaptive machine learning. Furthermore, scheduling function placement within a single workflow presents another challenge. Another challenge is designing practical training on open source serverless platforms to evaluate their performance and scalability. Moreover, subsequent research endeavours could focus on developing efficiency measures, analytical frameworks, and optimisation strategies for specific test parameters, incorporating sustainability measures derived from the efficient allocation of energy resources, and cross-implementing sustainable automated scaling strategies. Another priority is improving robustness and efficiency by leveraging DRL methods, such as MADQL, to reduce costs, latency, and energy consumption while maintaining throughput and scalability.

The findings obtained for research directions can include serverless computing management with federated and hybrid models, maintaining privacy in decentralized orchestration, adaptive deployment, and environmentally friendly. DRL can continue to be used as a method for cost and latency optimization, including in the context of multi-agent DRL as a determinant of optimal locations for running functions between the edge and the cloud. Other research can also focus on data security and integrity, operational security rules, and the development of modular multi-agent architectures for resource optimization. The need for automated SLA learning models, evaluation of the cost and time efficiency of process placement in multicloud environments, and the development of data protection mechanisms are other opportunities for further research to create adaptive, efficient, and sustainable solutions in complex multicloud environments.

The three-layer architecture in Figure 5 supports Deep Reinforcement Learning (DRL) and Federated Learning (FL)-based Resource Management and Orchestration in a distributed computing environment. The first layer is the Cloud Layer, designated as the training and orchestration layer. This layer addresses scalability, cost, sustainability, privacy and security of training. It trains DRL and FL models, which are then pushed down to lower layers. The Cloud Layer is responsible for global optimization of resource management policies and strategies. Figure 5 is not derived from a single study but synthesized from recurring architectural patterns identified across the reviewed literature.

The second of three tiers is the orchestration layer. This is the decision-making and processing engine consisting of three components: the SLA Reasoning Module, which ensures compliance with Service Level Agreements (SLAs); the Multi-Agent Coordinator, which oversees the complex orchestration of flows between agents; and the DRL/FL orchestrator, which provides adaptive decisions. The final tier is the Edge/Fog layer, which is the implementation and deployment center. Local operations are performed under the privacy, resource, and health monitoring of local DRL agents to enable rapid edge decision-making. This architectural structure enables bilateral communication from the cloud to the edge, facilitating secure, seamless, and adaptive orchestration across multi-cloud and edge computing environments.

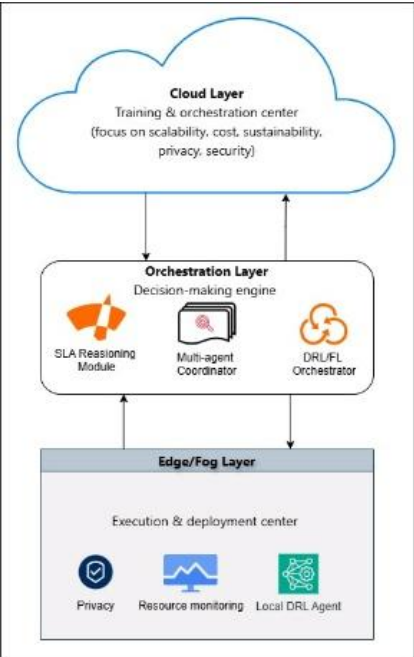


Figure 5. The Three-layer Architecture DRL and FL

3.5 Limitations and Open Challenges

The technical challenge lies in integrating both Machine Learning methods into practice in a serverless multicloud environment. One of the most significant technical limitations is infrastructure heterogeneity and poor interoperability among multicloud service providers. Studies have shown that there is no standard collaboration protocol for seamless integration between FL, DRL, and cloud platforms; model performance is hampered by varying latency, API formats, and policies across cloud providers.

Furthermore, the biggest drawback is that existing serverless architectures are too centralized and do not support the full functionality of edge-cloud-multicloud scenarios, making it difficult to adaptively move functions between clouds. Finally, these systems have proven difficult to scale and operate efficiently in highly dynamic multicloud environments. From an algorithmic perspective, the main technical limitations lie in the high computational complexity and slow convergence of DRL models, especially in highly dynamic serverless multicloud environments. For example, DRL requires intensive training per node, while edge resources are limited, and distributed training increases the risk of resource bottlenecks.

Table 6. Research Gap	
Author	Research Gap
Xiao, et al., 2025 [29]	Suboptimal computing resources, rapidly changing service requirements, deployment issues in hybrid multi-server MEC, shortcomings of classical optimization, need for collaborative privacy-preserving solutions, and performance issues with DRL approaches. Enhancing Adaptability and Autonomy, integrating with other technologies such as privacy-preserving techniques, and providing the decision-making process of the DRL agent and the aggregation mechanism (non-IID). Centralized learning increases communication overhead, and prior joint cache-recommendation studies overlook user preference heterogeneity, limiting accuracy and efficiency in real-world scenarios.
Singh & Adhikari, 2025[13]	
Zhou et al., 2024 [36]	

Z. Zhang, et al., 2024 and H. Zhang, et al., 2024 [47], [53]	While FL-DRL integration offers advanced solutions for tasks such as secure task offloading and resource allocation by leveraging FL for data privacy and DRL for optimization, further research is needed to refine these integrated approaches.
A. John, et al., 2025, [50]	Investigating federated learning approaches is crucial to enabling decentralized, privacy-preserving model training across multiple cloud environments. This would improve orchestration without compromising sensitive data.
G. R. Russo, et al., 2024 [55]	Previous Function-as-a-Service (FaaS) frameworks were developed primarily for use in cloud environments. New frameworks are needed to accommodate the edge-cloud continuums. Increased heterogeneity, greater network latency, and diminished computing resources. Although Sledge and tinyFaaS are designed for edge environments, they target single-node deployment and are unable to harness cloud resources optimally.
G. R. Russo, et al., 2024 [58]	Popular open-source FaaS frameworks often rely on centralized schedulers or gateway components, which can be a limitation in decentralized multicloud-to-edge continuum settings.
C. Jiang and L. Su, 2025 [48]	Techniques like F3WDS (FedAvg) have shown promising results in reducing job completion time, increasing resource utilization, improving energy efficiency, and enhancing quality of service compared to DRL-based baselines, suggesting avenues for further optimization in FL-DRL integrated systems.

Federated Learning still faces non-IID and imbalanced data, as well as system heterogeneity, which make global models in multicloud environments very difficult to stabilise and yield inconsistent performance. The challenge is even more severe when FL is paired with DRL, because both are equally sensitive to cold data, client environments, and client device variations. The spesifik Research Gap can see in Table 6.

3.6 Trade-offs in FL–DRL Integration for Multicloud Serverless Environments

The central constraints would be centred around the improvement in the use of DRL for efficiency, for the Latency for resource-elastic scheduling dynamically, and how the overhead would be for the algorithms. Additionally, the training for DRL is on the higher end of the computation load, and this can be especially true for multi-cloud or edge systems, when the processing power for the nodes is on the lower end. Federated learning is also training models, which means the local data must be sent, which causes packets to be sent, and the systems can become more heterogeneous. Therefore, the attempts to use FL-DRL are an improvement in systems, but the cost is high in terms of resources, and the complexity is very high in terms of training models.

The use of FL in multi-cloud will be the use of the edge of the microservices, which will be serverless. There is also the use of privacy, which can be high, but the efficiency of the systems cannot be high; there must be a balance between the two. Serverless can also not become Latency. There is a trade-off between privacy and operational efficiency, as single systems are more interconnected, the more efficiency must be given to the use of microservices and privacy. In addition, there can be more confidentiality, but there must be, in combination, more computing power, and the Latency on the systems can become greater. There can be more Latency as the overhead increases. The greater the number of micro services, the less efficient it becomes because of the number of controls on the systems to be serverless and on the microservices, and the more the Latency increases of the control systems. It becomes more efficient.

Integrating FL–DRL in heterogeneous multicloud environments requires balancing framework consistency and operational flexibility. While a consistent framework supports system integration, reproducibility, and cross-provider deployment, excessive rigidity may limit adaptability to dynamic workloads, multi-tenancy, and Service Level Agreements (SLAs). Conversely, highly flexible but unstructured approaches often result in fragmented designs that are difficult to implement and scale in practice. Therefore, the evolving nature of FL–DRL demands an interoperable orchestration framework that ensures integration while retaining sufficient flexibility to accommodate the complexity and variability of multicloud environments.

3.7 Future Work and Opportunities

Designing standardized frameworks that foster interoperability between the integration of Federated Learning and Deep Reinforcement Learning (DRL) in diverse multi-cloud and edge environments should be the primary focus of future research. As of now, there is no reference architecture to close the gaps in protocols, Application Programming Interfaces (APIs), security frameworks, and deployment policies from various cloud providers. Therefore, the focus of future research should be on developing distributed orchestration frameworks that are responsive to resource variability and dynamic workloads. These include developing a decentralized scheduling framework and creating standardized multicloud serverless workflows.

Additionally, research is needed to develop and test various approaches for integrating Federated Learning (FL), Deep Reinforcement Learning (DRL), and Reinforcement Learning (FRL) that decrease computational burden, improve training efficiency, and simplify complexity. The potential development of lightweight architectures for Deep Reinforcement Learning (DRL), techniques for gradient compression, adaptive aggregation, and partial model training methods to address the problems associated with Non-IID data, data imbalance, and the absence of sufficient synchronization in real time, are all promising. Furthermore, the FRL models developed for multi-cloud and edge environments should be the ones that facilitate real-time decision making and are robust to shifts in the environment to ensure reliable real-time decision making across edge and multi-cloud environments. Hybrid approaches, particularly the integration of DRL and other heuristics, such as optimization, meta-learning, and transfer learning, should be developed to improve generalization with little resource drain.

To validate the efficacy of FL-DRL systems at large-scale real-world networks and operational conditions, practical, large-scale, and real-world experiments must be performed in the future. Simulation-based research does not adequately depict the intricacies of multicloud environments, including SLA variability, latency changes, diverse serverless functions, and cross-provider autoscaling. To fill this void, researchers may construct multicloud-edge testbeds or combine their approaches with open-source serverless platforms to facilitate empirical assessments. Future research should also consider FL-DRL performance in real-world cloud ecosystems and explore the operational costs, sustainability, energy consumption, trust, and security to determine the potential advantages of integrating FL-DRL systems into production-grade systems.

4. Conclusion

This study examines and synthesizes recent research on Federated Learning (FL) and Deep Reinforcement Learning (DRL) in serverless multicloud environments to better understand their role in addressing privacy, scalability, and intelligent resource management challenges. This work provides a comparative synthesis of performance results, architectural patterns, and integration strategies across various studies.

Overall, the synthesized findings demonstrate that FL effectively addresses privacy and scalability challenges, DRL enables adaptive and real-time resource management, and their integration provides compounded benefits across latency, energy efficiency, and cost optimization in serverless multicloud environments.

In summary, this study provides a comparative synthesis and conceptual foundation for FL–DRL integration in serverless multicloud environments. The findings clearly demonstrate the need for intelligent resource orchestration models that coordinate federated learning, reinforcement learning-based decision-making, and deployment across heterogeneous infrastructures. Therefore, building on the identified gaps, future research should prioritize the design and validation of decentralized orchestration mechanisms that support adaptive scheduling, interoperability, and scalability in serverless multicloud systems.

Acknowledgement

Thank you to Institut Teknologi dan Bisnis STIKOM Bali which has become a place for researchers to develop this journal research. Hopefully, this research can make a major contribution to the advancement of technology especially in applying machine learning to multi-cloud serverless environments.

References

- [1] R. Vayyala, "Serverless Data Management Architectures for Multi Cloud Environments," *Proc. 7th Int. Conf. Intell. Sustain. Syst. ICISS 2025*, pp. 593–598, 2025. <https://doi.org/10.1109/ICISS63372.2025.11076490>
- [2] Y. Chen, B. Liu, W. Lin, Y. Guo, and Z. Peng, "CASR: Optimizing cold start and resources utilization in serverless computing," *Futur. Gener. Comput. Syst.*, vol. 170, Sep. 2025. <https://doi.org/10.1016/j.future.2025.107851>
- [3] T. Singh, V. Shreshth, A. Raj, S. Swain, A. Bandyopadhyay, and N. Sharma, "Incentive Mechanisms for Federated Learning in Multi-Cloud Environments," *Proc. - 2025 7th Int. Conf. Comput. Intell. Commun. Technol. CCICT 2025*, pp. 302–307, 2025. <https://doi.org/10.1109/CCICT65753.2025.00054>
- [4] P. Tam, R. Corrado, C. Eang, and S. Kim, "Applicability of Deep Reinforcement Learning for Efficient Federated Learning in Massive IoT Communications," *Appl. Sci.*, vol. 13, no. 5, Mar. 2023. <https://doi.org/10.3390/app13053083>
- [5] Z. Shojaei Rad and M. Ghobaei-Arani, "Federated serverless cloud approaches: A comprehensive review," *Comput. Electr. Eng.*, vol. 124, May 2025. <https://doi.org/10.1016/j.compeleceng.2025.110372>
- [6] K. A. Ali, O. A. Fadare, and F. Al-Turjman, "Dynamic Resource Allocation (DRA) in Cloud Computing," in *Sustainable Civil Infrastructures*, vol. Part F4042, Springer Science and Business Media B.V., 2025, pp. 1033–1049. https://doi.org/10.1007/978-3-031-72509-8_85
- [7] C. Prigent, A. Costan, G. Antoniu, and L. Cudennec, "Enabling federated learning across the computing continuum: Systems, challenges and future directions," *Futur. Gener. Comput. Syst.*, vol. 160, pp. 767–783, Nov. 2024. <https://doi.org/10.1016/j.future.2024.06.043>
- [8] D. Ritter, "Cost-aware process modeling in multiclouds," *Inf. Syst.*, vol. 108, p. 101969, 2022. <https://doi.org/10.1016/j.is.2021.101969>
- [9] W. Khalafat, W. Elmedany, and H. Alryalat, "Privacy and security of federated learning in resource-constrained Internet of Things environment: Systematic literature review," *Internet Things (The Netherlands)*, vol. 33, Sep. 2025. <https://doi.org/10.1016/j.iot.2025.101679>

- [10] H. Qiu *et al.*, "SIMPPO: A Scalable and Incremental Online Learning Framework for Serverless Resource Management," in *SoCC 2022 - Proceedings of the 13th Symposium on Cloud Computing*, Nov. 2022, pp. 306–322. <https://doi.org/10.1145/3542929.3563475>
- [11] Y. Liu, F. Li, and C. Kong, "A generic framework for minimizing cold start times in serverless applications via resource serialization," *J. Supercomput.*, vol. 81, no. 12, Aug. 2025. <https://doi.org/10.1007/s11227-025-07710-z>
- [12] L. Albshaier, S. Almarri, and A. Albuali, "Federated Learning for Cloud and Edge Security: A Systematic Review of Challenges and AI Opportunities," *Electron.*, vol. 14, no. 5, Mar. 2025. <https://doi.org/10.3390/electronics14051019>
- [13] N. Singh and M. Adhikari, "SelfFed : Self-adaptive Federated Learning with Non-IID data on Heterogeneous Edge Devices for Bias Mitigation and Enhance Training," *Inf. Fusion*, vol. 118, no. October 2024, p. 102932, 2025. <https://doi.org/10.1016/j.inffus.2025.102932>
- [14] S. Demil and M. R. Abdmeziem, "Internet of Things Incentivizing task offloading in IoT: A distributed auctions-based DRL approach," *Internet of Things*, vol. 30, no. December 2024, p. 101493, 2025. <https://doi.org/10.1016/j.iot.2025.101493>
- [15] O. Bushehrian and A. Moazeni, "Deep reinforcement learning - based optimal deployment of IoT machine learning jobs in fog computing architecture," *Computing*, vol. 107, no. 1, pp. 1–25, 2025. <https://doi.org/10.1007/s00607-024-01353-3>
- [16] J. Yu, R. Zhou, B. Li, L. Wu, and S. Member, "Intelligent Frameworks for Minimizing Job Completion Time in Clustered Federated Learning," pp. 1–14, 2025. <https://doi.org/10.1109/TON.2025.3600674>
- [17] D. Ayepah-mensah, G. Sun, G. O. Boateng, S. Anokye, and G. Liu, "Blockchain-Enabled Federated Learning-Based Resource Allocation and Trading for Network Slicing in 5G," *IEEE/ACM Trans. Netw.*, vol. 32, no. 1, pp. 654–669, 2024. <https://doi.org/10.1109/TNET.2023.3297390>
- [18] C. Wang, T. Yao, T. Fan, S. Peng, C. Xu, and S. Yu, "Modeling on Resource Allocation for Age-Sensitive Mobile-Edge Computing Using Federated Multiagent Reinforcement Learning," *IEEE Internet Things J.*, vol. 11, no. 2, pp. 3121–3131, 2024. <https://doi.org/10.1109/JIOT.2023.3294535>
- [19] S. M. Rajagopal, M. Supriya, and R. Buyya, "FedSDM: Federated learning based smart decision making module for ECG data in IoT integrated Edge–Fog–Cloud computing environments," *Internet of Things (Netherlands)*, vol. 22, no. April, p. 100784, 2023. <https://doi.org/10.1016/j.iot.2023.100784>
- [20] N. Yellas, B. Addis, S. Boumerdassi, R. Riggio, and S. Secci, "Function Placement for In-network Federated Learning," *Comput. Networks*, vol. 256, no. February 2024, p. 110900, 2025. <https://doi.org/10.1016/j.comnet.2024.110900>
- [21] N. Hudson, H. Khamfroush, M. Baughman, D. E. Lucani, K. Chard, and I. Foster, "QoS-aware edge AI placement and scheduling with multiple implementations in FaaS-based edge computing," vol. 157, no. September 2023, pp. 250–263, 2024. <https://doi.org/10.1016/j.future.2024.03.035>
- [22] W. Feng, X. Zuo, R. Zhang, Y. Zhu, and C. Wang, "Federated Deep Reinforcement Learning for Multimodal Content Caching in Edge-Cloud Networks," *IEEE Trans. Netw. Sci. Eng.*, vol. 12, no. 3, pp. 2188–2201, 2025. <https://doi.org/10.1109/TNSE.2025.3545924>
- [23] A. A. Okine, N. Adam, F. Naeem, and G. Kaddoum, "FedRoute : A Multi-Server Federated Meta-DRL Routing Scheme for Tactical Air-Ground WSNs," *IEEE Open J. Commun. Soc.*, vol. 6, no. April, pp. 4176–4193, 2025. <https://doi.org/10.1109/OJCOMS.2025.3567024>
- [24] A. S. M. S. Sagar, A. Haider, and H. S. Kim, "A hierarchical adaptive federated reinforcement learning for efficient resource allocation and task scheduling in hierarchical IoT network," *Comput. Commun.*, vol. 229, no. July 2024, p. 107969, 2025. <https://doi.org/10.1016/j.comcom.2024.107969>
- [25] W. Hou, H. Wen, S. Member, N. Zhang, and S. Member, "Adaptive Training and Aggregation for Federated Learning in Multi-Tier Computing Networks," *IEEE Trans. Mob. Comput.*, vol. 23, no. 5, pp. 4376–4388, 2024. <https://doi.org/10.1109/TMC.2023.3289940>
- [26] C. Sun, S. Member, X. Li, and J. Wen, "Federated Deep Reinforcement Learning for Recommendation-Enabled Edge Caching in Mobile Edge-Cloud Computing Networks," *IEEE J. Sel. Areas Commun.*, vol. 41, no. 3, pp. 690–705, 2023. <https://doi.org/10.1109/JSAC.2023.3235443>
- [27] W. Fan, P. Chen, X. Chun, and Y. Liu, "MADRRL-Based Model Partitioning, Aggregation Control, and Resource Allocation for Cloud-Edge-Device Collaborative Split Federated Learning," *IEEE Trans. Mob. Comput.*, vol. 24, no. 6, pp. 5324–5341, 2025. <https://doi.org/10.1109/TMC.2025.3530482>
- [28] A. R. Malipatil, M. E. Paramasivam, D. Gulyamova, and A. Saravanan, "Energy-Efficient Cloud Computing Through Reinforcement Learning-Based Workload Scheduling," vol. 16, no. 4, pp. 645–656, 2025. <https://doi.org/10.14569/IJACSA.2025.0160464>
- [29] L. Xiao, H. Shan, J. Zhu, R. Mao, and S. Pan, "FD3QN : A Federated Deep Reinforcement Learning Approach for Cross-Domain Resource Cooperative Scheduling in Hybrid Cloud Architecture," vol. 49, pp. 127–146, 2025. <https://doi.org/10.31449/inf.v49i10.7114>
- [30] M. C. Sekhar, P. Kumaraswamy, N. Yamsani, G. B. K, and R. Kotoju, "FedTaskRL : A Reinforcement Learning-Based Framework for Efficient Task Scheduling in Federated Cloud Environments," vol. 12, no. 7, pp. 74–89, 2025. <https://doi.org/10.14445/23488549/IJECE-V12I7P107>
- [31] N. Ma, A. Tang, Z. Xiong, and F. Jiang, "A deep multi-agent reinforcement learning approach for the micro-service migration problem with affinity in the cloud," *Expert Syst. Appl.*, vol. 273, no. February, p. 126856, 2025. <https://doi.org/10.1016/j.eswa.2025.126856>
- [32] S. B. Tadele, W. Yahya, B. Kar, Y. D. Lin, Y. C. Lai, and F. G. Wakgra, "Optimizing the Ratio-Based Offloading in Federated Cloud-Edge Systems: A MADRL Approach," *IEEE Trans. Netw. Sci. Eng.*, vol. 12, no. 1, pp. 463–475, 2025. <https://doi.org/10.1109/TNSE.2024.3501398>
- [33] S. Najafli, A. Toroghi, and H. Babak, "A novel reinforcement learning - based hybrid intrusion detection system on fog - to - cloud computing," *J. Supercomput.*, vol. 80, no. 18, pp. 26088–26110, 2024. <https://doi.org/10.1007/s11227-024-06417-x>
- [34] B. Brik, S. Member, and M. Essegir, "On Adjusting Data Throughput in IoT Networks : Game Approach," *IEEE Internet Things J.*, vol. 11, no. 7, pp. 11368–11380, 2024. <https://doi.org/10.1109/JIOT.2023.3330408>
- [35] S. Cho, "DRL-Enabled Hierarchical Federated Learning Optimization for Data Heterogeneity Management in Multi-Access Edge Computing," *IEEE Access*, vol. 12, no. September, pp. 147209–147219, 2024. <https://doi.org/10.1109/ACCESS.2024.3473008>
- [36] H. Zhou, H. Wang, Z. Yu, G. Bin, M. Xiao, and J. Wu, "Federated Distributed Deep Reinforcement Learning for Recommendation-Enabled Edge Caching," *IEEE Trans. Serv. Comput.*, vol. 17, no. 6, pp. 3640–3656, 2024. <https://doi.org/10.1109/TSC.2024.3433579>

- [37] S. Vadigi, K. Sethi, D. Mohanty, and S. Prasad, "Journal of Information Security and Applications Federated reinforcement learning based intrusion detection system using dynamic attention mechanism," vol. 78, no. September, 2023. <https://doi.org/10.1016/j.jisa.2023.103608>
- [38] J. Shi, C. Li, Y. Guan, P. Cong, and J. Li, "Multi-UAV-assisted computation offloading in DT-based networks : A distributed deep reinforcement learning approach," *Comput. Commun.*, vol. 210, no. April, pp. 217–228, 2023. <https://doi.org/10.1016/j.comcom.2023.07.041>
- [39] P. Zhang, N. Chen, G. S. Member, and S. Li, "Multi-Domain Virtual Network Embedding Algorithm Based on Horizontal Federated Learning," vol. 18, pp. 3363–3375, 2023. <https://doi.org/10.1109/TIFS.2023.3279587>
- [40] D. Qiao, S. Guo, S. Member, D. Liu, and S. Long, "Adaptive Federated Deep Reinforcement Learning for Proactive Content Caching in Edge Computing," *IEEE Trans. Parallel Distrib. Syst.*, vol. 33, no. 12, pp. 4767–4782, 2022. <https://doi.org/10.1109/TPDS.2022.3201983>
- [41] M. Bazargani, S. Tarkesh, E. Behnam, and H. Reza, "AFedSLL-LDL: a framework based-on federated self-supervised learning and lightweight deep learning for attack detection in serverless edge computing," 2025. <https://doi.org/10.1007/s10586-025-05673-7>
- [42] G. Nagabushnam and K. Hoon, *Faddeer: a deep multi-agent reinforcement learning-based scheduling algorithm for aperiodic tasks in heterogeneous fog computing networks*, vol. 28, no. 6. Springer US, 2025. <https://doi.org/10.1007/s10586-025-05435-5>
- [43] M. Abdullah, R. Munir, W. Iqbal, and S. Ahmad, "Cost and performance-effective dynamic VM type selection in auto-scaling using reinforcement learning," 2025. <https://doi.org/10.1007/s00607-025-01535-7>
- [44] P. Zhang, E. Wang, M. Guizani, K. Liu, J. Wang, and L. Tan, "Privacy-Preserving Task Offloading in Vehicular Edge Computing," *IEEE Trans. Veh. Technol.*, vol. PP, no. Xx, pp. 1–13, 2025. <https://doi.org/10.1109/TVT.2025.3588204>
- [45] T. Zeng, X. Zhang, J. Duan, C. Yu, C. Wu, and X. Chen, "An Offline-Transfer-Online Framework for Cloud-Edge Collaborative Distributed," *IEEE Trans. Parallel Distrib. Syst.*, vol. 35, no. 5, pp. 720–731, 2024. <https://doi.org/10.1109/TPDS.2024.3360438>
- [46] A. H. D. Mendes, M. J. F. Rosa, M. A. Marotta, A. Araujo, A. C. M. A. Melo, and C. G. Ralha, "MAS-Cloud+: A novel multi-agent architecture with reasoning models for resource management in multiple providers," *Futur. Gener. Comput. Syst.*, vol. 154, no. February 2023, pp. 16–34, 2024. <https://doi.org/10.1016/j.future.2023.12.022>
- [47] Z. Zhang, K. Ning, and G. Wu, "Enhancing multi-cloud service deployment with SkyCap: A loss-aware coordinator in sky computing," *Ad Hoc Networks*, vol. 157, no. January, p. 103460, 2024. <https://doi.org/10.1016/j.adhoc.2024.103460>
- [48] C. Jiang and L. Su, "Federated learning with three-way decisions for privacy-preserving multicloud resource scheduling," *Appl. Soft Comput.*, vol. 183, no. October 2024, p. 113634, 2025. <https://doi.org/10.1016/j.asoc.2025.113634>
- [49] S. Salar, S. Mobina, R. Craciunescu, S. Maiduc, M. Bayram, and B. Arasteh, "Adaptive Resource Scheduling in Multi-Cloud Computing Using Recurrent Neural Forecasting and Memory-Based Metaheuristic Optimization," 2025. <https://doi.org/10.1007/s10723-025-09812-7>
- [50] A. John, J. Kawash, and R. Alhajj, *Predictive container orchestration in the cloud using artificial intelligence techniques*, vol. 107, no. 7. Springer Vienna, 2025. <https://doi.org/10.1007/s00607-025-01505-z>
- [51] B. Kumar, A. Verma, and P. Verma, *A multivariate transformer - based monitor - analyze - plan - execute (MAPE) autoscaling framework for dynamic resource allocation in cloud environment*, vol. 107, no. 3. Springer Vienna, 2025. <https://doi.org/10.1007/s00607-025-01426-x>
- [52] D. Dong, "Agent-based cloud simulation model for resource management," *J. Cloud Comput.*, vol. 12, no. 1, 2023. <https://doi.org/10.1186/s13677-023-00540-5>
- [53] H. Zhang, J. Wang, H. Zhang, and C. Bu, "Security computing resource allocation based on deep reinforcement learning in serverless multi-cloud edge computing," *Futur. Gener. Comput. Syst.*, vol. 151, no. September 2023, pp. 152–161, 2024. <https://doi.org/10.1016/j.future.2023.09.016>
- [54] M. Emami Khansari and S. Sharifian, "A deep reinforcement learning approach towards distributed Function as a Service (FaaS) based edge application orchestration in cloud-edge continuum," *J. Netw. Comput. Appl.*, vol. 233, no. August 2024, p. 104042, 2025. <https://doi.org/10.1016/j.jnca.2024.104042>
- [55] G. R. Russo, V. Cardellini, and F. Lo Presti, "A framework for offloading and migration of serverless functions in the Edge – Cloud Continuum," vol. 100, no. March, 2024. <https://doi.org/10.1016/j.pmcj.2024.101915>
- [56] M. Zhou, B. Zheng, and L. Pan, "Balancing function performance and cluster load in serverless computing: A reinforcement learning solution," *J. Netw. Comput. Appl.*, vol. 243, no. March, 2025. <https://doi.org/10.1016/j.jnca.2025.104299>
- [57] A. Mampage, S. Karunasekera, and R. Buyya, "Deep reinforcement learning for application scheduling in resource-constrained , multi-tenant serverless computing environments," *Futur. Gener. Comput. Syst.*, vol. 143, pp. 277–292, 2023. <https://doi.org/10.1016/j.future.2023.02.006>
- [58] G. R. Russo, D. Ferrarelli, D. Pasquali, V. Cardellini, and F. Lo Presti, "QoS-aware offloading policies for serverless functions in the Cloud-to-Edge continuum," *Futur. Gener. Comput. Syst.*, vol. 156, no. July 2023, pp. 1–15, 2024. <https://doi.org/10.1016/j.future.2024.02.019>
- [59] J. Kaur, I. Chana, and A. Bala, "MADQL : Multi-Agent Deep Q-Learning for Optimized Job Scheduling in Serverless Computing," *Arab. J. Sci. Eng.*, 2025. <https://doi.org/10.1007/s13369-025-10461-x>
- [60] A. Z. M. Ghobaei-arani and L. Esmaili, "An efficient function placement approach in serverless edge computing," *Computing*, vol. 107, no. 3, pp. 1–58, 2025. <https://doi.org/10.1007/s00607-025-01438-7>
- [61] S. Nastic, "Self - Provisioning Infrastructures for the Next Generation Serverless Computing," *SN Comput. Sci.*, 2024. <https://doi.org/10.1007/s42979-024-03022-w>