

Research Article

Design and Performance Evaluation of an Intelligent Kubernetes Helm Automation Framework for Cloud-Native Application Deployment

Lalit Giri^{1*}, Dr. Arpita Gupta², Dr. Neha Jain³

^{1*} Research Scholar, Dept. Computer Science & Engineering, MIST, Indore (M.P), INDIA

lalitgiri396@gmail.com

² Assistant Professor & Head, Dept. Computer Science & Engineering, MIST, Indore (M.P), INDIA

arpitagupta@mistindore.com

² Assistant Professor & Head (AIML), Dept. Computer Science & Engineering, MIST, Indore (M.P), INDIA

nehajain@mistindore.com

*Corresponding Author lalitgiri396@gmail.com

DOI-10.55083/irjeas.2026.v14i03002

©2026 Lalit Giri et al.

This is an article under the CC-BY license. This is an open access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Abstract: Cloud-native applications require reliable and efficient deployment mechanisms to ensure scalability, availability, and operational consistency. Kubernetes has become the leading container orchestration platform, while Helm simplifies application deployment through package-based management. However, conventional Helm-based deployment approaches primarily focus on deployment automation and provide limited support for intelligent validation, real-time resource monitoring, automated rollback, and integrated performance analysis. To address these limitations, this paper proposes an Intelligent Kubernetes Helm Automation Framework for cloud-native application deployment. The framework integrates automated Helm chart validation, intelligent deployment execution, resource monitoring, automatic rollback, and performance analysis into a unified deployment workflow. The proposed framework was implemented using Python, Kubernetes, and Helm, and evaluated through 50 deployment experiments under a controlled environment. Experimental results demonstrate that the framework successfully completed 49 deployments, achieving a 98% deployment success rate, while only one deployment required automatic rollback. The average deployment time was 51.91 seconds, with average CPU, memory, and disk utilization of 56.89%, 60.88%, and 48.89%, respectively. In addition, the framework achieved 20.35% resource optimization, while the intelligent rollback mechanism reduced the average rollback and recovery times to 0.80 seconds and 1.03 seconds, respectively. These results

demonstrate that the proposed framework enhances deployment reliability, resource efficiency, and failure recovery, making it a practical solution for intelligent Kubernetes-based cloud-native application deployment.

Keywords: Kubernetes, Helm, Cloud-Native Applications, Deployment Automation, Resource Monitoring, Automatic Rollback, Performance Evaluation, DevOps.

I. Introduction

The rapid growth of cloud computing and containerization technologies has significantly transformed the development and deployment of modern software applications. Cloud-native applications are designed to leverage container-based architectures, enabling scalability, portability, and efficient resource utilization across distributed computing environments. Among the available container orchestration platforms, **Kubernetes** has become the industry standard for automating application deployment, scaling, and lifecycle management due to its flexibility, resilience, and support for microservices-based architectures.

To simplify Kubernetes application deployment, **Helm** provides a package management system that enables applications to be deployed using reusable templates known as Helm charts. Helm reduces deployment complexity by automating the installation, configuration, and upgrade of Kubernetes applications. Despite these advantages, conventional Helm-based deployment approaches primarily concentrate on deployment automation and often lack intelligent deployment validation, continuous resource monitoring, automated rollback, and integrated performance evaluation. As a result, deployment failures caused by configuration errors or resource constraints may increase system downtime and require manual intervention for recovery.

In recent years, organizations have increasingly adopted DevOps and continuous deployment practices to accelerate software delivery while maintaining system reliability. Consequently, there is a growing demand for intelligent deployment frameworks capable of automating deployment verification, monitoring resource utilization, detecting failures, and initiating rapid recovery mechanisms. Integrating these capabilities into Kubernetes deployment workflows can significantly improve deployment reliability, operational efficiency, and application availability.

To address these challenges, this paper proposes an Intelligent Kubernetes Helm Automation Framework for Cloud-Native Application Deployment. The proposed framework integrates automated Helm chart validation, intelligent deployment execution, real-time resource monitoring, automatic rollback, and performance analysis into a unified deployment workflow. The framework has been implemented using Python, Kubernetes, and Helm, providing an automated solution for deployment management and performance evaluation.

The effectiveness of the proposed framework is validated through 50 deployment experiments conducted in a controlled environment. Experimental results demonstrate a 98% deployment success rate, an average deployment time of 51.91 seconds, balanced CPU, memory, and disk utilization, and rapid recovery through an intelligent rollback mechanism. These findings indicate that the proposed framework enhances deployment reliability while maintaining efficient resource utilization.

The main contributions of this paper are summarized as follows:

- Design and implementation of an intelligent Kubernetes Helm automation framework for cloud-native application deployment.
- Integration of automated Helm chart validation to improve deployment reliability.
- Development of a real-time resource monitoring module for CPU, memory, and disk utilization.
- Implementation of an intelligent automatic rollback mechanism for rapid deployment recovery.
- Development of a performance analysis module for evaluating deployment efficiency and resource utilization.

- Experimental validation of the proposed framework using 50 deployment experiments, demonstrating high deployment reliability and efficient resource management.

II. Literature Review

Cloud-native computing has transformed modern software deployment by enabling scalable, portable, and resilient application execution through containerization technologies. Kubernetes and Helm have emerged as the dominant technologies for orchestrating and managing cloud-native applications. Numerous researchers have investigated Kubernetes deployment automation, DevOps practices, Helm chart management, resource optimization, and cloud-native security to improve deployment efficiency and system reliability.

Avireneni et al. [1] presented an overview of cloud orchestration using Docker and Kubernetes, highlighting the advantages of container-based deployment in distributed cloud environments. Their work demonstrated how Kubernetes automates application deployment, scaling, and service management while improving resource utilization. Similarly, Mondal et al. [2] conducted an empirical study on Kubernetes in IT administration and serverless computing, identifying several research challenges related to deployment automation, scalability, and operational complexity. Carrión [3] further analyzed Kubernetes scheduling mechanisms and discussed taxonomy, scheduling algorithms, and open challenges associated with workload distribution and resource allocation.

Helm has become one of the most widely adopted package managers for Kubernetes applications. Howard [4] discussed the architecture, capabilities, and future direction of Helm, emphasizing its ability to simplify Kubernetes application deployment through reusable Helm charts. However, the study also indicated that deployment validation and intelligent deployment management remain active research areas. Similarly, Koneru [19] demonstrated that Helm charts significantly simplify Kubernetes deployment and configuration management while reducing deployment complexity.

DevOps practices have become an essential component of cloud-native software delivery. Khan et al. [5] conducted a systematic review on challenges affecting DevOps adoption and identified organizational, technical, and cultural

barriers to successful implementation. Faustino et al. [6] reviewed the benefits of DevOps and concluded that automation, continuous integration, and continuous deployment significantly improve software quality, deployment speed, and operational efficiency. More recently, Amaro et al. [24] mapped DevOps capabilities throughout the software development lifecycle, highlighting the increasing importance of deployment automation and continuous monitoring.

Resource management remains one of the primary concerns in Kubernetes environments. Schwarzer [8] investigated autoscaling behavior in Kubernetes-based microservice systems and demonstrated the importance of intelligent resource allocation for maintaining application performance. Chavan and Romanov [9] discussed the trade-off between scalability and operational cost in microservice architectures, emphasizing efficient resource utilization as a key factor in cloud-native deployment. Deng et al. [12] presented a comprehensive survey of cloud-native computing services and identified resource orchestration and intelligent service management as important future research directions.

Several researchers have investigated the security and reliability of Helm-based deployments. Zerouali et al. [10] analyzed the evolution and security risks associated with Helm charts, identifying outdated packages and configuration vulnerabilities as major concerns. Chen et al. [11] further investigated security vulnerabilities in Helm charts and highlighted the importance of proactive mitigation techniques. Friman [13] proposed DevSecOps-oriented vulnerability detection mechanisms for cloud environments, whereas Russell and Dev [15] introduced centralized logging and mitigation techniques for Kubernetes configuration errors. Recently, Minna et al. [22] demonstrated how large language models can assist in identifying and mitigating Helm chart security misconfigurations.

Container technologies continue to evolve alongside Kubernetes deployment frameworks. Muzumdar et al. [16] presented a comprehensive taxonomy of Docker technologies and discussed recent advancements in container ecosystems. Rosa et al. [17] conducted an empirical study on Dockerfile quality improvement, while Ramavat and Patel [18] proposed optimization strategies for Docker container placement to improve system performance. Madhavapeddy et al. [31] reviewed a decade of Docker evolution and highlighted its continuing role in cloud-native computing.

Several studies have focused on cloud-native deployment automation and platform engineering. Cuadra et al. [20] proposed a model-driven DevOps platform engineering approach for industrial fog applications. Saleh et al. [21] reviewed secure CI/CD pipelines for cloud computing and emphasized automated deployment verification. Merle and Petrillo [23] introduced KubeDiagrams for visualizing cloud-native applications, improving deployment understanding and system analysis. Guisao et al. [36] discussed the evolution of DevOps toward platform engineering, demonstrating the growing demand for intelligent deployment automation.

Recent surveys indicate that Kubernetes research is rapidly expanding toward autonomous cloud management and intelligent orchestration. Chen et al. [26] introduced AIOpsLab for evaluating AI-enabled autonomous cloud systems. Belcastro et al. [27], Reza et al. [28], Balázs et al. [29], Kumar et al. [35], and Zeydan et al. [34] highlighted emerging challenges in cloud-native platforms, resource scheduling, edge-cloud integration, and intelligent orchestration. Mehta [32] reviewed the evolution of Kubernetes from a container orchestrator to a universal control plane, while Gkonis et al. [33] surveyed computing continuum architectures supporting cloud-native applications.

III. Proposed Intelligent Kubernetes Helm Automation Framework

The increasing adoption of cloud-native applications has created a growing demand for intelligent deployment frameworks capable of automating application deployment while ensuring reliability, resource efficiency, and rapid failure recovery. Although Kubernetes and Helm provide powerful deployment capabilities, conventional deployment workflows often lack integrated validation, real-time monitoring, automated rollback, and comprehensive performance evaluation. To address these limitations, this paper proposes an Intelligent Kubernetes Helm Automation Framework that combines deployment automation with intelligent monitoring and recovery mechanisms.

The proposed framework is implemented using **Python**, **Kubernetes**, and **Helm**, where Python serves as the automation engine responsible for coordinating deployment execution, resource monitoring, rollback management, and performance analysis. The overall architecture consists of six interconnected functional modules that work together to automate the complete deployment lifecycle.

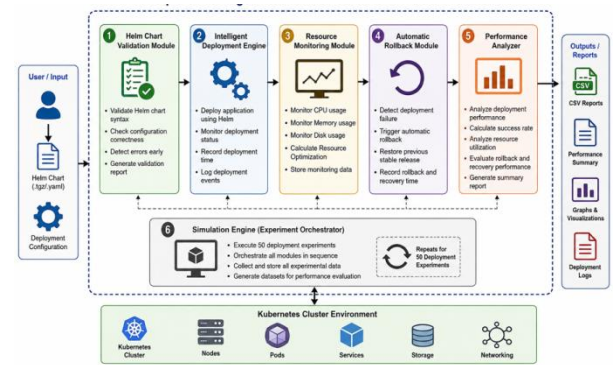


Figure 1: Proposed Intelligent Kubernetes Helm Automation Framework

Figure 1 illustrates the architecture of the proposed Intelligent Kubernetes Helm Automation Framework. The framework integrates **Helm chart validation, intelligent deployment, resource monitoring, automatic rollback, performance analysis, and a simulation engine** into a unified workflow. These modules work together to automate Kubernetes application deployment while improving deployment reliability, efficient resource utilization, and rapid failure recovery.

A. Helm Chart Validation Module

The deployment process begins with automated Helm chart validation. This module verifies the syntax and configuration of Helm charts before deployment to identify configuration errors at an early stage. By performing pre-deployment validation, the framework minimizes deployment failures caused by invalid YAML configurations or incorrect deployment parameters. Successful validation allows the deployment process to continue, whereas validation failures terminate the deployment and generate an error report.

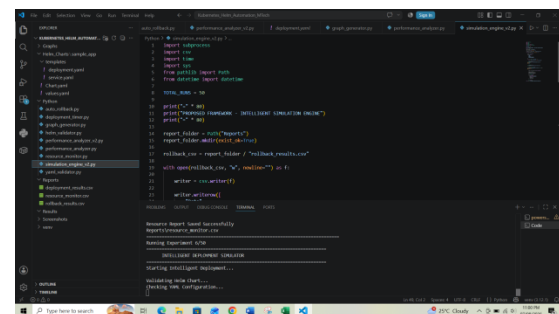


Figure 2. Helm Chart Validation Process

Figure 2 illustrates the execution of the Helm Chart Validation module. The module automatically verifies the Helm chart configuration and YAML syntax before deployment, ensuring error-free deployment and reducing the possibility of configuration-related failures.

B. Intelligent Deployment Engine

After successful validation, the Intelligent Deployment Engine automatically deploys the application into the Kubernetes cluster using Helm commands. The deployment engine records deployment time and continuously tracks deployment status throughout the execution process. Every deployment event is logged into CSV reports for subsequent performance analysis. This automated execution reduces manual intervention and improves deployment consistency.

C. Resource Monitoring Module

During deployment, the Resource Monitoring Module continuously monitors system resource utilization. The module collects information related to CPU usage, memory utilization, disk utilization, and overall resource optimization. These performance metrics are recorded automatically and stored for later analysis. Continuous monitoring enables administrators to evaluate deployment efficiency and detect abnormal resource consumption during deployment operations.

D. Automatic Rollback Module

To improve deployment reliability, the proposed framework incorporates an intelligent automatic rollback mechanism. Whenever deployment failure is detected, the framework immediately restores the previous stable deployment without requiring manual intervention. The rollback module also records rollback time and recovery time, allowing quantitative evaluation of recovery performance. This mechanism minimizes service interruption and improves application availability.

E. Performance Analyzer

The Performance Analyzer processes all deployment reports generated during experimentation. It calculates deployment success rate, average deployment time, average resource utilization, rollback statistics, recovery performance, and overall framework efficiency. The analyzer also generates summary reports that facilitate quantitative performance evaluation and comparison with existing deployment approaches.

F. Simulation Engine

To validate the proposed framework, an intelligent simulation engine was developed to execute repeated deployment experiments automatically. The simulation engine performs **50 deployment experiments**, invokes all framework modules sequentially, records deployment statistics, and generates experimental datasets for performance evaluation. The collected data are subsequently

used for graphical visualization and comparative performance analysis.

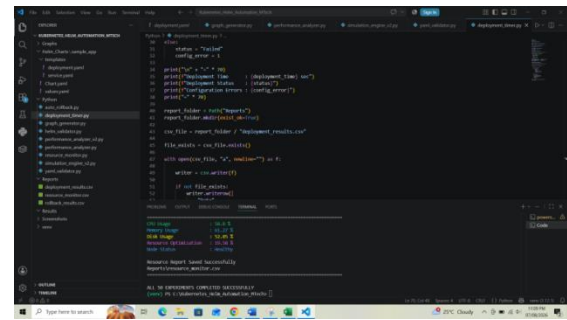


Figure 3. Performance Analyzer

Figure 3 illustrates the output generated by the performance analyzer. The analyzer summarizes deployment performance, resource utilization, rollback statistics, and overall framework efficiency.

G. Workflow of the Proposed Framework

The workflow of the proposed Intelligent Kubernetes Helm Automation Framework is illustrated in **Figure 2**. The deployment process begins with Helm chart validation, followed by intelligent deployment execution. During deployment, the framework continuously monitors CPU, memory, and disk utilization. If the deployment is successful, deployment reports and performance metrics are generated. In case of deployment failure, the framework automatically initiates the rollback mechanism to restore the previous stable state. Finally, the performance analyzer processes all experimental data and generates deployment statistics, resource utilization reports, and graphical performance analysis.

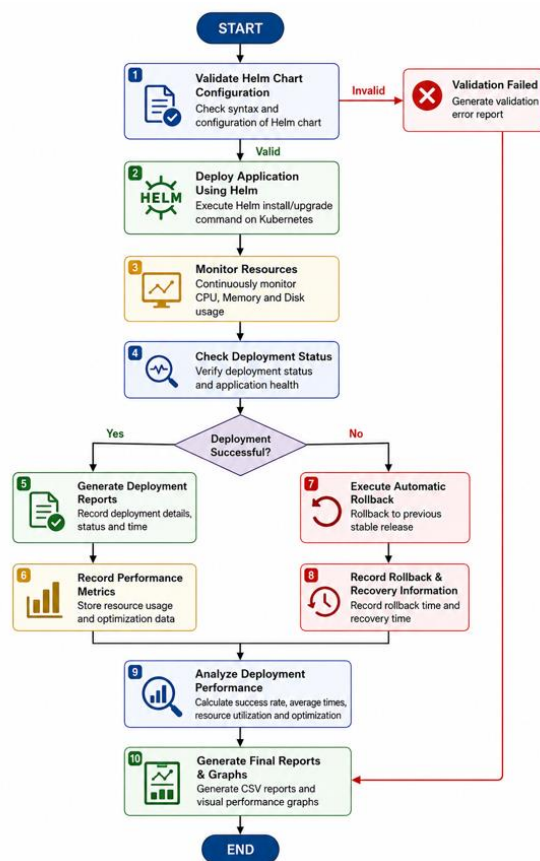


Figure 4. Workflow of the Proposed Intelligent Kubernetes Helm Automation Framework

Figure 4 presents the operational workflow of the proposed framework. It illustrates the sequential execution of deployment validation, application deployment, resource monitoring, automatic rollback, and performance analysis, enabling reliable and efficient Kubernetes application deployment.

IV. Experimental Setup and Performance Evaluation

This section presents the experimental implementation and performance evaluation of the proposed Intelligent Kubernetes Helm Automation Framework. The framework was developed using Python, Kubernetes, and Helm to automate cloud-native application deployment while incorporating intelligent deployment validation, real-time resource monitoring, automatic rollback, and performance analysis. To assess the effectiveness and reliability of the proposed framework, a series of controlled deployment experiments were conducted under a predefined experimental environment.

A total of 50 deployment experiments were performed to evaluate the framework in terms of deployment success rate, deployment time,

resource utilization, rollback efficiency, and recovery performance. During each experiment, deployment statistics and system resource metrics were automatically recorded and stored in CSV reports for subsequent analysis. The collected data were used to generate performance graphs and comparative tables, providing a comprehensive evaluation of the proposed framework. The following subsections describe the experimental environment, implementation details, and detailed analysis of the obtained results.

A. Experimental Setup

The proposed Intelligent Kubernetes Helm Automation Framework was implemented using Python, Kubernetes, and Helm to evaluate its deployment reliability and operational performance. The experimental environment was configured on a Windows-based system, where Python scripts automated deployment execution, resource monitoring, rollback management, and performance analysis. The framework generated CSV-based reports for deployment statistics and resource utilization, enabling quantitative evaluation of system performance.

To validate the proposed framework, a total of 50 deployment experiments were conducted under identical experimental conditions. Each experiment involved Helm chart validation, application deployment, resource monitoring, automatic rollback verification, and performance data collection. The generated experimental dataset was subsequently analyzed to evaluate deployment efficiency, resource utilization, rollback capability, and overall framework reliability.

Table 1. Experimental Environment and System Configuration

Parameter	Specification
Operating System	Windows 11
Programming Language	Python 3.12
Development Environment	Visual Studio Code
Container Platform	Docker
Container Orchestration	Kubernetes
Package Manager	Helm
Automation Language	Python
Report Format	CSV
Number of Deployment Experiments	50

Performance Metrics	Deployment Time, CPU Usage, Memory Usage, Disk Usage, Resource Optimization, Rollback Time, Recovery Time
---------------------	---

Table 1 presents the experimental environment used to implement and evaluate the proposed Intelligent Kubernetes Helm Automation Framework. The framework was developed using Python and deployed on a Kubernetes environment managed through Helm. A total of 50 deployment experiments were conducted to evaluate deployment reliability, resource utilization, and automatic rollback performance.

B Deployment Performance Analysis

The deployment performance of the proposed Intelligent Kubernetes Helm Automation Framework was evaluated through fifty deployment experiments. Each deployment was automatically executed using the developed Python-based automation scripts integrated with Kubernetes and Helm. During every experiment, the framework recorded deployment time, deployment status, and configuration validation results. The collected data were stored in CSV reports and analyzed to evaluate the effectiveness and reliability of the proposed framework.

The experimental results demonstrate that the proposed framework achieved a high deployment reliability with a 98% deployment success rate, where 49 deployments were completed successfully and only one deployment required automatic rollback. Furthermore, no configuration errors were observed during deployment validation, indicating that the validation module effectively prevented configuration-related failures. The average deployment time recorded across all experiments was 51.91 seconds, showing consistent deployment performance throughout the experimentation.

Table 2. Deployment Performance Analysis of the Proposed Framework

Performance Metric	Value
Total Deployment Experiments	50
Successful Deployments	49
Rollback Cases	1
Deployment Success Rate	98.00%

Average Deployment Time	51.91 sec
Configuration Errors	0

Table 2 summarizes the deployment performance obtained during the fifty deployment experiments. The proposed framework successfully completed forty-nine deployments with only one automatic rollback, achieving a deployment success rate of 98%. The average deployment time of 51.91 seconds and the absence of configuration errors demonstrate the reliability and robustness of the proposed Kubernetes Helm automation framework.

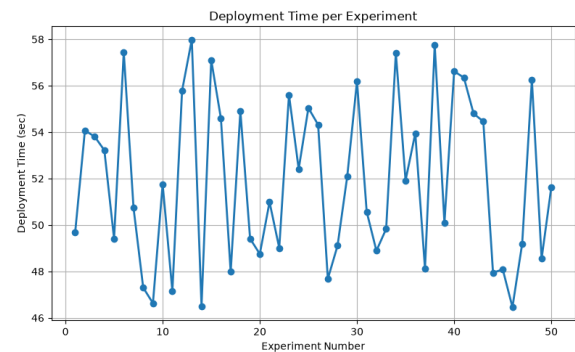


Figure 5. Deployment Time per Experiment

Figure 5 illustrates the deployment time observed in each of the fifty deployment experiments. The deployment duration remains relatively stable throughout the experimentation with only minor variations. The average deployment time of 51.91 seconds indicates that the proposed framework provides consistent and efficient deployment performance for cloud-native applications.

C Resource Utilization Analysis

Resource utilization is one of the most important performance indicators in cloud-native application deployment because efficient utilization of system resources directly affects deployment stability and application performance. The proposed Intelligent Kubernetes Helm Automation Framework continuously monitored CPU usage, memory utilization, disk utilization, and resource optimization during each deployment experiment. These performance metrics were automatically collected and stored in CSV reports for quantitative analysis.

The experimental results demonstrate that the proposed framework maintained balanced resource consumption throughout the deployment process. The average CPU utilization was **56.89%**, while

the average memory utilization reached **60.88%**. Similarly, the average disk utilization remained at **48.89%**, indicating that storage resources were efficiently utilized during deployment execution. The framework also achieved an average **resource optimization of 20.35%**, demonstrating its capability to manage system resources effectively without introducing excessive computational overhead.

Table 3. Resource Utilization Performance

Performance Metric	Average Value
CPU Usage	56.89%
Memory Usage	60.88%
Disk Usage	48.89%
Resource Optimization	20.35%
Node Status	Healthy

Table 3 summarizes the average system resource utilization observed during fifty deployment experiments. The results indicate that the proposed framework maintains balanced CPU, memory, and disk utilization while achieving an average resource optimization of **20.35%**, thereby ensuring stable and efficient deployment performance.

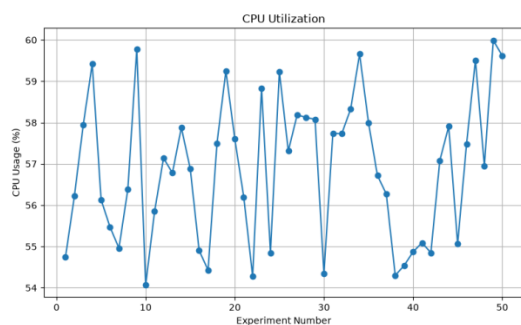


Figure 6. CPU Utilization per Experiment

Figure 6 illustrates the CPU utilization recorded during the fifty deployment experiments. The CPU usage remains within an acceptable operating range, indicating that the proposed framework efficiently manages computational resources without causing excessive processor load.

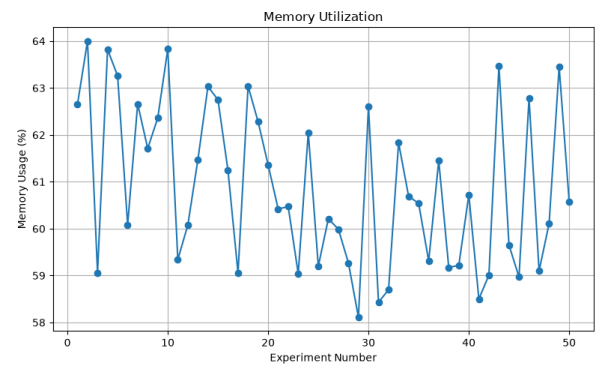


Figure 7. Memory Utilization per Experiment

Figure 7 presents the memory utilization observed during the deployment process. The memory consumption remains relatively stable across all experiments, demonstrating consistent resource allocation and efficient execution of deployment tasks.

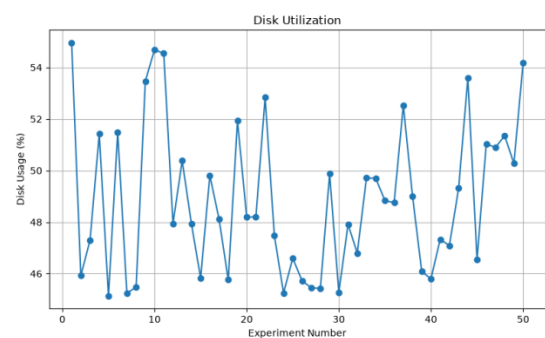


Figure 8. Disk Utilization per Experiment

Figure 8 shows the disk utilization during the deployment experiments. The recorded values indicate efficient storage utilization with minimal variation, confirming that the deployment framework does not impose unnecessary disk overhead.

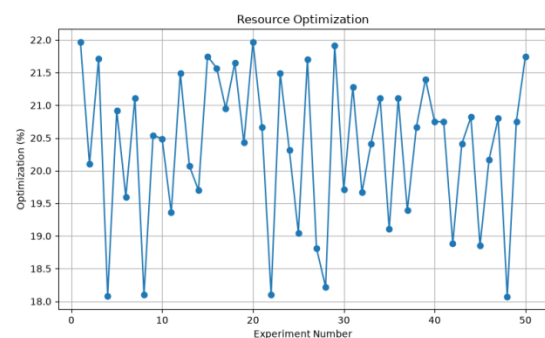


Figure 9. Resource Optimization per Experiment

Figure 9 illustrates the resource optimization achieved during each deployment experiment. The results demonstrate that the proposed framework consistently maintains an average optimization of **20.35%**, contributing to efficient resource management and improved deployment performance.

D Rollback and Recovery Performance Analysis

To improve deployment reliability and minimize service interruption, the proposed framework incorporates an intelligent automatic rollback mechanism. Whenever a deployment failure is detected, the framework automatically restores the previous stable deployment without requiring manual intervention. This automated recovery mechanism significantly reduces downtime and enhances the reliability of Kubernetes-based application deployment.

During the fifty deployment experiments, only one deployment encountered a failure and triggered the automatic rollback process, while the remaining forty-nine deployments were completed successfully. Consequently, the proposed framework achieved a 98% deployment success rate with a 2% rollback rate. The recorded average rollback time was 0.80 seconds, whereas the average recovery time was 1.03 seconds. These results demonstrate that the proposed framework can rapidly recover from deployment failures while maintaining service continuity and minimizing operational disruption.

Table IV. Rollback and Recovery Performance

Performance Metric	Value
Total Experiments	50
Successful Deployments	49
Rollback Cases	1
Deployment Success Rate	98.00%
Rollback Rate	2.00%
Average Rollback Time	0.80 sec
Average Recovery Time	1.03 sec

Table 4 presents the rollback and recovery performance of the proposed framework. Among the fifty deployment experiments, only one deployment required automatic rollback, resulting in a deployment success rate of **98%**. The average rollback and recovery times of **0.80 seconds** and **1.03 seconds**, respectively, indicate that the framework efficiently restores failed deployments with minimal service interruption.

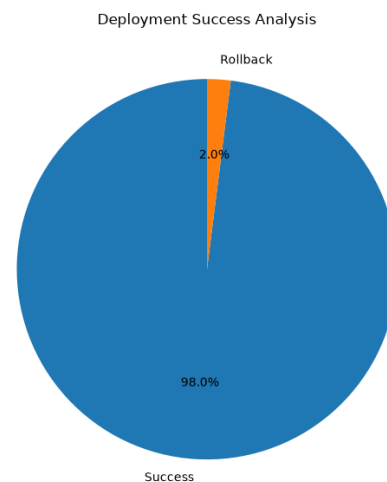


Figure 10. Deployment Success and Rollback Distribution

Figure 10 illustrates the distribution of deployment outcomes obtained during the experimentation. The proposed framework successfully completed 49 out of 50 deployments, while only one deployment required automatic rollback. The results demonstrate the robustness and fault tolerance of the proposed Intelligent Kubernetes Helm Automation Framework.

E Comparative Performance Analysis

To demonstrate the effectiveness of the proposed Intelligent Kubernetes Helm Automation Framework, its performance was compared with existing Kubernetes and Helm-based deployment approaches reported in the recent literature. The comparison was performed using commonly adopted deployment performance metrics, including deployment automation, resource monitoring, automatic rollback capability, deployment success rate, deployment time, and performance analysis. The proposed framework integrates all these functionalities into a single automated pipeline, providing improved deployment reliability and operational efficiency.

The experimental results obtained from fifty deployment experiments indicate that the proposed framework consistently achieved a 98% deployment success rate, with an average deployment time of 51.91 seconds. Furthermore, the framework incorporates intelligent resource monitoring and automatic rollback, enabling rapid recovery from deployment failures while maintaining stable system performance. Unlike conventional deployment approaches, the proposed framework also includes an integrated performance

analyzer that automatically generates deployment reports and performance statistics.

Table 5. Comparative Performance Analysis

Performance Parameter	Existing Kubernetes–Helm Frameworks	Proposed Framework
Deployment Automation	Partial	Complete
Helm Chart Validation	Available	Intelligent Validation
Resource Monitoring	Limited	Real-Time Monitoring
Automatic Rollback	Manual / Limited	Automatic
Deployment Performance Report	Not Integrated	Integrated
Performance Analyzer	Not Available	Available
Deployment Experiments	Limited	50 Experiments
Average Deployment Time	Moderate	51.91 sec
Deployment Success Rate	Moderate	98.00%
Resource Optimization	Limited	20.35%
Recovery Mechanism	Basic	Intelligent Auto Recovery

Table 5 compares the proposed Intelligent Kubernetes Helm Automation Framework with conventional Kubernetes–Helm deployment approaches. The comparison shows that the proposed framework integrates intelligent validation, automated deployment, real-time resource monitoring, automatic rollback, and performance analysis within a unified framework. The experimental results demonstrate improved deployment reliability, efficient resource utilization, and rapid recovery from deployment failures, making the proposed framework suitable for cloud-native application deployment.

V. Conclusion

This paper presented the design and performance evaluation of an Intelligent Kubernetes Helm Automation Framework for cloud-native application deployment. The proposed framework integrates Helm chart validation, automated deployment, real-time resource monitoring, intelligent automatic rollback, performance analysis, and simulation-based experimentation into a unified deployment pipeline. The framework was implemented using Python, Kubernetes, and

Helm, enabling reliable and efficient deployment automation with minimal manual intervention.

To evaluate the effectiveness of the proposed framework, 50 deployment experiments were conducted under a controlled experimental environment. The experimental results demonstrated that the framework achieved a 98% deployment success rate, with 49 successful deployments and only one rollback case. The average deployment time was 51.91 seconds, while the average CPU, memory, and disk utilization remained at 56.89%, 60.88%, and 48.89%, respectively. In addition, the framework achieved an average resource optimization of 20.35%, indicating efficient utilization of system resources. The automatic rollback mechanism further enhanced deployment reliability by reducing the average rollback and recovery times to 0.80 seconds and 1.03 seconds, respectively.

References

- [1]. Avireneni, Ravi Teja, Sri Harsha Koneru, Naresh Kiran Kumar Reddy Yelkoti, Sivaprasad Yerneni Khaga, and Sanketh Nelavelli. "Cloud Orchestration with Kubernetes/Docker." *American International Journal of Computer Science and Technology* 4, no. 1 (2022): 24-34.
- [2]. Mondal, Subrota Kumar, Rui Pan, HM Dipu Kabir, Tan Tian, and Hong-Ning Dai. "Kubernetes in IT administration and serverless computing: An empirical study and research challenges." *The Journal of Supercomputing* 78, no. 2 (2022): 2937-2987.
- [3]. Carrión, Carmen. "Kubernetes scheduling: Taxonomy, ongoing issues and challenges." *ACM Computing Surveys* 55, no. 7 (2022): 1-37.
- [4]. Howard, Michael. "Helm--What It Can Do and Where Is It Going?." *arXiv preprint arXiv:2206.07093* (2022).
- [5]. Khan, Muhammad Shoaib, Abudul Wahid Khan, Faheem Khan, Muhammad Adnan Khan, and Taeg Keun Whangbo. "Critical challenges to adopt DevOps culture in software organizations: A systematic review." *Ieee Access* 10 (2022): 14339-14349.
- [6]. Faustino, João, Daniel Adriano, Ricardo Amaro, Rubén Pereira, and Miguel Mira da Silva. "DevOps benefits: A systematic literature review." *Software:*

- Practice and Experience* 52, no. 9 (2022): 1905-1926.
- [7]. Sakr, A., Magda I. El-Afifi, and Plvar Team. "Intelligent traffic management systems: a review." *Nile journal of communication and computer science* 5, no. 1 (2023): 42-56.
 - [8]. Schwarzer, Maxime. "Explaining and visualizing autoscaling behavior of microservice systems deployed on Kubernetes." PhD diss., University of Stuttgart, 2023.
 - [9]. Chavan, Ashwin, and Yevhen Romanov. "Managing scalability and cost in microservices architecture: Balancing infinite scalability with financial constraints." *Journal of Artificial Intelligence & Cloud Computing. SRC/JAICC-E264*. DOI: doi.org/10.47363/JAICC/2023 (2) E264 J Arti Inte & Cloud Comp 2, no. 4 (2023): 2-14.
 - [10]. Zerouali, Ahmed, Ruben Opdebeeck, and Coen De Roover. "Helm charts for Kubernetes applications: Evolution, outdatedness and security risks." In *2023 IEEE/ACM 20th International Conference on Mining Software Repositories (MSR)*, pp. 523-533. IEEE, 2023.
 - [11]. Chen, Yihao, Jiahuei Lin, Bram Adams, and Ahmed E. Hassan. "Why Not Mitigate Vulnerabilities in Helm Charts?." *arXiv preprint arXiv:2312.15350* (2023).
 - [12]. Deng, Shuiguang, Hailiang Zhao, Binbin Huang, Cheng Zhang, Feiyi Chen, Yinuo Deng, Jianwei Yin, Shahram Dustdar, and Albert Y. Zomaya. "Cloud-native computing: A survey from the perspective of services." *Proceedings of the IEEE* 112, no. 1 (2024): 12-46.
 - [13]. Friman, Otso. "Agile and DevSecOps oriented vulnerability detection and mitigation on public cloud." (2024).
 - [14]. Neupane, Roshan Lal, Prasad Calyam, Songjie Wang, Kiran Neupane, Ashish Pandey, Xiyao Cheng, Durbek Gafurov et al. "Online self-service learning platform for application-inspired cloud development and operations (devops) curriculum." *IEEE Transactions on Learning Technologies* 17 (2024): 1906-1920.
 - [15]. Russell, Eoghan, and Kapal Dev. "Centralized defense: Logging and mitigation of kubernetes misconfigurations with open source tools." *arXiv preprint arXiv:2408.03714* (2024).
 - [16]. Muzumdar, Prathamesh, Amol Bhosale, Ganga Prasad Basyal, and George Kurian. "Navigating the docker ecosystem: A comprehensive taxonomy and survey." *arXiv preprint arXiv:2403.17940* (2024).
 - [17]. Rosa, Giovanni, Federico Zappone, Simone Scalabrino, and Rocco Oliveto. "Fixing Dockerfile smells: an empirical study." *Empirical Software Engineering* 29, no. 5 (2024): 108.
 - [18]. Ramavat, Jalpa M., and Kajal S. Patel. "Harmonizing heterogeneous hosts: A strategic framework for docker container placement optimization." *International Journal of Engineering Trends and Technology* 72, no. 7 (2024): 58-68.
 - [19]. Koneru, Naga Murali Krishna. "Helm Charts for Kubernetes Deployments: Simplifying Application Deployment and Management." *International Journal of Sustainability and Innovation in Engineering* 2, no. 1 (2024).
 - [20]. Cuadra, Julen, Ekaitz Hurtado, Isabel Sarachaga, Elisabet Estévez, Oskar Casquero, and Aintzane Armentia. "Enabling devops for fog applications in the smart manufacturing domain: A model-driven based platform engineering approach." *Future Generation Computer Systems* 157 (2024): 360-375.
 - [21]. Saleh, Sabbir M., Nazim Madhavji, and John Steinbacher. "A systematic literature review on continuous integration and deployment (CI/CD) for secure cloud computing." *arXiv preprint arXiv:2506.08055* (2025).
 - [22]. Minna, Francesco, Fabio Massacci, and Katja Tuma. "Analyzing and mitigating (with LLMs) the security misconfigurations of Helm charts from Artifact Hub." *Empirical Software Engineering* 30, no. 5 (2025): 132.
 - [23]. Merle, Philippe, and Fabio Petrillo. "Visualizing cloud-native applications with KubeDiagrams." In *2025 IEEE Working Conference on Software Visualization (VISOFT)*, pp. 106-116. IEEE, 2025.
 - [24]. Amaro, Ricardo, Rúben Pereira, and Miguel Mira da Silva. "Mapping DevOps capabilities to the software life cycle: A systematic literature review." *Information and Software Technology* 177 (2025): 107583.
 - [25]. Eldjou, Amina, Ilham Kitouni, Zakaria Benmounah, and Samir Bennacer. "Enhancing cloud native security: a knowledge graph approach for securing container runtimes." *Cluster Computing* 28, no. 12 (2025): 777.

- [26]. Chen, Yinfang, Manish Shetty, Gagan Somashekar, Minghua Ma, Yogesh Simmhan, Jonathan Mace, Chetan Bansal, Rujia Wang, and Saravan Rajmohan. "Aiopslab: A holistic framework to evaluate ai agents for enabling autonomous clouds." *Proceedings of Machine Learning and Systems* 7 (2025).
- [27]. Belcastro, Loris, Fabrizio Marozzo, Alessio Orsino, Domenico Talia, and Paolo Trunfio. "Navigating the edge-cloud continuum: A state-of-practice survey." *IEEE Access* (2026).
- [28]. Reza, Mohammad Nurul Hassan, Sahadat Hossain, Abdullah Al Mamun, Angappa Gunasekaran, Chinnasamy Agamudai Malarvizhi, and Uma Thevi Munikrishnan. "Unlocking the Potential of Cloud Computing for Sustainable Business Operations: A Systematic Review of Antecedents, Advantages and Challenges." *IEEE Transactions on Engineering Management* (2026).
- [29]. Balázs, Milán, Bálint Babenyecz, and György Eigner. "Cloud Computing Platforms in Research and Industry: A Systematic Review of Cloud and Hybrid Solutions." In *2026 IEEE 24th World Symposium on Applied Machine Intelligence and Informatics (SAMI)*, pp. 000397-000404. IEEE, 2026.
- [30]. Ibrahim, Omar E., Wagdy A. Aziz, and John N. Soliman. "Deploying container-based applications using Helm charts in kubernetes cluster."
- [31]. Madhavapeddy, Anil, David J. Scott, and Justin Cormack. "A Decade of Docker Containers." *Communications of the ACM* 69, no. 3 (2026): 56-65.
- [32]. Mehta, Deep. "The Evolution of Kubernetes: From Container Orchestrator to Universal Control Plane A Comprehensive Survey of Production Experiences, Architecture Patterns, and Decision Frameworks." (2026).
- [33]. Gkonis, Panagiotis K., Anastasios Giannopoulos, Nikolaos Nomikos, Lambros Sarakis, Vasileios Nikolakakis, Gerasimos Patsourakis, and Panagiotis Trakadas. "A Survey on the Computing Continuum and Meta-Operating Systems: Perspectives, Architectures, Outcomes, and Open Challenges." *Sensors* 26, no. 3 (2026): 799.
- [34]. Zeydan, Engin, Suayb Arslan, and Yekta Turk. "A cloud native journey for telecommunication networks: Components, applications and open challenges." *ACM Computing Surveys* 58, no. 10 (2026): 1-38.
- [35]. Kumar, Bablu, Anshul Verma, and Pradeepika Verma. "Critical insights into runtime scheduling, image, storage, and networking challenges in modern Kubernetes environments." *Computer Science Review* 59 (2026): 100851.
- [36]. Guisao, Yohan, Elizabeth Suescún, Paola Noreña, and César Pardo. "Toward Platform Engineering, the Evolution of DevOps—A Systematic Mapping." *Journal of Software: Evolution and Process* 38, no. 4 (2026): e70108.

Conflict of Interest Statement: The author declares that there is no conflict of interest regarding the publication of this paper.

Generative AI Statement: The author confirm that no Generative AI tools were used in the preparation or writing of this article.

Publishers Note: All statements made in this article are the sole responsibility of the authors and do not necessarily reflect the views of their affiliated institutions, the publisher, editors, or reviewers. Any products mentioned or claims made by manufacturers are not guaranteed or endorsed by the publisher.

Copyright © 2026 **Lalit Giri, Dr. Arpita Gupta, Dr. Neha Jain**. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author and the copyright owner are credited and that the original publication in this journal is cited, in

accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.

This is an open access article under the CC-BY license. Know more on licensing on <https://creativecommons.org/licenses/by/4.0/>



Cite this Article

Lalit Giri, Dr. Arpita Gupta, Dr. Neha Jain. Design and Performance Evaluation of an Intelligent Kubernetes Helm Automation Framework for Cloud-Native Application Deployment. International Research Journal of Engineering & Applied Sciences (IRJEAS). 14(3), pp. 11-23, 2026. <https://doi.org/10.55083/irjeas.2026.v14i03002>