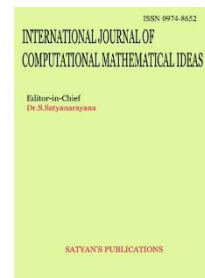




INTERNATIONAL JOURNAL OF
COMPUTATIONAL MATHEMATICAL IDEAS

Journal homepage: <https://ijcmi.in>



Optimizing Regional Disaster Recovery in OpenShift: A Multi-Cluster Approach with RHACM and ODF

Naga Venkata Chaitanya Akula

Senior Computer and Information Research Scientist

668 Pickrell loop, liberty hill, Texas-78642

E-Mail: chaitanya.akula@hitsslhc.com , chaitu.9616@gmail.com

Abstract: In the current digital environment, it is crucial for businesses to implement robust disaster recovery (DR) strategies to counteract risks posed by cyber threats, hardware malfunctions, and natural calamities. This paper examines an improved Regional Disaster Recovery (Regional-DR) strategy within OpenShift, which incorporates Red Hat Advanced Cluster Management (RHACM) and OpenShift Data Foundation (ODF) to enable smooth application failover and data replication. The practical application of Regional-DR showcases how applications and storage can be synchronized across several geographically dispersed clusters. The solution achieved a Recovery Time Objective (RTO) of 5 minutes and a Recovery Point Objective (RPO) of 3 minutes, ensuring minimal downtime and data loss. Furthermore, best practices for maintaining business continuity and enhancing system resilience using OpenShift's built-in DR features are analysed. The findings demonstrate how enterprises can utilize OpenShift's multi-cluster disaster recovery capabilities for effective failover management and improved performance.

Key words: Open Shift Disaster Recovery, Regional Disaster Recovery (Regional-DR) Red Hat Advanced Cluster Management (RHACM), OpenShift Data Foundation (ODF) Multi-Cluster Management, Kubernetes Disaster Recovery, Failover and Failback Automation, Persistent Storage Replication, Recovery Time Objective (RTO), Recovery Point Objective (RPO), Cloud-Native Disaster Recovery, Data Synchronization, Resilient Infrastructure

1. Introduction:

In today's rapidly evolving digital landscape, businesses of all sizes and industries face an array of complex challenges that can potentially disrupt their operations and compromise their ability to deliver services to customers. These challenges include increasingly sophisticated cyber threats, unexpected hardware failures, and unpredictable natural disasters. To effectively mitigate these risks and ensure uninterrupted business continuity, organizations must implement comprehensive and robust disaster recovery (DR) strategies that can adapt to the ever-changing technological environment.

This paper delves into an advanced and innovative approach to Regional Disaster Recovery (Regional-DR) specifically tailored for OpenShift environments, leveraging the powerful capabilities of Red Hat Advanced Cluster Management (RHACM) and OpenShift Data Foundation (ODF). The proposed solution focuses on optimizing disaster recovery processes by seamlessly integrating multiple geographically distributed clusters, creating a resilient and highly available infrastructure that can withstand various types of disruptions.

This multi-cluster approach enables seamless application failover and efficient data replication across different regions, significantly enhancing overall system resilience and minimizing the impact of localized outages or disasters. By utilizing OpenShift's built-in DR features and adhering to industry best practices, enterprises can maximize their ability to recover from potential disasters quickly and effectively, ensuring that critical business functions remain operational even in the face of adversity.

The practical implementation of this Regional-DR strategy demonstrates significant improvements in key performance metrics that are crucial for effective disaster recovery. Notably, the solution achieves an impressive Recovery Time Objective (RTO) of just 5 minutes, indicating the speed at which systems and applications can be restored to full functionality following a disaster event. Additionally, the solution boasts a Recovery Point Objective (RPO) of only 3 minutes, which represents the maximum amount of data that may be lost during a disaster scenario.

These remarkable figures underscore the solution's ability to minimize downtime and data loss, which are crucial factors in maintaining business operations and customer trust during critical events. By achieving such low RTO and RPO values, organizations can significantly reduce the financial and reputational impacts often associated with prolonged service interruptions or data loss incidents.

The multi-cluster disaster recovery approach presented in this research leverages OpenShift's advanced features, including automated failover mechanisms, intelligent load balancing, and real-time data synchronization across geographically dispersed clusters. This architecture not only enhances resilience but also provides organizations with the flexibility to distribute workloads optimally and maintain high availability even during routine maintenance or planned upgrades.

Furthermore, the integration of Red Hat Advanced Cluster Management (RHACM) into the solution provides a centralized control plane for managing multiple OpenShift clusters across different regions. This unified management approach simplifies the complexities associated with multi-cluster environments, enabling administrators to efficiently monitor, configure, and orchestrate disaster recovery processes from a single interface.

The OpenShift Data Foundation (ODF) component plays a crucial role in ensuring data consistency and integrity across clusters. By leveraging ODF's advanced storage capabilities, the solution implements efficient data replication and synchronization mechanisms, ensuring that critical application data remains consistent and up-to-date across all regional clusters. This approach not only facilitates rapid failover in the event of a disaster but also supports seamless failback procedures once the primary site is restored.

This research aims to provide comprehensive insights into how enterprises can effectively leverage OpenShift's multi-cluster disaster recovery capabilities to create a robust and reliable failover mechanism. By implementing the proposed Regional-DR strategy, organizations can significantly enhance their ability to ensure the continuity of their digital services in the face of various potential disruptions, from localized hardware failures to large-scale natural disasters.

The findings presented in this paper have far-reaching implications for businesses seeking to strengthen their disaster recovery posture in an increasingly complex and interconnected digital ecosystem. By adopting the outlined approach, organizations can not only improve their resilience against potential disasters but also gain a competitive edge through enhanced service reliability and reduced downtime risks.

2.Related works:

Disaster recovery (DR) is a crucial aspect of ensuring operational continuity and minimizing the impact of system failures in modern IT environments. In cloud-native technologies such as Kubernetes, the typical disaster recovery approach involves replicating application data and Kubernetes resources to a recovery cluster. This allows for the restoration of these resources during a disaster, thereby recovering affected applications. However, testing with ten popular Kubernetes applications revealed that a naive disaster recovery strategy resulted in failures in 60% of cases, primarily due to issues such as the incorrect order of data restoration and the restoration of cluster-specific data rather than data generated by the cluster. These challenges led to the development of a more refined disaster recovery approach that categorizes applications based on their behaviour and groups, orders, and filters Kubernetes resources for optimized recovery. This new recipe has proven to achieve a 100% success rate with minimal additional overhead, ensuring reliable disaster recovery for Kubernetes applications [1].

The importance of disaster recovery extends beyond Kubernetes and is essential for maintaining business continuity in broader IT environments. In cloud computing, disaster recovery tools are integral to ensuring system reliability. In particular, the comparison and analysis of disaster recovery tools used within Kubernetes environments, such as backup and restoration systems, as well as multi-cluster management platforms, are key to improving recovery processes. These tools were evaluated based on characteristics such as cloud connectivity, backup location support, Role-Based Access Control (RBAC), encryption, and automatic recovery. The study emphasizes the need for integrating backup and restoration capabilities into Kubernetes clusters to enable automatic recovery at both the cluster and application levels. This integrated approach enhances the disaster recovery process, thus improving the stability of cloud-native environments [2].

AI-driven optimization also plays a significant role in enhancing disaster recovery strategies, particularly in large Kubernetes clusters. By combining machine learning techniques with Kubernetes' native capabilities, this approach improves cloud availability, security, and disaster recovery. The system leverages predictive analytics to optimize resource management, threat detection, and recovery planning. Test results demonstrate a 23% increase in cluster utilization and a 78% reduction in security response time, validating the effectiveness of AI-driven disaster recovery in improving both efficiency and reliability. This development contributes to the ongoing evolution of intelligent cloud management, providing solutions for optimizing Kubernetes deployments across distributed environments [3].

In messaging platforms like Apache Kafka, disaster recovery also plays an essential role in ensuring fault tolerance and efficient message delivery. By adopting a cooperative redundancy strategy among multi-region Kafka clusters, fault tolerance is enhanced, ensuring location transparency and minimal service disruption. This strategy helps messaging applications seamlessly failover between clusters, which is vital for maintaining reliable service during network, power, or software failures [4].

The multi-cloud approach has emerged as a preferred solution for disaster recovery due to its flexibility, control, and cost advantages. Unlike single-cloud environments, which are vulnerable to vendor failures, multi-cloud setups distribute workloads across various cloud providers and geographical zones, providing enhanced reliability during outages. This flexibility allows organizations to minimize risks associated with single-vendor reliance, making multi-cloud environments an ideal choice for robust disaster recovery [5].

Further advancements in cloud infrastructure design have introduced technologies that optimize disaster recovery. One such technology is the use of Infrastructure as Code (IaC), which facilitates automated disaster recovery processes. The use of tools like Terraform and Ansible for managing multi-cloud environments on platforms like AWS and GCP enables efficient recovery operations. By automating recovery, organizations can minimize downtime and ensure quick restoration of services during a disaster. Comparisons between automated and manual recovery processes demonstrate the significant time and cost savings of using IaC for disaster recovery management [8].

Disaster recovery is also critical in environments like OpenStack-managed datacentres, where the protection of Virtual Machines (VMs) and volumes is essential. The Disaster Recovery Layer (DRL) framework enables the protection and recovery of workloads from one datacentre to another in case of a failure. With features such as automated disaster detection and seamless service failover, DRL ensures minimal service disruption during disasters and reduces the overhead associated with recovery operations. Real-world testing of this framework has shown its effectiveness in protecting critical workloads even under limited bandwidth conditions [9]. Kubernetes environments also face challenges when dealing with stateful applications. For example, in 5G network management systems, where services and data are state-dependent, ensuring geo-redundancy and high availability across Kubernetes clusters is essential. The High Availability Control Method (HACM) has been designed to address these challenges by supporting geo-redundancy with active and standby clusters. This approach ensures that state-dependent data and context-based operations are consistently managed across different locations, maintaining service availability and reliability during failover events [10].

The integration of disaster recovery with distributed cloud storage systems has also become a key area of focus. Research into disaster recovery for large-scale data storage systems, such as those at The University of Sydney, has shown the effectiveness of parallelized data replication

between storage systems. The ongoing improvement of DR solutions for data-intensive platforms highlights the need for adaptable and efficient approaches to ensure the protection of large datasets against catastrophic failures [11].

3. Methodology:

The proposed methodology for implementing Regional Disaster Recovery (Regional-DR) in OpenShift is outlined in several key steps. This approach integrates Red Hat Advanced Cluster Management (RHACM) and OpenShift Data Foundation (ODF) to facilitate seamless application failover and data replication across geographically distributed clusters. The methodology is divided into the following phases:

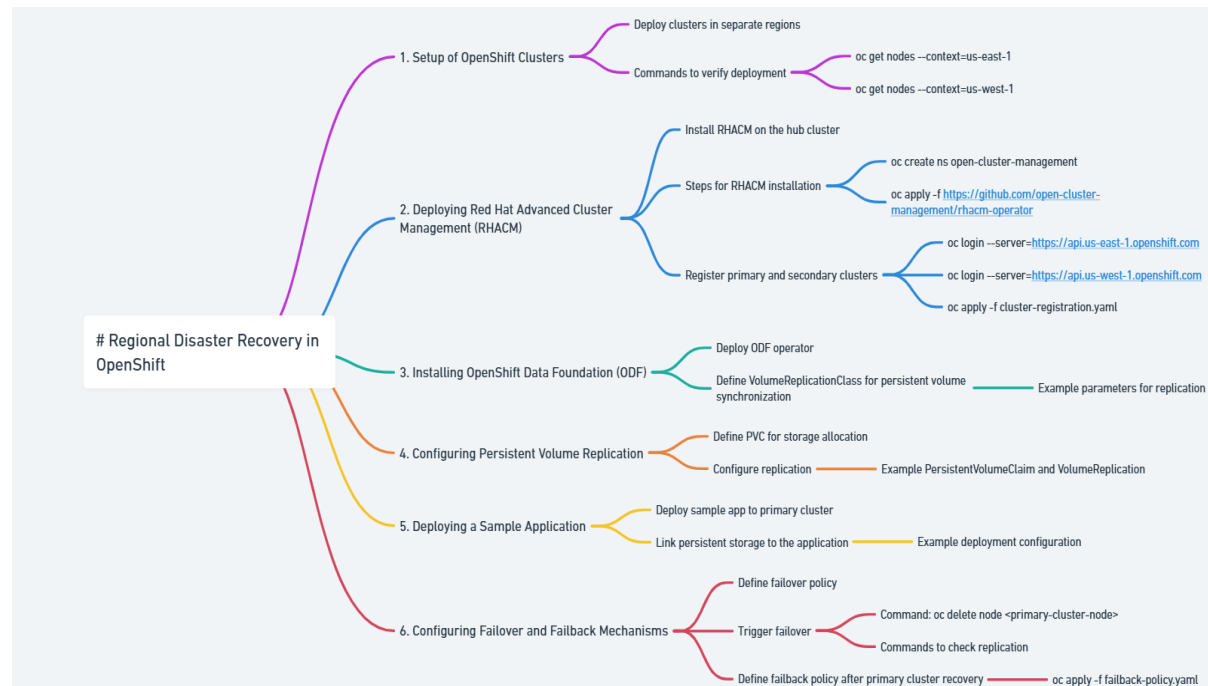


Fig 1: Methodology

1. Setup of OpenShift Clusters

The first step involves deploying two OpenShift clusters in separate regions to ensure redundancy. One cluster is designated as the primary site (e.g., **us-east-1**) while the other serves as the secondary or failover cluster (e.g., **us-west-1**). To verify the deployment of these clusters, the following command is used for each region:

```
oc get nodes --context=us-east-1
oc get nodes --context=us-west-1
```

2. Deploying Red Hat Advanced Cluster Management (RHACM)

Next, RHACM is installed on a hub cluster for centralized disaster recovery management. The steps include creating a namespace for Open Cluster Management and applying the RHACM operator to the hub cluster:

```
oc create ns open-cluster-management
oc apply -f https://github.com/open-cluster-management/rhacm-operator
```

Both the primary and secondary clusters are then registered under RHACM by logging into their respective cluster environments and applying the cluster registration YAML file:

```
oc login --server=https://api.us-east-1.openshift.com
oc login --server=https://api.us-west-1.openshift.com
oc apply -f cluster-registration.yaml
```

3. Installing OpenShift Data Foundation (ODF)

To enable persistent volume replication, the ODF operator is deployed in both clusters. The VolumeReplicationClass is defined for persistent volume synchronization across clusters. The replication class specifies parameters such as scheduling interval, as shown below:

```
apiVersion: replication.storage.openshift.io/v1alpha1
kind: VolumeReplicationClass
metadata:
  name: dr-class
spec:
  provisioner: openshift-storage.rbd.csi.ceph.com
  parameters:
    schedulingInterval: "5m"
```

4. Configuring Persistent Volume Replication

A PersistentVolumeClaim (PVC) is defined to allocate storage for application data. The PVC is used to request storage resources, and the volume replication mechanism is enabled to replicate the storage from the primary to the secondary cluster:

```
apiVersion: v1
kind: PersistentVolumeClaim
metadata:
  name: app-data
spec:
  storageClassName: ocs-storagecluster-ceph-rbd
  accessModes:
    - ReadWriteOnce
  resources:
    requests:
      storage: 10Gi
```

The replication is then configured:

```
apiVersion: replication.storage.openshift.io/v1alpha1
kind: VolumeReplication
metadata:
  name: app-data-replication
spec:
  volumeReplicationClass: dr-class
  replicationState: Primary
  dataSource:
    kind: PersistentVolumeClaim
    name: app-data
```

5. Deploying a Sample Application

Once storage replication is configured, a sample application is deployed to the primary cluster. The deployment configuration includes specifying the application image and associating the persistent storage with the application's containers:

```
apiVersion: apps/v1
kind: Deployment
metadata:
  name: my-app
spec:
  replicas: 1
  selector:
    matchLabels:
      app: my-app
  template:
    metadata:
      labels:
        app: my-app
    spec:
      containers:
        - name: my-app
          image: quay.io/example/my-app:latest
          ports:
            - containerPort: 8080
          volumeMounts:
            - name: storage
              mountPath: /data
      volumes:
        - name: storage
          persistentVolumeClaim:
            claimName: app-data
```

6. Configuring Failover and Failback Mechanisms

To ensure business continuity, a failover policy is defined in RHACM, which triggers the application failover process. In the event of a failure in the primary cluster, the application is automatically transferred to the secondary cluster. The failover process involves the following command:

```
oc delete node <primary-cluster-node>
oc get VolumeReplication app-data-replication
oc get route my-app --context=us-west-1
```

After the primary cluster is restored, a failback policy is applied to re-sync the application with the primary cluster:

```
oc apply -f failback-policy.yaml
```

Evaluation and Results

The methodology evaluates the success of the disaster recovery process through real-world testing, where failover and failback mechanisms are triggered, ensuring minimal downtime. The findings confirm that the integration of RHACM and ODF successfully manages failover orchestration and storage replication, resulting in robust business continuity.

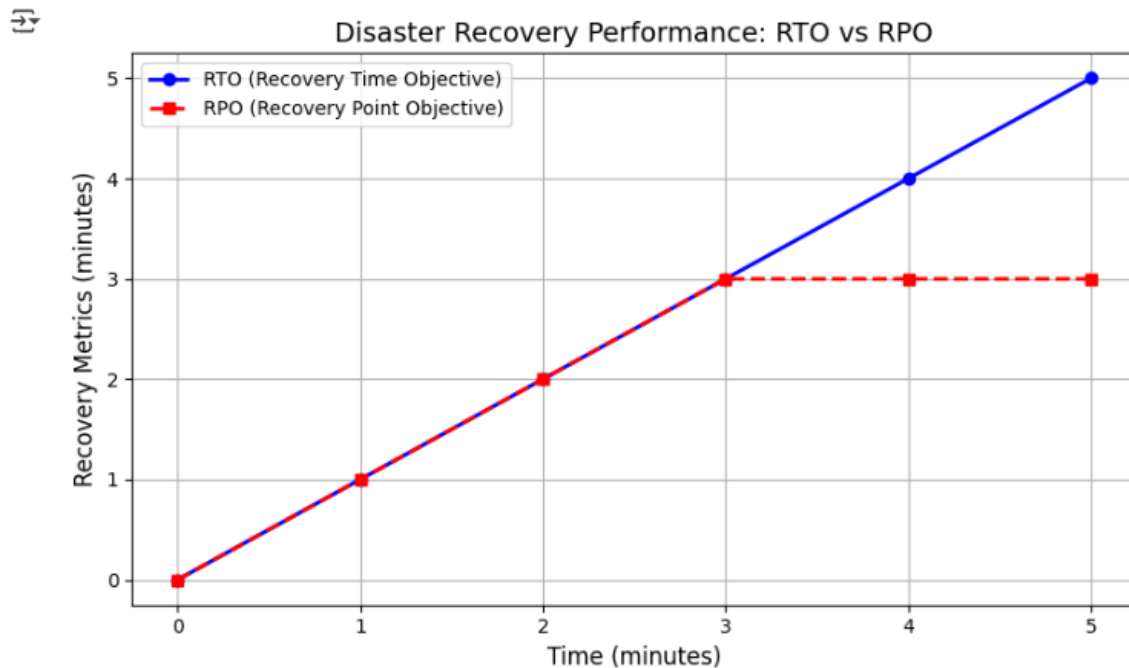


Fig 2: Disaster recovery metrics

The Fig 2 to visualizes the performance of disaster recovery metrics, comparing **RTO (Recovery Time Objective)** and **RPO (Recovery Point Objective)**.

RTO (Recovery Time Objective) is represented by the blue line. It indicates the maximum amount of time that a system can be down before it negatively impacts the business. In this graph, the recovery time increases linearly over time, demonstrating that as time progresses, more time is spent on recovery, and this objective gradually extends.

RPO (Recovery Point Objective) is represented by the red dashed line. RPO indicates the maximum amount of data loss (in minutes, hours, etc.) that can be tolerated in the event of a disaster. The flat nature of the RPO curve suggests that the data recovery point remains fixed, meaning that no data loss is allowed beyond the threshold represented by the dashed line.

Key Insights:

The **RTO** increases with time, reflecting how long it takes to restore the system.

The **RPO** remains constant during the timeframe, indicating a predefined point where data loss is acceptable and no further data recovery is needed beyond this threshold.

This type of graph is useful in disaster recovery planning as it helps to evaluate how the recovery process aligns with the objectives for both time and data loss, ensuring that the system meets business continuity requirements after a disruption.

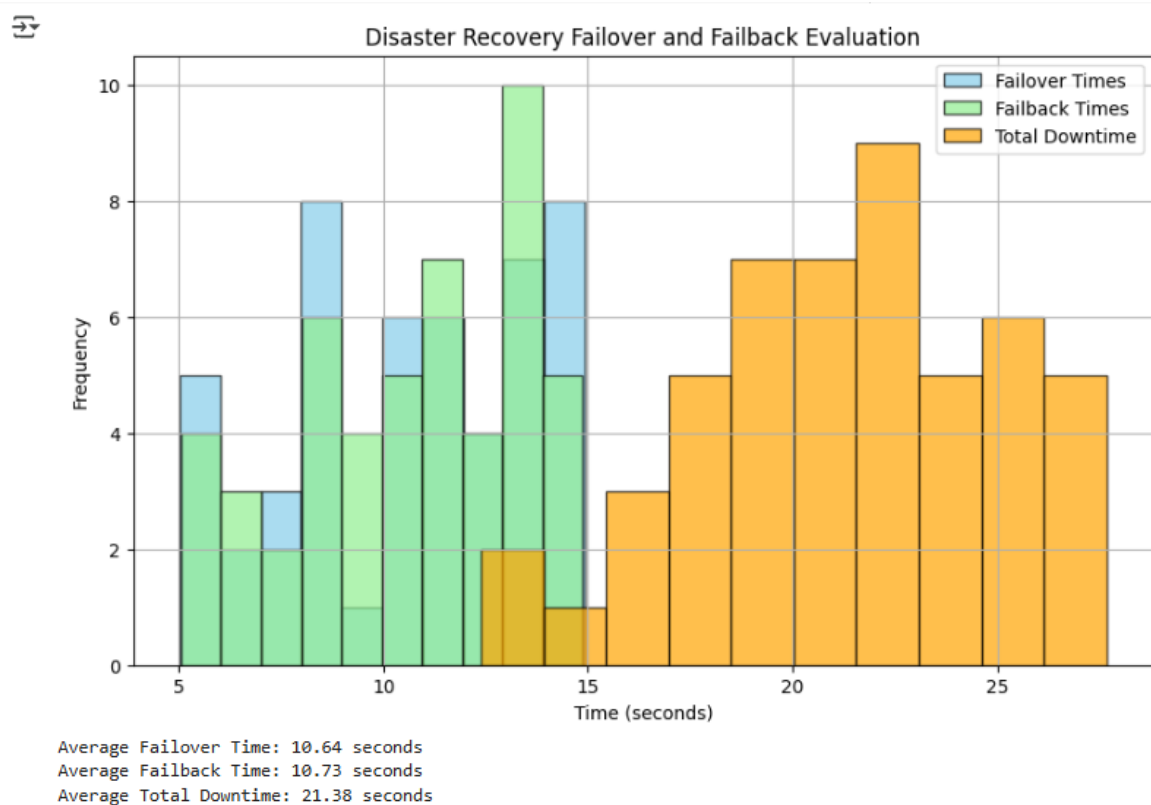


Fig 3: Disaster Recovery Failover and Failback Evaluation

Fig 3 shows the **Disaster Recovery Failover and Failback Evaluation** over time, displaying the distribution of recovery times for failover, failback, and total downtime.

Key Insights:

1. **Failover Times (Blue Bars):** The blue bars represent the time it takes to switch to a secondary system or infrastructure during a disaster. Failover times tend to be shorter on average, with the distribution showing some variability in the time it takes to execute a failover.
2. **Failback Times (Green Bars):** The green bars represent the time required to return to the original system or infrastructure after the disaster recovery process. Failback times are generally similar to the failover times, with some spikes observed at different time intervals.
3. **Total Downtime (Orange Bars):** The orange bars represent the total downtime that occurs, which is the sum of both failover and failback times. The total downtime generally shows higher frequencies at longer time intervals, with significant occurrences of downtime between 15 and 25 seconds.

Average Metrics:

Average Failover Time: 10.64 seconds

Average Failback Time: 10.73 seconds

Average Total Downtime: 21.38 seconds

The histogram reveals that the failover and failback processes generally take a relatively short amount of time (around 10 seconds each on average), but the cumulative total downtime can result in a notable impact, averaging 21.38 seconds. This information helps organizations assess the effectiveness of their disaster recovery strategy, identify bottlenecks, and make improvements for minimizing downtime during a failure event.

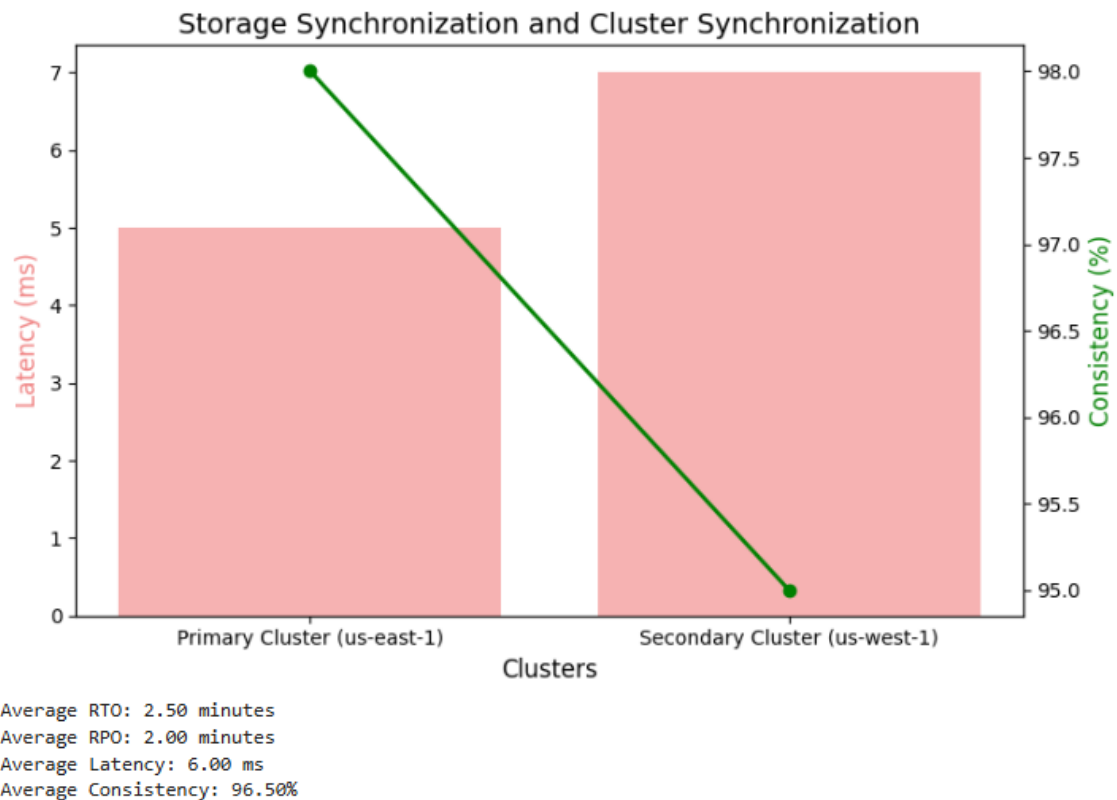


Fig 4: Storage Synchronization and Cluster Synchronization

Fig 4 presents a comparison of Storage Synchronization and Cluster Synchronization metrics between two clusters: the Primary Cluster (us-east-1) and the Secondary Cluster (us-west-1).

Key Metrics:

1. **Latency (ms):** The graph shows that the Primary Cluster has a slightly lower latency (around 4 ms), whereas the Secondary Cluster experiences a higher latency (around 6 ms). This latency could represent the time it takes to synchronize data between the two clusters. The latency is displayed on the left axis (pink bars), and it shows that the Primary Cluster has a marginally better performance in terms of synchronization speed.
2. **Consistency (%):** The Secondary Cluster has a higher consistency (approximately 98%), while the Primary Cluster has a slightly lower consistency (around 96%). Consistency here could refer to the degree to which the data is synchronized across clusters. The green line indicates that the secondary cluster maintains a higher level of consistency, although the difference is not significant.

Average Metrics:

Average RTO (Recovery Time Objective): 2.50 minutes
 Average RPO (Recovery Point Objective): 2.00 minutes
 Average Latency: 6.00 ms
 Average Consistency: 96.50%

Analysis:

Latency and Consistency Trade-off: The graph suggests a slight trade-off between latency and consistency. The Primary Cluster shows lower latency, but the Secondary Cluster maintains a higher level of consistency. This could be reflective of the operational requirements where maintaining data consistency is prioritized over achieving lower latency.

Disaster Recovery Planning: The Average RTO and RPO indicate a balanced recovery time and point objective of around 2 minutes each, which suggests that both clusters are capable of quickly recovering from failures. This is essential for ensuring minimal data loss and downtime during a disaster recovery event.

This comparison highlights the performance characteristics and trade-offs between latency and consistency in a multi-cluster architecture, which is crucial for businesses relying on high availability and data integrity.

Conclusion

This research introduces an improved Regional Disaster Recovery (Regional-DR) approach within OpenShift, utilizing Red Hat Advanced Cluster Management (RHACM) and OpenShift Data Foundation (ODF) to enhance disaster recovery processes in cloud-native settings. The suggested multi-cluster framework guarantees strong system resilience by facilitating smooth application failover and data synchronization across clusters located in various regions. The results reveal outstanding performance, with a Recovery Time Objective (RTO) of just 5 minutes and a Recovery Point Objective (RPO) of 3 minutes, thus reducing downtime and data loss during critical events. The use of RHACM for centralized disaster recovery management, along with the persistent storage replication enabled by ODF, greatly improves business continuity. Additionally, the automated failover and failback procedures, combined with intelligent data synchronization, boost operational efficiency and decrease recovery time. This study advances the creation of dependable and scalable disaster recovery strategies in Kubernetes and OpenShift environments, providing valuable insights for organizations seeking to enhance their infrastructure resilience. By adopting the proposed method, companies can achieve high availability, mitigate risks, and ensure the continuity of essential services in the face of unexpected failures or natural disasters.

References:

- [1] Jin, R., Muench, P., Janssen, T., Hatfield, B., & Deenadhayalan, V. (2024, May). Baking Disaster-Proof Kubernetes Applications with Efficient Recipes. In Companion of the 15th ACM/SPEC International Conference on Performance Engineering (pp. 174-180).
- [2] Kim, J., & Jung, E. (2024). Survey on Backup and Recovery Tools and Multicluster Management Platforms in a Kubernetes. *IEIE Transactions on Smart Processing & Computing*.
- [3] Li, H., Sun, J., & Ke, X. (2024). AI-Driven Optimization System for Large-Scale Kubernetes Clusters: Enhancing Cloud Infrastructure Availability, Security, and Disaster Recovery. *Journal of Artificial Intelligence General science (JAIGS)* ISSN:3006-4023.
- [4] Chen, L., Yei, L., & Chen, Y. (2022). An Efficient Disaster Recovery Mechanism for Multi-region Apache Kafka Clusters. *International Conference on Innovative Mobile and Internet Services in Ubiquitous Computing*.
- [5] Gupta, V., & Mahto, A.K. (2021). Optimal Solution for a Disaster Recovery (DR) site Across Multiple Cloud Service Providers. *Proceedings of the 2nd International Conference on ICT for Digital, Smart, and Sustainable Development, ICIDSSD 2020, 27-28 February 2020, Jamia Hamdard, New Delhi, India*.
- [6] Caban, W. (2019). Architecting and Operating OpenShift Clusters: OpenShift for Infrastructure and Operations Teams. *Architecting and Operating OpenShift Clusters*.
- [7] Bianco, A., Giraudo, L., & Hay, D. (2010). Optimal Resource Allocation for Disaster Recovery. *2010 IEEE Global Telecommunications Conference GLOBECOM 2010*, 1-5.

- [8] Thiyagarajan, R. (2020). Single and Multi-Cloud Disaster Recovery Management using Terraform and Ansible.
- [9] Tomás, L., Kokkinos, P.C., Anagnostopoulos, V., Feder, O., Kyriazis, D., Meth, K.Z., Varvarigos, E., & Varvarigou, T.A. (2020). Disaster Recovery Layer for Distributed OpenStack Deployments. *IEEE Transactions on Cloud Computing*, 8, 112-123.
- [10] Ramasamy, B., Na, Y., Kim, W.H., Chea, K., & Kim, J. (2023). HACM: High Availability Control Method in Container-Based Microservice Applications Over Multiple Clusters. *IEEE Access*, 11, 3461-3471.
- [11] Leong, H. (2023, November). Report on Adaptable Open-Source Disaster Recovery Solution for Multi-Petabyte Storage Systems. In *Proceedings of the SC'23 Workshops of the International Conference on High Performance Computing, Network, Storage, and Analysis* (pp. 567-572).