

# Accelerating IBM watsonx.data with IBM Fusion HCI

Larry Coyne

Paulina Acevedo

Eduardo Daniel Ibarra Alvarez

Gabriela Valencia Castillo

Kenneth Hartsoe

Mike Kieran

Alberto Larios

Savitha H N

Khanh Ngo

AshaRani G R

Shyamala Rajagopalan

Ben Randall

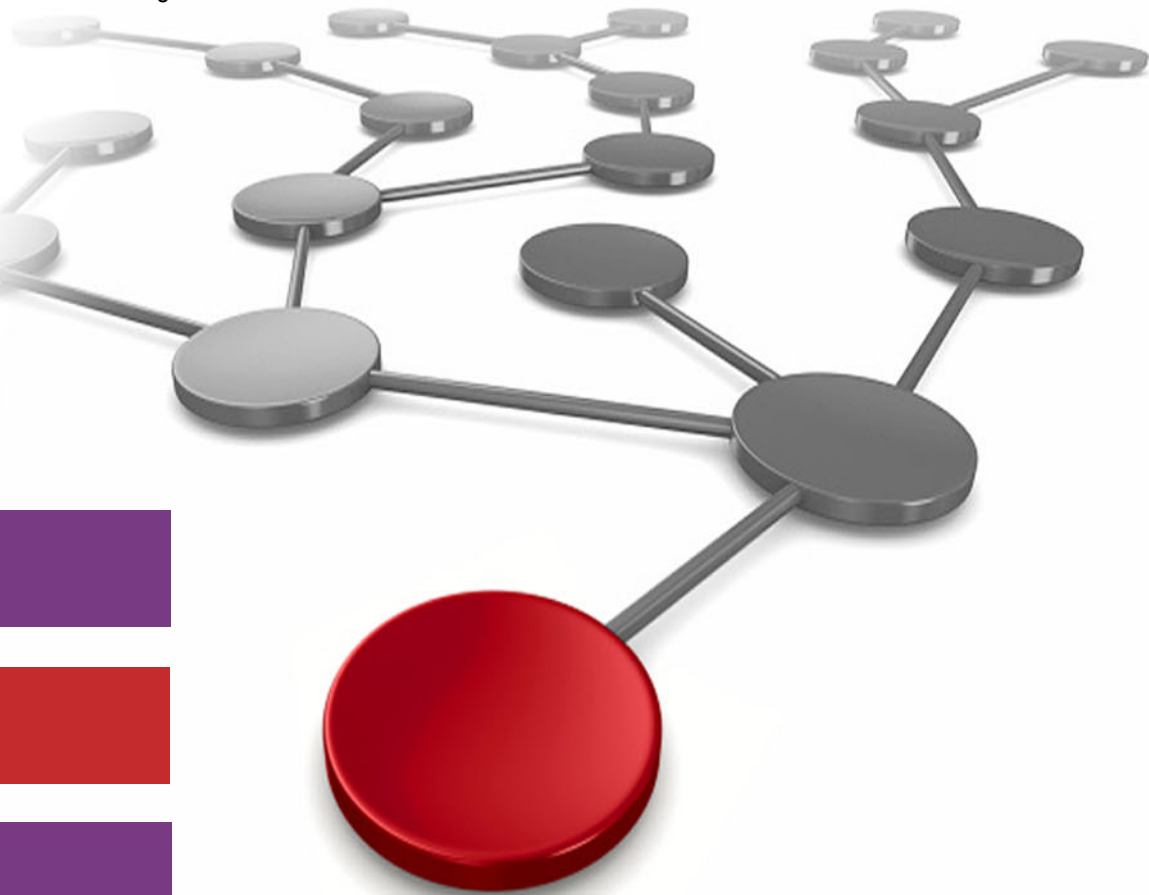
Hemalatha B T

Todd Tosseth

Jayson Tsingine

Israel Vizcarra

Zhi Yong Xue



**Data and AI**

 **Cloud**

**Storage**





IBM Redbooks

## **Accelerating IBM watsonx.data with IBM Fusion HCI**

March 2024

**Note:** Before using this information and the product it supports, read the information in “Notices” on page v.

**First Edition (March 2024)**

This edition applies to Version 2, Release 7, Modification x of IBM Fusion HCI

**© Copyright International Business Machines Corporation 2024. All rights reserved.**

Note to U.S. Government Users Restricted Rights -- Use, duplication or disclosure restricted by GSA ADP Schedule Contract with IBM Corp.

# Contents

|                                                                                                      |     |
|------------------------------------------------------------------------------------------------------|-----|
| <b>Notices</b> .....                                                                                 | v   |
| Trademarks .....                                                                                     | vi  |
| <b>Preface</b> .....                                                                                 | vii |
| Authors .....                                                                                        | vii |
| Now you can become a published author, too! .....                                                    | ix  |
| Comments welcome .....                                                                               | ix  |
| Stay connected to IBM Redbooks .....                                                                 | x   |
| <b>Chapter 1. Solution overview</b> .....                                                            | 1   |
| 1.1 Overview .....                                                                                   | 2   |
| 1.2 Benefits .....                                                                                   | 4   |
| 1.3 Architecture, components, and functional characteristics .....                                   | 5   |
| 1.3.1 Integrated solution architecture .....                                                         | 5   |
| 1.3.2 Solution component architectures .....                                                         | 6   |
| <b>Chapter 2. Sizing and planning</b> .....                                                          | 11  |
| 2.1 Sizing guidelines .....                                                                          | 12  |
| 2.1.1 Licensing .....                                                                                | 12  |
| 2.1.2 IBM Fusion HCI with IBM watsonx.data .....                                                     | 12  |
| 2.2 Planning .....                                                                                   | 15  |
| 2.2.1 Network access to object storage .....                                                         | 15  |
| 2.2.2 AFM .....                                                                                      | 15  |
| 2.3 Customer use cases .....                                                                         | 16  |
| 2.3.1 Data sharing .....                                                                             | 16  |
| <b>Chapter 3. Implementation</b> .....                                                               | 17  |
| 3.1 IBM Fusion HCI installation .....                                                                | 18  |
| 3.2 Installing Red Hat OpenShift Data Foundation and configuring the Multicloud Object Gateway ..... | 18  |
| 3.2.1 Configuring Advanced File Management nodes .....                                               | 19  |
| 3.2.2 Creating the static PV and PVC .....                                                           | 21  |
| 3.2.3 Performance tuning .....                                                                       | 27  |
| 3.3 Installing IBM watsonx.data on IBM Fusion HCI .....                                              | 29  |
| <b>Chapter 4. Monitoring</b> .....                                                                   | 31  |
| <b>Chapter 5. Backup and restore of IBM Cloud Pak for Data</b> .....                                 | 33  |
| 5.1 Considerations and requirements .....                                                            | 34  |
| 5.2 Getting the prerequisites ready .....                                                            | 34  |
| 5.2.1 Installing the Cloud Pak Backup and Restore service .....                                      | 34  |
| 5.2.2 Installing the cpdbr service on the target cluster .....                                       | 35  |
| 5.2.3 Backup policies for Cloud Pak for Data applications .....                                      | 36  |
| 5.2.4 Creating and assigning a backup policy .....                                                   | 37  |
| 5.3 Backing up the source cluster .....                                                              | 39  |
| 5.4 Restoring to an alternative cluster .....                                                        | 39  |
| <b>Related publications</b> .....                                                                    | 45  |
| IBM Redbooks .....                                                                                   | 45  |
| Other publications .....                                                                             | 45  |

|                        |    |
|------------------------|----|
| Online resources ..... | 45 |
| Help from IBM .....    | 46 |

# Notices

This information was developed for products and services offered in the US. This material might be available from IBM in other languages. However, you may be required to own a copy of the product or product version in that language in order to access it.

IBM may not offer the products, services, or features discussed in this document in other countries. Consult your local IBM representative for information on the products and services currently available in your area. Any reference to an IBM product, program, or service is not intended to state or imply that only that IBM product, program, or service may be used. Any functionally equivalent product, program, or service that does not infringe any IBM intellectual property right may be used instead. However, it is the user's responsibility to evaluate and verify the operation of any non-IBM product, program, or service.

IBM may have patents or pending patent applications covering subject matter described in this document. The furnishing of this document does not grant you any license to these patents. You can send license inquiries, in writing, to:

*IBM Director of Licensing, IBM Corporation, North Castle Drive, MD-NC119, Armonk, NY 10504-1785, US*

INTERNATIONAL BUSINESS MACHINES CORPORATION PROVIDES THIS PUBLICATION "AS IS" WITHOUT WARRANTY OF ANY KIND, EITHER EXPRESS OR IMPLIED, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF NON-INFRINGEMENT, MERCHANTABILITY OR FITNESS FOR A PARTICULAR PURPOSE. Some jurisdictions do not allow disclaimer of express or implied warranties in certain transactions, therefore, this statement may not apply to you.

This information could include technical inaccuracies or typographical errors. Changes are periodically made to the information herein; these changes will be incorporated in new editions of the publication. IBM may make improvements and/or changes in the product(s) and/or the program(s) described in this publication at any time without notice.

Any references in this information to non-IBM websites are provided for convenience only and do not in any manner serve as an endorsement of those websites. The materials at those websites are not part of the materials for this IBM product and use of those websites is at your own risk.

IBM may use or distribute any of the information you provide in any way it believes appropriate without incurring any obligation to you.

The performance data and client examples cited are presented for illustrative purposes only. Actual performance results may vary depending on specific configurations and operating conditions.

Information concerning non-IBM products was obtained from the suppliers of those products, their published announcements or other publicly available sources. IBM has not tested those products and cannot confirm the accuracy of performance, compatibility or any other claims related to non-IBM products. Questions on the capabilities of non-IBM products should be addressed to the suppliers of those products.

Statements regarding IBM's future direction or intent are subject to change or withdrawal without notice, and represent goals and objectives only.

This information contains examples of data and reports used in daily business operations. To illustrate them as completely as possible, the examples include the names of individuals, companies, brands, and products. All of these names are fictitious and any similarity to actual people or business enterprises is entirely coincidental.

## COPYRIGHT LICENSE:

This information contains sample application programs in source language, which illustrate programming techniques on various operating platforms. You may copy, modify, and distribute these sample programs in any form without payment to IBM, for the purposes of developing, using, marketing or distributing application programs conforming to the application programming interface for the operating platform for which the sample programs are written. These examples have not been thoroughly tested under all conditions. IBM, therefore, cannot guarantee or imply reliability, serviceability, or function of these programs. The sample programs are provided "AS IS", without warranty of any kind. IBM shall not be liable for any damages arising out of your use of the sample programs.

# Trademarks

IBM, the IBM logo, and ibm.com are trademarks or registered trademarks of International Business Machines Corporation, registered in many jurisdictions worldwide. Other product and service names might be trademarks of IBM or other companies. A current list of IBM trademarks is available on the web at “Copyright and trademark information” at <https://www.ibm.com/legal/copytrade.shtml>

The following terms are trademarks or registered trademarks of International Business Machines Corporation, and might also be trademarks or registered trademarks in other countries.

Db2®  
DS8000®  
Enterprise Storage Server®

IBM®  
IBM Cloud®  
IBM Cloud Pak®

Redbooks®  
Redbooks (logo) ®  
XIV®

The following terms are trademarks of other companies:

Red Hat, Ceph, OpenShift, are trademarks or registered trademarks of Red Hat, Inc. or its subsidiaries in the United States and other countries.

Other company, product, or service names may be trademarks or service marks of others.

# Preface

Organizations that are expanding from AI pilot projects to full-scale production systems typically need a set of tools for building and deploying foundation models, a container-based application platform, software-defined storage, and hardware on which to run it all. This IBM Redpaper publication describes the IBM® solution for running IBM watsonx.data on premises, with IBM Fusion HCI providing an appliance-based hosting platform, and IBM Storage Ceph providing cloud-scale object storage.

This publication shows how to set up the Storage Acceleration feature, so IBM watsonx.data queries can benefit from a shareable on-premises high-performance cache acceleration. The Storage Acceleration feature is available only on an IBM Fusion HCI.

This paper is targeted toward technical professionals, consultants, technical support staff, IT Architects, and IT specialists who are responsible for delivering data lakehouse solutions optimized for data, analytics, and AI workloads.

## Authors

This paper was produced by a team of specialists from around the world working with the IBM Redbooks, Tucson Center.

**Larry Coyne** is a Project Leader at the IBM International Technical Support Organization, Tucson, Arizona, center. He has over 35 years of IBM experience, with 23 years in IBM storage software management. He holds degrees in Software Engineering from the University of Texas at El Paso and Project Management from George Washington University. His areas of expertise include client relationship management, quality assurance, development management, and support management for IBM storage management software.

**Paulina Acevedo** is a System Test Architect for the Application and Resiliency Fusion Team. Paulina has been with IBM for more than 17 years and has held several different positions within the Systems organization. She is a certified Project Manager and has been the System Test Project manager for several IBM Storage products.

**Eduardo Daniel Ibarra Alvarez** is a Systems Test Engineer for the Application and Resiliency Fusion Team. Eduardo has a strong automation and testing background.

**Gabriela Valencia Castillo** is a System Test Engineer for the Application and Resiliency Fusion Team. Gabriela has extensive experience in functional and non-functional testing, as well as test plan creation, design scenarios, and creation of test cases.

**Kenneth Hartsoe** is the documentation Content Strategist for IBM Storage Ceph and IBM Fusion Data Foundation, as well as providing content strategy collaboration across multiple solutions within IBM Storage. Kenneth has more than twenty years experience within the storage documentation area, both as a senior technical writer and content strategist, including several years as the storage content strategist at Red Hat.

**Mike Kieran** is a product marketer currently focusing on Storage for IBM watsonx.data. He has more than a decade of storage marketing experience at Pure Storage, NetApp, and Nimble Storage. Mike is the author of four books on digital color and the recipient of an Emmy for science and technology writing. He is an avid stargazer and astrophotographer.

**Alberto Larios** is a System Test Engineer for the Application and Resiliency Fusion Team. Alberto collaborates with his team as a systems associate and tests Cloud Pak products on IBM Storage. His expertise lies on the data science field with knowledge on machine learning and modelling.

**Savitha H N** is a System Test Engineer for the Application and Resiliency Fusion Team. Savitha has been with IBM from past one year and is a systems associate with Cloud Pak. She holds a Devops engineer certification, and specialist in Q/A testing by executing scenarios to Ensures the product is robust and failure scenarios are considered and refactored.

**Khanh Ngo** is a leader in the IBM Storage CTO office specializing in Data and AI integration with IBM Storage products including the optimization of IBM watsonx.data.

**AshaRani G R** is a System Test Engineer for the Application and Resiliency Fusion Team. Asha possesses extensive expertise encompassing server hardware, management software, and storage solutions within SAN and DAS domains. She has expertise in bringing up compute nodes for Infrastructure as a Service and performing end-to-end system testing. Additionally, she has a strong background in virtualization technology.

**Shyamala Rajagopalan** is the senior lead technical content and information architect for IBM Fusion product. She has over 24 years of industry experience, with a significant track record of successfully delivering end-to-end IT documentation solutions to global clients across diverse industries. Her area of contribution includes digital transformation, legacy modernization, integration platforms (middleware, EAI), semiconductor, Cloud, and storage systems. She has performed diverse roles in the information development domain, which includes Information Architect, Technical Writing Manager, Lead Content Developer, Senior technical writer, and Competency Area Mentor.

**Ben Randall** is the User Experience Architect for IBM Fusion and works on product development and design. He has worked in the enterprise storage industry for 21 years, focusing on technologies such as disaster recovery, backup and restore, container native storage, software defined storage, high performance computing, and SAN monitoring.

**Hemalatha B T** is a System Test lead for Application and Resiliency Fusion Team. Hema has been with IBM for over 15 years and was associated with Power System Performance before moving to the functional/system test area working for products like IBM Cloud Pak® for Systems, IBM Fusion. An enthusiastic quality assurance person who thrives to ensure high quality and perfectly functioning systems are delivered to customers.

**Todd Tosseth** is a Software Engineer for IBM in Tucson, Arizona. Joining IBM in 2001, he has worked as a test and development engineer on several IBM storage products, such as IBM DS8000®, IBM Storage Scale, and IBM Enterprise Storage Server®. He is working on IBM Cloud Pak as a system test engineer, with an emphasis on Cloud Pak storage integration.

**Jayson Tsingine** is an Advisory Software Engineer in the IBM Systems Storage Group based in Tucson, Arizona. Jayson has worked in numerous test roles since he joined IBM in 2003, providing test support for storage products, including IBM XIV®, DS8000, FlashSystems, Hydra, Spectrum Scale, Spectrum NAS, and Cloud Pak Solutions. He holds a BS degree in Computer Science from the University of Arizona.

**Israel Vizcarra** is a test specialist that works for the Cloud Pak Storage Test Team. He graduated as a Mechatronics Engineer and he has 8 years of experience in the Quality Assurance area for the storage organization. Israel has actively participated in different roles from Functional Verification, System level, and Automation Testing.

**Zhi Yong Xue** is an Architect of Fusion Storage in China. He has 15 years of experience in software design and development as a developer and architect at IBM. He holds a Bachelor's degree in Exploration Technology and Engineering from the University of Petroleum. His areas of expertise include Storage and Cloud computing.

Thanks to the following people for their contributions to this project:

Patrik Hysky

**IBM Redbooks®, IBM Infrastructure**

Pramod Thekkepat Achutha, Tara Astigarraga, Scott Colbeck, Karli Collins, Marcel Hergaarden, Ted Hoover, Lisa Huston, Kedar Karmarkar, Moushumi Kalita, Joshua Kim, Ana Lilia Sanchez Llamas, Ajay Lunawat, Boda Devi Manikanta, Kevin Shen, Bill Stoddard, Chris Tan, Thiha Than, Garth Tschetter, Hans Uhlig

**IBM Infrastructure**

David Wohlford

**IBM CHQ, Marketing**

## Now you can become a published author, too!

Here's an opportunity to spotlight your skills, grow your career, and become a published author—all at the same time! Join an IBM Redbooks residency project and help write a book in your area of expertise, while honing your experience using leading-edge technologies. Your efforts will help to increase product acceptance and customer satisfaction, as you expand your network of technical contacts and relationships. Residencies run from two to six weeks in length, and you can participate either in person or as a remote resident working from your home base.

Find out more about the residency program, browse the residency index, and apply online at:

[ibm.com/redbooks/residencies.html](http://ibm.com/redbooks/residencies.html)

## Comments welcome

Your comments are important to us!

We want our papers to be as helpful as possible. Send us your comments about this paper or other IBM Redbooks publications in one of the following ways:

- Use the online **Contact us** review Redbooks form found at:

[ibm.com/redbooks](http://ibm.com/redbooks)

- Send your comments in an email to:

[redbooks@us.ibm.com](mailto:redbooks@us.ibm.com)

- Mail your comments to:

IBM Corporation, IBM Redbooks  
Dept. HYTD Mail Station P099  
2455 South Road  
Poughkeepsie, NY 12601-5400

## Stay connected to IBM Redbooks

- ▶ Find us on LinkedIn:  
<https://www.linkedin.com/groups/2130806>
- ▶ Explore new Redbooks publications, residencies, and workshops with the IBM Redbooks weekly newsletter:  
<https://www.redbooks.ibm.com/subscribe>
- ▶ Stay current on recent Redbooks publications with RSS Feeds:  
<https://www.redbooks.ibm.com/rss.html>



## Solution overview

Organizations that are expanding from AI pilot projects to full-scale production systems typically need the following components:

- ▶ A set of tools for building and deploying foundation models
- ▶ A container-based application platform
- ▶ Software-defined storage
- ▶ Hardware on which to run it

This publication describes the IBM solution for running IBM watsonx.data on premises, with IBM Fusion HCI providing an appliance-based hosting platform, and IBM Storage Ceph providing cloud-scale object storage. This publication shows how to set up the Storage Acceleration feature, which is only available on IBM Fusion HCI, so IBM watsonx.data queries can benefit from a shareable on-premises, high-performance cache acceleration.

This paper is targeted toward technical professionals including consultants, technical support staff, IT Architects, and IT specialists who are responsible for delivering optimized for data, analytics, and AI workloads.

This chapter includes an overview covering the background of data lakes and how the IBM solution of IBM watsonx.data, IBM Storage Ceph, and IBM Fusion HCI accelerated infrastructure works to improve on-premises performance and improves cost efficiency. The architecture of the solution and components are also described.

## 1.1 Overview

This section describes the evolution of data lakes, the emergence of data lakehouses, and IBM watsonx.data lakehouse, IBM Storage Ceph, and the IBM Fusion HCI accelerated infrastructure solution.

### From data warehouse to data lake

During the past 20 years, large organizations have changed the way they aggregate data for analytics and business intelligence (BI) purposes. The original approach was to build a single monolithic database, or data warehouse, and then analyze specific subsets of the data through an extract, transform, load (ETL) process based on queries by using structured query language (SQL).

Data warehouses are often used for repeatable reporting and analysis workloads such as monthly sales reports, tracking of sales per region, and website traffic. But building and maintaining a data warehouse is a costly, time-consuming process, and data warehouses work only with structured data.

Moving data warehouses to the cloud doesn't solve the problem. Sometimes, it makes them even more expensive, and they're still not well suited to machine learning or AI applications.

These limitations led to the concept of the data lake, which is a centralized repository that can store massive volumes of data in its original form so that it's consolidated, integrated, secure, and accessible. Data lakes are designed to accommodate all types of data from many different sources:

- ▶ Structured data, such as database tables and Excel sheets
- ▶ Semi-structured data, such as herbages and XML files
- ▶ Unstructured data, such as images, video, audio, and social media posts

Because data lakes are massively scalable and can handle all types of data, they are ideal for real-time analytics, predictive analytics, and machine learning or AI. They are also typically less costly than data warehouses.

### Data lakehouse architecture

The data lakehouse is an emerging architecture that offers the flexibility of a data lake with the performance and structure of a data warehouse. Lakehouse solutions typically provide a high-performance query engine over low-cost object storage along with a metadata governance layer. Data lakehouses are based around open-standard object storage and enable multiple analytics and AI workloads to operate simultaneously on top of the data lake without requiring that the data be duplicated and converted.

A key benefit of data lakehouses is that they address the needs of both traditional data warehouse analysts who curate and publish data for business intelligence and reporting purposes; and of data scientists and engineers who run more complex data analysis and processing workloads.

IBM watsonx.data, shown in Figure 1-1, is built on an open lakehouse architecture, supported by querying, governance, and open data formats for accessing and sharing data.

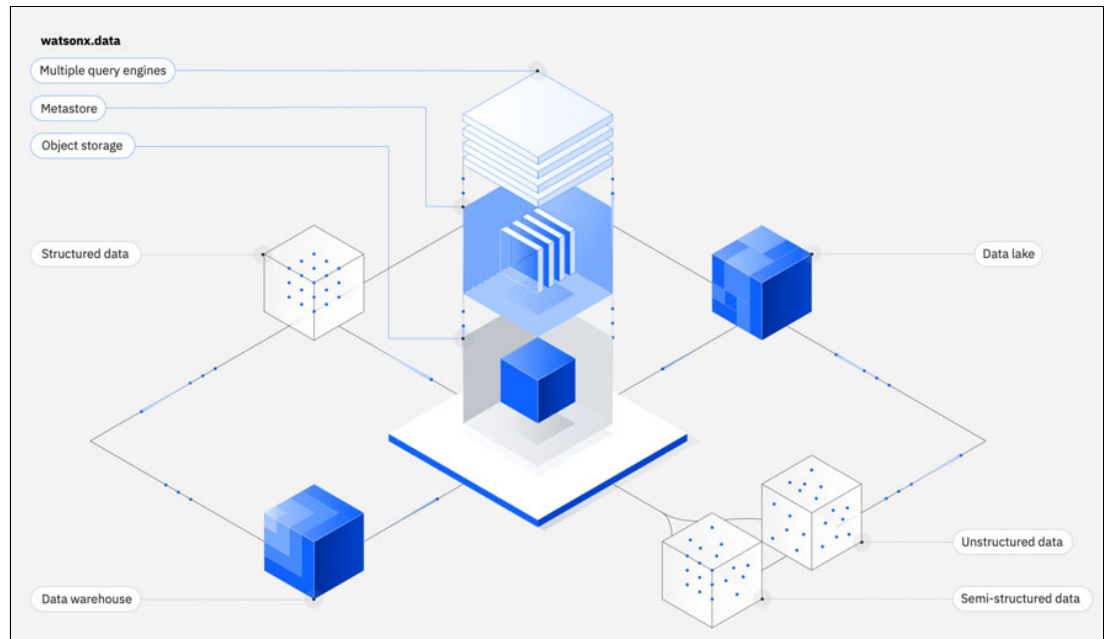


Figure 1-1 IBM watsonx.data provides an ideal platform for building and scaling AI applications

## IBM watsonx.data, IBM Storage Ceph, and the IBM Fusion HCI accelerated infrastructure solution

Administrators of today's modern data lakehouses are required to think about storage optimizations as a top priority and a two-tiered approach. The first tier is an on-premises high-performance acceleration layer, which provides superior storage bandwidth with a cost-effective caching approach for the hybrid cloud object storage. The second tier is the low-cost persistent storage for your on-premises storage needs. With the combination of IBM Fusion HCI as your first tier solution and IBM Storage Ceph as your second tier solution, an organization can improve query performance with Storage Acceleration, significant cost advantage, and superior data management capabilities. IBM watsonx.data can take advantage of both of these tiers when using the IBM Fusion HCI and IBM Storage Ceph.

## 1.2 Benefits

IBM Fusion HCI is a hosting platform for IBM watsonx.data and provides the following benefits and features:

- ▶ Hosting platform for IBM watsonx products, starting with IBM watsonx.data:
  - Provides an automated deployment of Red Hat OpenShift on top of resilient compute, network, and storage in an appliance form-factor.
  - Provides all the storage classes that are needed by IBM Cloud Pak for Data (CP4D) and IBM watsonx.data
- ▶ Storage acceleration feature for Tier 1 data caching to accelerate IBM watsonx.data query performance to 5–15x improvement:
  - Connects to multiple object buckets
  - Uses intelligent caching to accelerate data access including automatic eviction
  - High-performance persistent object cache with low-capacity requirements:
    - Cache once concept for faster performance
    - Shareable across all engines and projects and namespaces
    - Cache available to all nodes
    - Multi-protocol (including virtualization) support
    - Supports IBM Cloud Object Storage, Amazon Web Services, Seagate Lyve Cloud, Google Cloud Platform
- ▶ watsonx platform in a box:
  - Install efforts of a few days
  - Support for a maximum of 2 dedicated GPU nodes with optimizations
  - Support for a maximum of 2 dedicated gateway nodes for data access services
  - Scalable by adding nodes and disk capacity
- ▶ Shared run-time platform:
  - Multiple solutions in a box:
    - IBM Db2® Warehouse
    - watsonx.data
  - Shared resources across multiple engines:
    - Presto
    - Spark
  - Compute-storage nodes provide high core-to-memory ratio. A C05 node with 64 cores and 2 TB memory yields a 1:32 core-to-memory ratio.
- ▶ Global Data Platform:
  - Data access services provides better performance across multiple parallel paths with single source of truth.
  - Data virtualization, collaboration and orchestration services for a true global namespace and data sharing across geo-distributed locations.

- Supports compression at storage class level for space savings for various open data formats.
- Encryption ensures both secure storage and secure deletion of data (at file system level).
- ▶ Local S3 object storage
- ▶ IBM Storage Ceph as an external cloud-scale S3 object store
- ▶ Ability to integrate GPUs into the IBM watsonx solution

It is worth noting that the Storage Acceleration feature providing the data caching for improved query performance is very different from your traditional local caching. The IBM Fusion HCI has a global data platform which allows for a cache only once concept to achieve faster performance and transparency. After an object has been cached, it is available and shareable to every engine with IBM watsonx.data across all nodes within the cluster. The Storage Acceleration provides a persistent data cache for all engines. Newly provisioned engines also begin with a warm or hot cache.

## 1.3 Architecture, components, and functional characteristics

This section provides an architecture overview of IBM watsonx.data with IBM Fusion HCI and the IBM technologies integrated within the solution.

### 1.3.1 Integrated solution architecture

This integrated solution, as shown in Figure 1-2, consists of IBM watsonx.data deployed on Red Hat OpenShift hosted by the IBM Fusion HCI. IBM watsonx.data is connected to accelerated buckets hosted in either the public cloud, which includes IBM Cloud, Amazon S3, and Google Cloud Storage, or on-premises infrastructure such as IBM Storage Ceph. By connecting IBM Fusion HCI to external object buckets, high-performance object access is delivered by intelligent caching that is provided by IBM Fusion HCI's storage infrastructure. IBM Fusion HCI exposes the accelerated buckets to IBM watsonx.data for attachment to a query engine (Presto, Spark). Persistent cache is immediately available for newly provisioned engines.

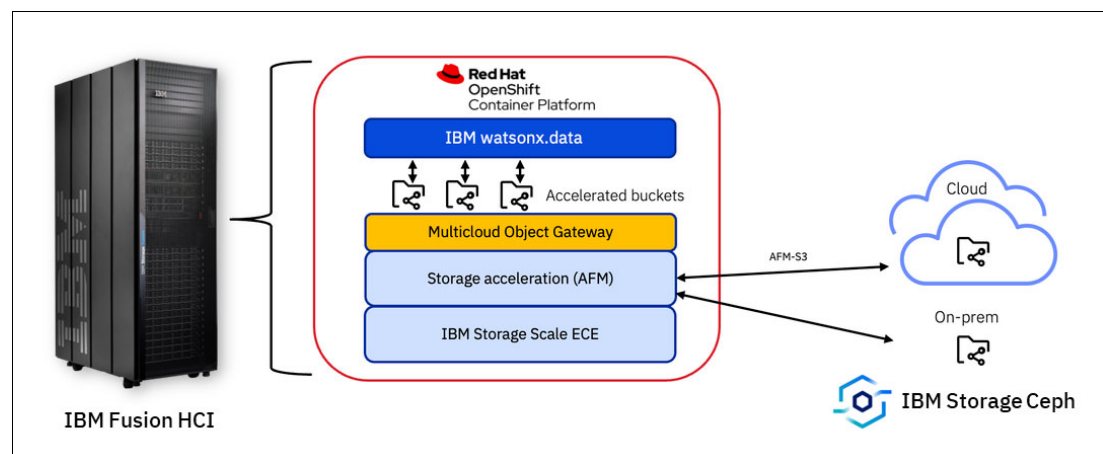


Figure 1-2 IBM watsonx.data Storage Acceleration hosted on IBM Fusion HCI

### 1.3.2 Solution component architectures

This section describes the architectures of the solution components.

#### IBM watsonx.data

IBM watsonx.data is an open, hybrid, and governed data lakehouse optimized for all data and AI workloads. It combines the high performance and usability of a data warehouse with the flexibility and scalability of data lakes. IBM watsonx.data is a unique solution that allows co-existence of open source technologies and proprietary products. It offers a single point of entry where you can store the data or attach data sources for managing and analyzing structured, semi-structured, and unstructured enterprise data, which enables access to all data across cloud and on-premises environments.

The following components as shown in Figure 1-3 provide the foundation of IBM watsonx.data:

- ▶ Open table formats, such as Apache Iceberg provide structure and deliver the reliability of SQL with big data. They allow different engines to access the same data at the same time, and enable data sharing across multiple repositories including data warehouses and data lakes.
- ▶ Query engines access data in an open table format. IBM watsonx.data query engines are fully modular and can be dynamically scaled to meet workload demands and concurrency.
- ▶ The technical metadata service enables the query engine to know the location, format, and read capabilities of the data.
- ▶ Data catalogs assist with finding the correct data and deliver semantic information for policies and rules.
- ▶ The policy engine enables users to define and enforce data protection.

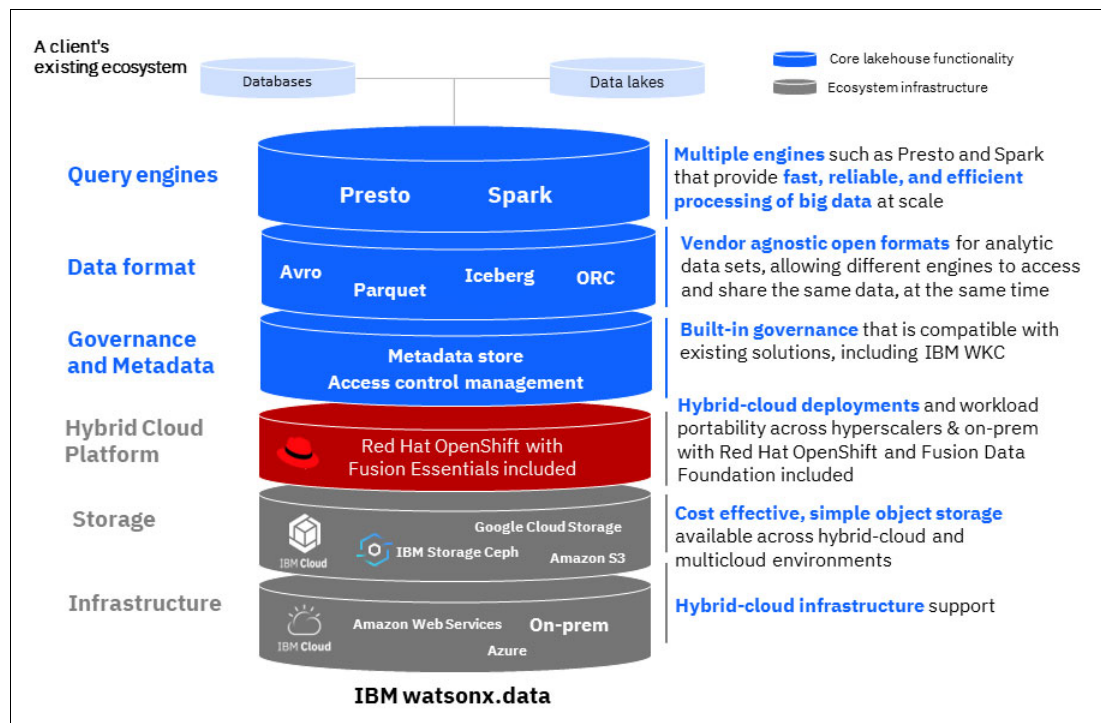


Figure 1-3 IBM watsonx.data architecture

The IBM watsonx.data software stack is built on IBM Fusion HCI and IBM Storage Ceph to provide high-performance infrastructure and storage.

## **IBM Fusion HCI**

IBM Fusion HCI is a hyper-converged appliance that delivers all of the infrastructure needed to run Red Hat OpenShift on bare metal, which eliminates the complexity of designing, deploying, and maintaining an on-premises architecture for IBM watsonx.data. See Figure 1-4 on page 8. The appliance is delivered as a rack with all components mounted, cabled, and tested. It provides all the infrastructure resources that are required to host the Red Hat OpenShift cluster, such as storage nodes, compute nodes, and network switches.

S3 object storage, IBM Storage Ceph, IBM Storage Scale, and NAS file arrays are available in a single namespace. Access by applications is unaffected by the type of storage behind the namespace. Intelligent global data caching enables accessing remote data at local file system speeds. IBM Fusion HCI uses a dedicated network for Red Hat OpenShift traffic and a dedicated, high-performance network for the storage cluster, and provides scalability for expanding workloads. Online migration of data from remote storage systems to the IBM Fusion HCI file system is included.

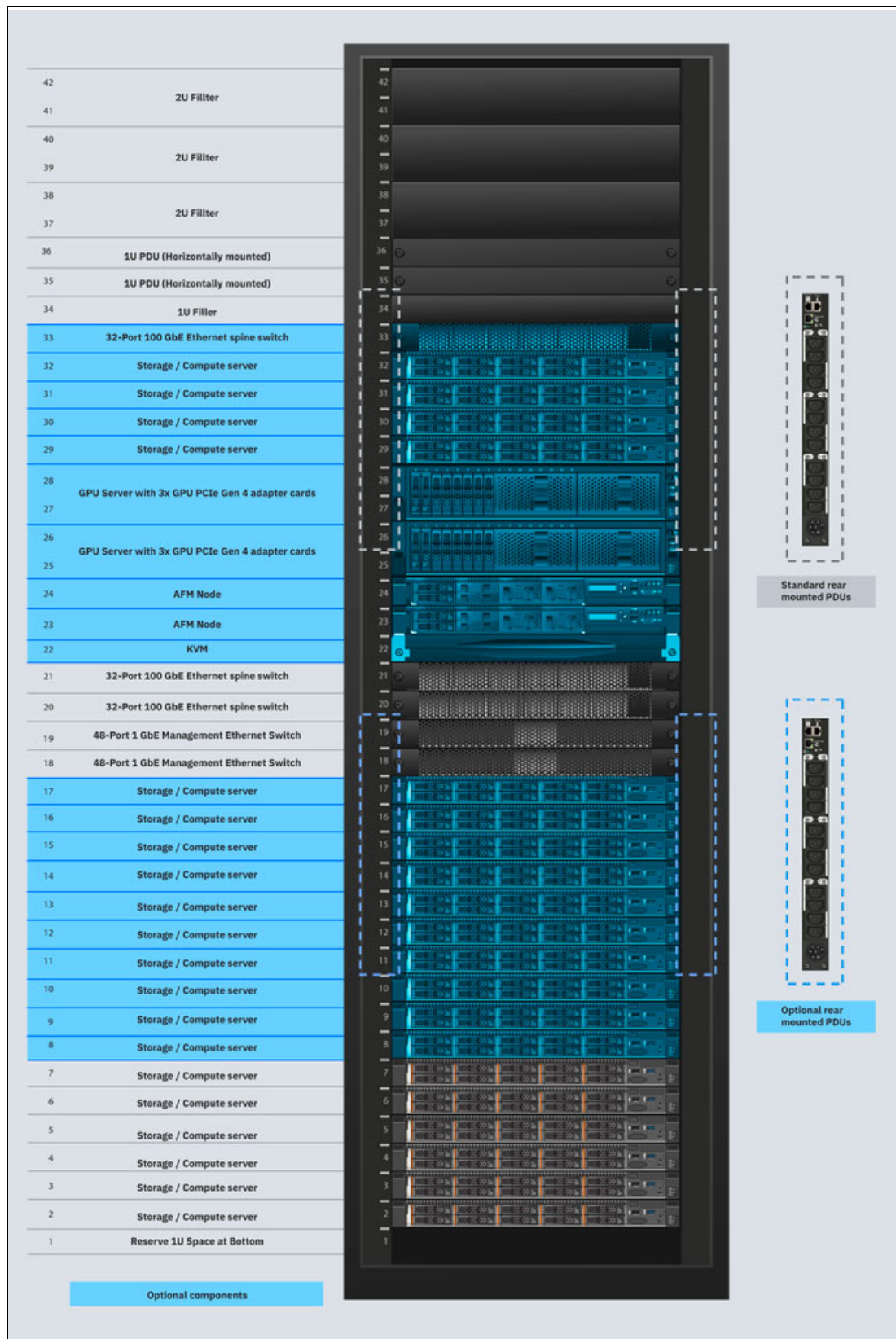


Figure 1-4 IBM Fusion HCI architecture

## IBM Storage Ceph

IBM Storage Ceph is a software-defined storage platform that is based on an open source development model and can be deployed on industry-standard x86 hardware. It provides non-disruptive, horizontal scaling of object, block, and file storage to thousands of clients accessing exabytes of data. It is ideal for modern AI frameworks that require data lake capabilities.

IBM Storage Ceph provides an external S3 object store for IBM watsonx.data. This S3 object store can be the main S3 object store for IBM watsonx.data, or an additional S3 object store with other on-premises or public-cloud object stores. The IBM Storage Ceph object storage interface, the Ceph Object Gateway, is compatible with a large subset of the Amazon S3 RESTful API. See Figure 1-5.

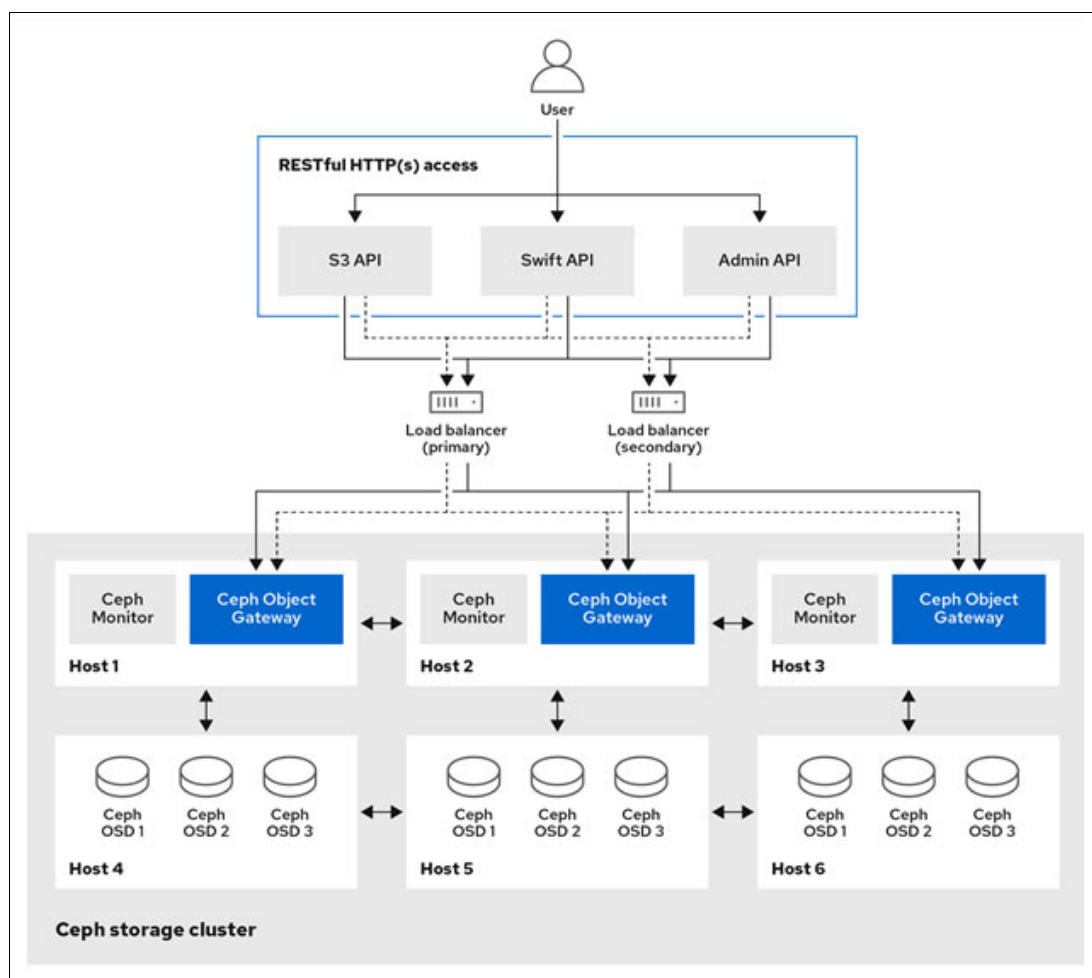


Figure 1-5 Ceph Object Gateway architecture

Multiple Red Hat OpenShift clusters can share storage from the same Ceph S3 object store.

## IBM Storage Scale Erasure Code Edition Active File Management

IBM Fusion HCI uses IBM Storage Scale Erasure Coded Edition (ECE) as an underlying storage platform. IBM Fusion ECE is a high-performance parallel file system that is used in High Performance Computing (HCP) and maximizes storage I/O within a clustered compute environment. This high-performance storage layer provides storage for IBM Cloud Pak for Data and IBM watsonx.data internal services and serves as a cache for storage accelerated buckets.

This solution achieves storage acceleration by using Storage Scale's Active File Management (AFM) technology to connect to existing object storage buckets. These buckets reside in an on-premises object storage solution, such as IBM Storage Ceph or IBM Cloud Object Storage or in a public cloud provider such as AWS, Azure, or IBM Cloud®. AFM is a high-speed cache for buckets it is attached to, allowing for data access that is significantly faster than I/O on the buckets directly. See Figure 1-6.

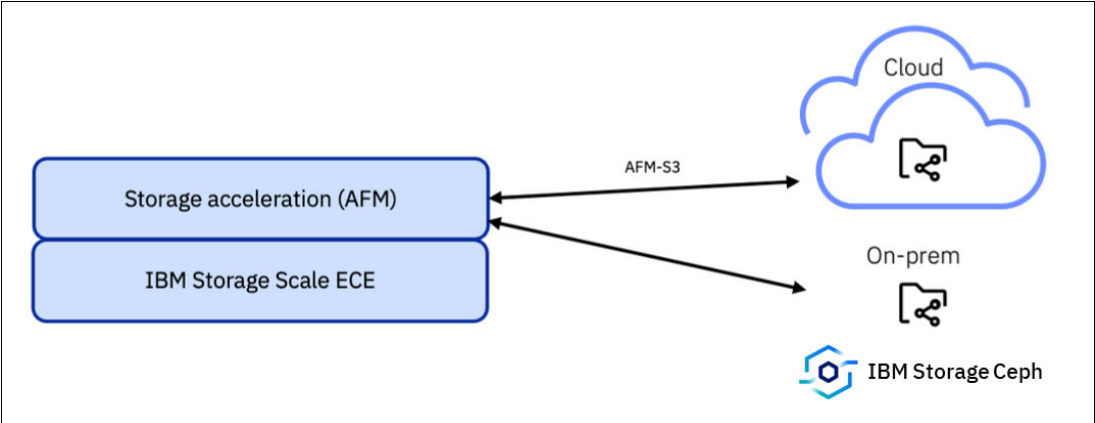


Figure 1-6 Storage acceleration with Active File Management

### Multicloud Object Gateway

Multicloud Object Gateway (MCG) is a lightweight object storage service for Red Hat OpenShift. Although MCG can function as an object storage provider using storage from the IBM Fusion HCI appliance, in this solution, MCG functions as a gateway between IBM watsonx.data and storage accelerated buckets provided by AFM. MCG connects directly to filesets in the Storage Scale ECE file system that map back to object buckets attached to AFM. When IBM watsonx.data reads from buckets, the reads pass through MCG to the AFM cache for the bucket. A cache hit results in a high-performance read. With a cache miss, the object is fetched from the external bucket hosting the data. See Figure 1-7.

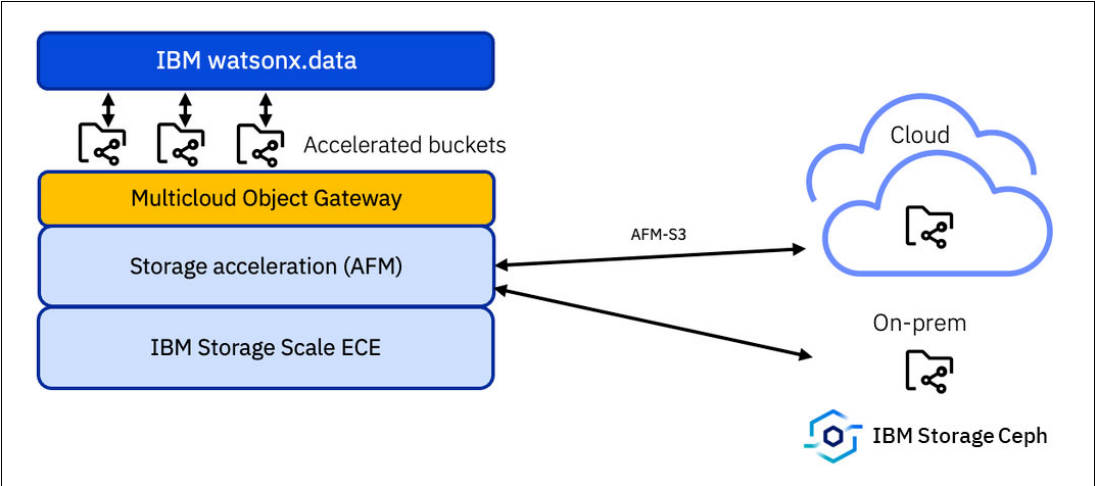


Figure 1-7 Multicloud Object Gateway and Storage Acceleration



# Sizing and planning

This chapter describes sizing guidelines for the licensed components and highlights several planning activities that are related to the solution in this publication.

## 2.1 Sizing guidelines

This section provides sample sizing configurations for the licensed components.

### 2.1.1 Licensing

The following list highlights the software and hardware licensing for IBM Fusion HCI:

- ▶ watsonx is licensed by available cores not total cores:
  - A Fusion HCI 32 core worker node has 20 available cores. Watsonx needs 20 VPC <sup>1</sup>of entitlement per Fusion HCI 32 core server.
  - A Fusion HCI 64 core worker node has 52 available cores. watsonx needs 52 VPC of entitlement per Fusion HCI 64 core server.
  - The other 12 cores in each server are reserved for Red Hat OpenShift and Fusion.
- ▶ Fusion HCI CPUs can support SMT<sup>2</sup>=2:
  - A Fusion HCI 32 core worker node has 40 available vCPU.
  - A Fusion HCI 64 core worker node has 104 available vCPU.
- ▶ Fusion HCI 9155 [Expert Care](#) is required:
  - Expert care provides Fusion hardware support starting at beginning of year 1.
  - Expert care provides Fusion software support starting at beginning of year 1.
- ▶ Strategies for managing excess Fusion HCI hardware capacity:
  - License watsonx to a sub-capacity of the Fusion HCI cluster size.
  - Use excess cluster capacity for other Cloud Paks and workloads. If segregation of workload is required workloads can be isolated in separate namespaces.

Learn more about container licensing [here](#).

### 2.1.2 IBM Fusion HCI with IBM watsonx.data

This section provides guidance to plan for IBM Fusion HCI configurations using watsonx.data to meet your requirements. There are two types of configurations: Standard and Performance:

- ▶ Standard configuration use the following:
  - 32 core worker nodes, each with 512 GB of memory.
  - Fusion HCI 32-core servers have 20 cores (40 vCPU) that are available for workloads after subtracting overhead for Red Hat OpenShift and Fusion software.
- ▶ Performance configurations use the following:
  - 64 core worker nodes each with 2048 GB of memory.
  - Fusion HCI 64-core servers have 52 cores (104 vCPU) that are available for workloads after subtracting overhead for Red Hat OpenShift and Fusion software.

---

<sup>1</sup> Note that Virtual Processor Core (VPC) is a unit of measure by which a Program can be licensed.

<sup>2</sup> SMT is Simultaneous Multi-Threading, See [Cores versus vCPUs and hyperthreading](#)

Detailed configurations follow:

- ▶ “Fusion HCI Standard configurations” on page 13
- ▶ “Fusion HCI Performance configurations” on page 14
- ▶ “Multi Rack Performance and Standard configurations” on page 14

## Fusion HCI Standard configurations

Fusion HCI Standard configurations are shown in Table 2-1.

Table 2-1 Standard configuration details

| Fusion HCI                  |                                    | E03      | E04      | E05      | E06      | E07      | E08      | E09      | E10      | E11      |
|-----------------------------|------------------------------------|----------|----------|----------|----------|----------|----------|----------|----------|----------|
| Fusion hardware<br>M/T 9155 | Control Nodes                      | 3        | 3        | 3        | 3        | 3        | 3        | 3        | 3        | 3        |
|                             | Worker nodes                       | 3        | 4        | 5        | 6        | 7        | 8        | 9        | 10       | 11       |
|                             | Active NVMe drives                 | 12       | 14       | 16       | 18       | 20       | 22       | 24       | 26       | 28       |
|                             | Query accelerator nodes            | Optional | Optional | Optional | Optional | Optional | Optional | Optional | Optional | Optional |
|                             | Total worker cores                 | 96       | 128      | 160      | 192      | 224      | 256      | 288      | 320      | 352      |
|                             | Available cores <sup>a</sup>       | 60       | 80       | 100      | 120      | 140      | 160      | 180      | 200      | 220      |
|                             | Available memory <sup>b</sup> (GB) | 1,296    | 1,728    | 2,160    | 2,592    | 3,024    | 3,456    | 3,888    | 4,320    | 4,752    |
|                             | Usable NVMe <sup>c</sup> (TB)      | 59       | 69       | 79       | 89       | 99       | 109      | 119      | 129      | 139      |
| Fusion software<br>5771-PP7 | VPCs to license                    | 96       | 128      | 160      | 192      | 224      | 256      | 288      | 320      | 352      |
| watsonx<br>D0F4SZX          | VPCs <sup>d</sup> to license       | 60       | 80       | 100      | 120      | 140      | 160      | 180      | 200      | 220      |

a. SMT=1.

b. This is a double memory configuration with 16 GB RAM per core. You can reduce cost by configuring with 8 GB RAM per core.

c. Decimal TB. The formula is 7.68 TB x number of drives x 0.65. The 0.65 approximates 4+2p erasure coding overhead.

d. watsonx.data is licensed per usable / available VPC not raw cores.

## Fusion HCI Performance configurations

Fusion HCI Performance configurations are shown in Table 2-2.

Table 2-2 Performance configuration details

| Fusion HCI                  |                                    | P03      | P04      | P05      | P06      | P07      | P08      | P09      | P10      | P11      |
|-----------------------------|------------------------------------|----------|----------|----------|----------|----------|----------|----------|----------|----------|
| Fusion hardware<br>M/T 9155 | Control Nodes                      | 3        | 3        | 3        | 3        | 3        | 3        | 3        | 3        | 3        |
|                             | Worker nodes                       | 3        | 4        | 5        | 6        | 7        | 8        | 9        | 10       | 11       |
|                             | Active NVMe drives                 | 12       | 14       | 16       | 18       | 20       | 22       | 24       | 26       | 28       |
|                             | Query accelerator nodes            | Optional | Optional | Optional | Optional | Optional | Optional | Optional | Optional | Optional |
|                             | Total worker cores                 | 192      | 256      | 320      | 384      | 448      | 512      | 576      | 640      | 704      |
|                             | Available cores <sup>a</sup>       | 156      | 208      | 260      | 312      | 364      | 416      | 468      | 520      | 572      |
|                             | Available memory <sup>b</sup> (GB) | 5,904    | 7,872    | 9,840    | 11,808   | 13,776   | 15,744   | 17,712   | 19,680   | 21,648   |
|                             | Usable NVMe <sup>c</sup> (TB)      | 59       | 69       | 79       | 89       | 99       | 109      | 119      | 129      | 139      |
| Fusion software<br>5771-PP7 | VPCs to license                    | 192      | 256      | 320      | 384      | 448      | 512      | 576      | 640      | 704      |
| watsonx<br>D0F4SZX          | VPCs <sup>d</sup> to license       | 156      | 208      | 260      | 312      | 364      | 416      | 468      | 520      | 572      |

a. SMT=1.

b. This is a double memory configuration with 16 GB RAM per core. You can reduce cost by configuring with 8 GB RAM per core.

c. Decimal TB. The formula is 7.68 TB x number of drives x 0.65. The 0.65 approximates 4+2p erasure coding overhead.

d. watsonx.data is licensed per usable / available VPC not raw cores.

## Multi Rack Performance and Standard configurations

Table 2-3 shows the Multi Rack Performance configurations.

Table 2-3 Multi Rack Performance configurations

| Size                     | 64-core worker nodes | Available vCPU (SMT=1) | Available vCPU (SMT=2) | Total Memory in GB |
|--------------------------|----------------------|------------------------|------------------------|--------------------|
| <b>Two racks</b>         |                      |                        |                        |                    |
| P11 + P03 <sup>a</sup>   | 17                   | 884                    | 1768                   | 33,456             |
| P11 + P11                | 25                   | 1300                   | 2600                   | 49,200             |
| <b>Three racks</b>       |                      |                        |                        |                    |
| P11+P11+P03 <sup>a</sup> | 31                   | 1612                   | 3224                   | 61,008             |
| P11 + P11 + P11          | 39                   | 2028                   | 4056                   | 76,752             |

a. The 3 control nodes in rack 2 and rack 3 are converted to worker nodes.

Table 2-4 shows the Multi Rack Standard configurations.

Table 2-4 Multi rack Standard configuration

| Size                     | 32-core worker nodes | Available vCPU (SMT=1) | Available vCPU (SMT=2) | Total Memory in GB |
|--------------------------|----------------------|------------------------|------------------------|--------------------|
| <b>Two racks</b>         |                      |                        |                        |                    |
| E11 + E03 <sup>a</sup>   | 17                   | 340                    | 680                    | 7,344              |
| E11 + E05                | 19                   | 380                    | 760                    | 8,208              |
| E11 + E07                | 21                   | 420                    | 840                    | 9,072              |
| E11 + E09                | 23                   | 460                    | 920                    | 9,936              |
| E11 + E11                | 25                   | 500                    | 1000                   | 10,800             |
| <b>Three racks</b>       |                      |                        |                        |                    |
| E11+E11+E03 <sup>a</sup> | 31                   | 620                    | 1240                   | 13,392             |
| E11+E11+E05              | 33                   | 660                    | 1320                   | 14,256             |
| E11 + E11 + E07          | 35                   | 700                    | 1400                   | 15,120             |
| E11 + E11 + E09          | 37                   | 740                    | 1480                   | 15,984             |
| E11 + E11 + E11          | 39                   | 780                    | 1560                   | 16,848             |

a. The 3 control nodes in rack 2 and rack 3 are converted to worker nodes.

## 2.2 Planning

Planning tasks help ensure that the IBM Fusion HCI is accurately integrated with IBM watsonx.data and configured properly for your operations.

For more information, see [Planning and prerequisites](#).

### 2.2.1 Network access to object storage

The IBM Fusion HCI connects to the data center network during the initial appliance setup. The appliance includes two high-speed switches that are connected to the core network through one port channel. This connection acts as the gateway between the IBM Fusion appliance and the network. It enables administration of the appliance and Red Hat OpenShift and is also used for network traffic in and out of the cluster. Network resources and applicable configuration settings are applied during this setup phase.

For more information, see [Network planning](#).

### 2.2.2 AFM

AFM creates associations between clusters and the data source. It provides a single, global namespace across sites to automate the flow of data. AFM is enabled on the fileset that connects to the remote S3 endpoint to access the cache.

The cache fileset is served by the AFM node, which functions as a gateway. As a gateway, the AFM node owns the fileset and communicates regarding data transfers.

Make sure the IBM Fusion HCI has AFM nodes installed and configured as part of the Red Hat OpenShift cluster. AFM nodes must be installed and configured before you begin any storage acceleration operations. In addition, determine how large the tier 1 cache must be for storage acceleration. Ensure there is sufficient usable storage for that cache, in addition to the storage that is used for the Cloud Pak for Data and IBM watsonx.data PVCs.

For more information, see [Planning for AFM](#) and [Sharing Data](#).

## 2.3 Customer use cases

IBM watsonx provides the crucial data analytics and AI capabilities that all large organizations require. The strength of IBM Fusion HCI for IBM watsonx.data is the appliance-like experience where Red Hat OpenShift cluster, compute, storage, and network is all in a single box. It allows for a shorter time to value by having everything you need, so you can immediately start performing queries by using IBM watsonx.data.

Consider the following key use cases for IBM watsonx.data:

- ▶ **AI and machine learning (ML) at scale**  
Build, train, tune, deploy, and monitor trusted AI and ML models for mission-critical workloads with governed data in IBM watsonx.data and ensure compliance with lineage and reproducibility of data used for AI.
- ▶ **Real-time analytical and business intelligence**  
Combine data from existing sources with new data to unlock new, faster insights without the cost and complexity of duplicating and moving data across different environments.
- ▶ **Streamline data engineering**  
Reduce data pipelines, simplify data transformation, and enrich data for consumption using SQL, Python, or an AI-infused conversational interface.
- ▶ **Responsible data sharing**  
Enable self-service access for more users to more data while ensuring security and compliance through centralized governance and local automated policy enforcement.

### 2.3.1 Data sharing

IBM Db2 Warehouse has the option to write and read to and from a cloud bucket using open formats such as parquet and iceberg. This allows for seamless integrating and sharing of data between IBM Db2 Warehouse and IBM watsonx.data without the need for deduplication or additional extract, transform, load operations. This might reduce your costs for storage that is used by IBM Db2 Warehouse and offload some of the workloads to IBM watsonx.data.



# Implementation

Implementation involves the combination of IBM Fusion HCI and IBM Storage Ceph, which provide all the infrastructure for a stand-alone data lakehouse. The implementation includes installation of IBM Fusion HCI, Multicloud Object Gateway (MCG) configuration, Active File Management (AFM) configuration, performance tuning, and the installation of IBM watsonx.data.

## 3.1 IBM Fusion HCI installation

If you already have an IBM Fusion HCI installation, ensure that your Red Hat OpenShift Container Platform is at Version 4.12.7.

For the procedure to install IBM Fusion HCI 2.7.x, see [Deploying IBM Fusion HCI](#).

## 3.2 Installing Red Hat OpenShift Data Foundation and configuring the Multicloud Object Gateway

The Multicloud Object Gateway (MCG) provides an object endpoint to which IBM watsonx.data and other workloads can connect to access multiple buckets, including Storage Acceleration buckets. The MCG is provided by the Red Hat OpenShift Data Foundation operator.

Install the Red Hat OpenShift Data Foundation operator into Red Hat OpenShift Container Platform:

1. Log in to your Red Hat OpenShift Container Platform.
2. Go to OperatorHub and search for Data Foundation operator.
3. Type **Data Foundation** in the Filter by keyword field to find the Data Foundation operator.
4. Click **Install**. See Figure 3-1.
5. Enter ibm-storage-fusion-cp-sc StorageClass that is configured by default in IBM Fusion HCI, and click **Install**.
6. Click **Install**.

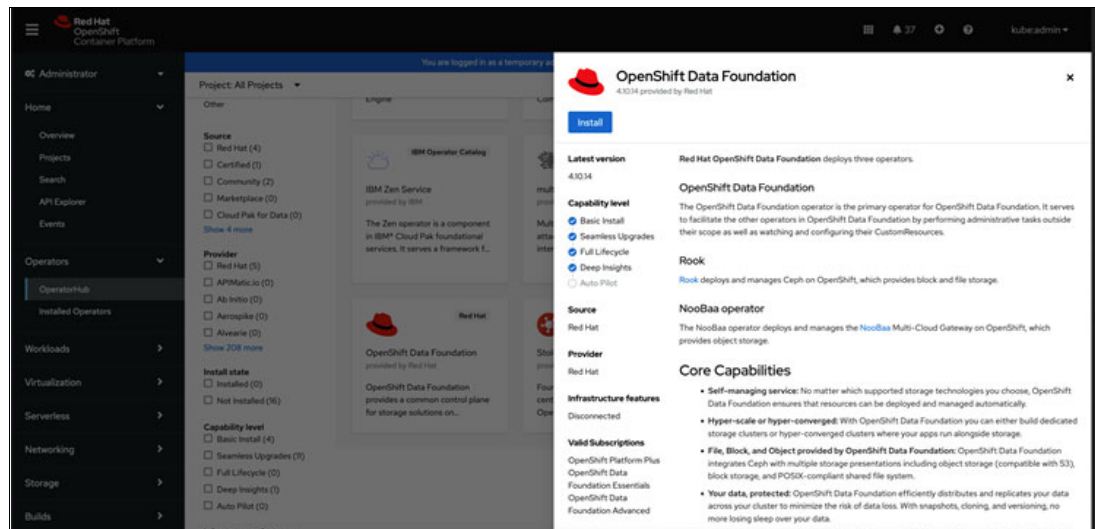


Figure 3-1 Installing Red Hat OpenShift Data Foundation

Create StorageSystem for Red Hat OpenShift Data Foundation:

1. Ensure that you select the Installation Mode as A specific namespace on the cluster and click **Install**.
2. Click **Create StorageSystem** for Red Hat OpenShift Data Foundation.
3. In the Deployment type field, select **MultiCloud Object Gateway**. For more information about MCG deployment, see [Deploy stand-alone Multicloud Object Gateway](#).

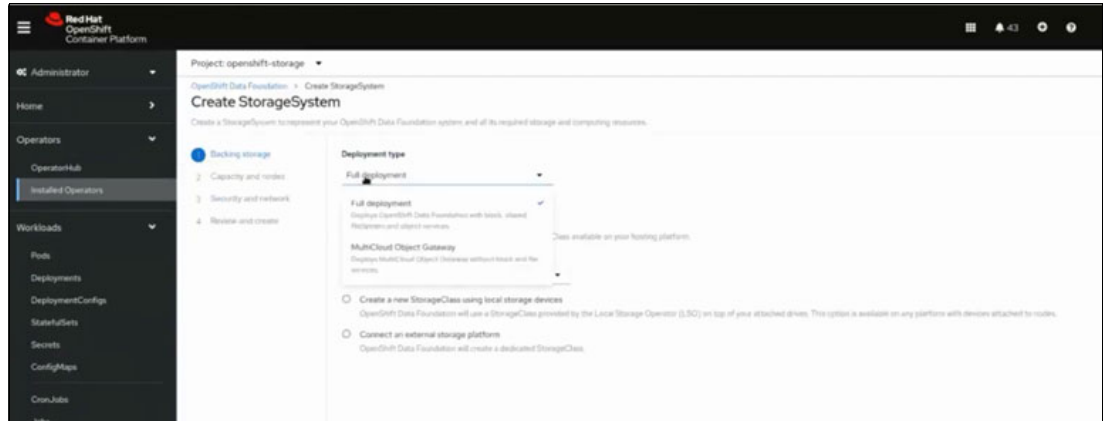


Figure 3-2 Deployment type selection in Create StorageSystem

### 3.2.1 Configuring Advanced File Management nodes

Configure Storage Acceleration on IBM Fusion HCI to connect your remote object bucket to the Storage Scale Advanced File Management (AFM) accelerator and then expose an accelerated bucket through the MCG. This workflow has two steps:

1. Create a Storage Scale AFM fileset that is connected to the remote S3 endpoint.
2. Connect the MCG to the local Scale AFM fileset.

After you complete the configuration, you can access the AFM cached remote S3 data.

#### Prerequisites

This list describes the tools and environment prerequisites:

- ▶ Ensure that the Fusion HCI rack has AFM nodes installed, and ensure they are configured to be part of the Red Hat OpenShift cluster.
- ▶ The **oc** command is used to issue commands to the Red Hat OpenShift Container Platform cluster.
- ▶ Noobaa is installed as a part of MCG. It is used to access the Red Hat OpenShift Data Foundation MCG, and the AWS/S3 CLI is used to access the noobaa API endpoint. If the installation is not available, see [noobaa-operator](#).
- ▶ Check whether you can log in to the Red Hat OpenShift Container Platform cluster.
- ▶ Retrieve the values for bucket name, access key, secret key, and S3 endpoint for the object bucket that you want to accelerate. The remote S3 can be IBM Storage Ceph for an on-premises environment. The S3 can be a cloud provider, such as IBM Cloud Object Storage (COS) or AWS S3. For more information, see [IBM Storage Ceph](#) and IBM Redpaper: IBM Storage Ceph Solutions Guide, [REDP-5715](#).

## Configuring the AFM nodes

Before you configure the AFM nodes, collect the following information:

- ▶ Endpoint as REMOTE\_S3\_ENDPOINT
- ▶ Access Key as REMOTE\_S3\_ACCESS\_KEY
- ▶ Access Secret as REMOTE\_S3\_ACCESS\_SECRET
- ▶ Bucket Name as REMOTE\_S3\_BUCKET

### Createing an AFM fileset and connecting to a remote S3 bucket

1. Get the Scale core pod name on AFM node:

```
AFM_NODE_POD_NAME=$(oc get node -l scale.spectrum.ibm.com/role=afm -o json | jq  
-r '.items[0].metadata.name' | awk -F '.' '{print $1}')
```

2. Go into the Scale core pod:

```
oc exec -it $AFM_NODE_POD_NAME bash -n ibm-spectrum-scale
```

3. Input REMOTE\_S3\_ENDPOINT, REMOTE\_S3\_BUCKET, REMOTE\_S3\_ACCESS\_KEY, and REMOTE\_S3\_ACCESS\_SECRET:

```
- REMOTE_S3_ENDPOINT=http://s3.us-south.cloud-object-storage.appdomain.cloud  
- REMOTE_S3_BUCKET=afm-s3test  
- REMOTE_S3_ACCESS_KEY=e48acxxxx750  
- REMOTE_S3_ACCESS_SECRET=8ffc2xxxxx1e22d
```

4. Create an access key:

```
mmafmcoskeys $REMOTE_S3_BUCKET set $REMOTE_S3_ACCESS_KEY  
$REMOTE_S3_ACCESS_SECRET
```

5. Input the AFM fileset, endpoint, bucket, and afm\_mode.

6. Create an AFM fileset, endpoint, bucket, and afm\_mode:

For example, create an AFM fileset, and the AFM node is in mode \*lu\*. See Example 3-1.

---

#### Example 3-1 Creating an AFM fileset

```
FILE_SYSTEM=ibmspectrum-fs  
fileset=afm-cos-s3-fileset  
AFM_MODE=lu
```

```
mmafmcosconfig $FILE_SYSTEM $fileset --endpoint $REMOTE_S3_ENDPOINT --object-fs  
--bucket $REMOTE_S3_BUCKET --cleanup --debug --mode $AFM_MODE --tmpdir .noobaa%
```

```
mmchfileset $FILE_SYSTEM $fileset -p afmPrefetchThreshold=100
```

```
mmafmcosctl $FILE_SYSTEM $fileset /mnt/${FILE_SYSTEM}/${fileset} download --all  
--metadata
```

---

The following list describes the available modes and their purposes:

- ▶ Independent Writer (IW) is for changes made from the cache and server. This option must be configured when you are setting up both read and write cache. As you set the accelerator on top of the object bucket, the accelerator works both as a read and write cache for the object. It is the default setting.
- ▶ Local Update (LU) is for changes that are made on only the server. In this mode, you can use it for testing of your model. You do not want the changes you are making to go to the backend object bucket. After the test is complete, you can change to the IW mode.
- ▶ Single Writer (SW) is for changes made only from cache. In this mode, only the cache fileset does all the writing and the cache does not check home for file or directory updates.

### ***Performance tuning***

Do configuration settings for performance tuning. For more information about tuning configuration, see 3.2.3, “Performance tuning” on page 27.

### ***Evicting cache data manually***

You can evict cache manually, or you can control eviction by defining quotas.

Perform the following steps to evict cache manually:

- ▶ Evict all cache data manually:  

```
mmafmcctl fs1 fileset1 /gpfs/fs1/new1 evict -all
```
- ▶ Evict all cache data and metadata:  

```
mmafmcctl fs1 fileset1 /gpfs/fs1/new1 evict -all --metadata
```
- ▶ Manual eviction after quota limit is set. Evict data by using a criteria:  

```
mmafmcctl fs1 evict -j fileset1 --order LRU  
mmafmcctl fs1 evict -j fileset1 --order Size
```

### ***Evicting cache by using quota enabled eviction***

Eviction can also be automatically controlled by using quotas. After the soft quota is exceeded, AFM automatically evicts the files based on LUR fashion. If a policy is not set for eviction, after the limit is reached, then the requests will fail with no space return error code. For more information, see [Enabling quotas](#).

In this example, set the soft quota to be 1 TB and the hard quota to be 2 TB. Adjust the values based on the PV/PVC size that you plan to create.

```
mmsetquota $FILE_SYSTEM:$fileset --block 1000G:2000G
```

Use either `mlsquota` or `mmrepquota` to view your quotas.

## **3.2.2 Creating the static PV and PVC**

The steps in this section describe how to create the static PV and PVC.

### **Defining the variables**

Assign the scale cluster ID to the variable `CLUSTER_ID`:

```
CLUSTER_ID=$(oc exec $AFM_NODE_POD_NAME -n ibm-spectrum-scale "mmlscluster" |  
grep "GPFS cluster id:" | awk '{print $4}')
```

Assign the file system ID to the variable `FILESYSTEM_ID`:

```
FILESYSTEM_ID=$(oc exec $AFM_NODE_POD_NAME -n ibm-spectrum-scale -- bash -c 'mmlsfs ibmspectrum-fs --uid' | grep "uid" | awk '{print $2}')
```

## Updating and applying the YAML template

Complete the following steps:

1. Update the YAML template for PV name, capacity, volumeHandle and PVC name, as shown in Example 3-2.

### *Example 3-2 Updating the YAML template*

---

```
apiVersion: v1
kind: PersistentVolume
metadata:
  name: {{PV_Name}}
spec:
  accessModes:
    - ReadWriteMany
  capacity:
    storage: {{Capacity}} csi:
      driver: spectrumscale.csi.ibm.com
      volumeHandle:
0;2;{{CLUSTER_ID}};{{FILESYSTEM_ID}};{{fileset}};/mnt/ibmspectrum-fs/{{fileset}}
  persistentVolumeReclaimPolicy: Retain
  volumeMode: Filesystem
---
apiVersion: v1
kind: PersistentVolumeClaim
metadata:
  name: {{PVC_Name}}
  namespace: openshift-storage
spec:
  accessModes:
    - ReadWriteMany
  resources:
    requests:
      storage: {{Capacity}}
  storageClassName: ""
  volumeMode: Filesystem
  volumeName: {{PV_Name}}
```

---

Example 3-3 is an example YAML file is an example using PV name, capacity, volumeHandle and PVC name. You will need to use your naming conventions for your organization.

### *Example 3-3 Example of updating the YAML template*

---

```
Example yaml as file pv_pvc.yaml
apiVersion: v1
kind: PersistentVolume
metadata:
  name: afm-cos-s3-remote-pv
spec:
  accessModes:
    - ReadWriteMany
  capacity:
```

```

    storage: 1000Gi
  csi:
    driver: spectrumscale.csi.ibm.com
    volumeHandle:
0;2;6734170828145876673;1180A8C0:646EEE59;;afm-cos-s3-fileset;/mnt/ibmspectrum-fs/
afm-cos-s3-fileset
    persistentVolumeReclaimPolicy: Retain
    volumeMode: Filesystem
---
apiVersion: v1
kind: PersistentVolumeClaim
metadata:
  name: afm-cos-s3-remote-pvc
  namespace: openshift-storage
spec:
  accessModes:
  - ReadWriteMany
  resources:
    requests:
      storage: 1000Gi
    storageClassName: ""
    volumeMode: Filesystem
    volumeName: afm-cos-s3-remote-pv

```

---

2. Apply the yaml to create static PV and PVC. The following command is an example:

```
oc apply -f pv_pvc.yaml
```

## Configuring the MCG bucket

Configure an MCG bucket that uses the PVC to access the AFM cache.

As a prerequisite, install ODF operator and create an ODF cluster with MCG only mode.

1. Create a NamespaceStore with the PVC, as shown in Example 3-4.

*Example 3-4 Creating a NamespaceStore with the PVC*

---

```

apiVersion: noobaa.io/v1alpha1
kind: NamespaceStore
metadata:
  name: {{NSS_Name}}
  namespace: openshift-storage
spec:
  nsfs:
    fsBackend: GPFS
    pvcName: {{PVC_Name}}
    subPath: ""
  type: nsfs

```

---

2. Create a bucket class with the NamespaceStore, as shown in Example 3-5.

*Example 3-5 Creating a bucket class with the NamespaceStore*

---

```
apiVersion: noobaa.io/v1alpha1
kind: BucketClass
metadata:
  name: {{Bucket_Class_Name}}
  namespace: openshift-storage
spec:
  namespacePolicy:
    single:
      resource: {{NSS_Name}}
      type: Single
```

---

3. Create a Noobaa Account for the NamespaceStore, as shown in Example 3-6.

*Example 3-6 Creating a Noobaa Account for the NamespaceStore*

---

```
apiVersion: noobaa.io/v1alpha1
kind: NooBaaAccount
metadata:
  name: {{Noobaa_Account_Name}}
  namespace: openshift-storage
spec:
  allow_bucket_creation: true
  default_resource: {{NSS_Name}}
  nsfs_account_config:
    gid: 0
    new_buckets_path: /
    nsfs_only: true
    uid: 0
```

---

4. Create an ObjectBucketClaim with the BucketClass as shown in Example 3-7.

*Example 3-7 Creating an ObjectBucketClaim with the BucketClass*

---

```
apiVersion: objectbucket.io/v1alpha1
kind: ObjectBucketClaim
metadata:
  name: {{Object_Bucket_Claim_Name}}
  namespace: openshift-storage
spec:
  additionalConfig:
    bucketclass: {{Bucket_Class_Name}}
  bucketName: {{Bucket_Name}}
  storageClassName: openshift-storage.noobaa.io
```

---

Example 3-8 is an example of creating an ObjectBucketClaim with the BucketClass. You will need to use your organizations naming conventions.

*Example 3-8 Example showing YAML creating an ObjectBucketClaim with the BucketClass*

---

```
Example YAML:
apiVersion: noobaa.io/v1alpha1
kind: NamespaceStore
metadata:
  name: afm-cos-s3-nss
```

```

    namespace: openshift-storage
spec:
  nsfs:
    fsBackend: GPFS
    pvcName: afm-cos-s3-remote-pvc
    subPath: data
    type: nsfs
---
apiVersion: noobaa.io/v1alpha1
kind: BucketClass
metadata:
  name: afm-cos-s3-bc
  namespace: openshift-storage
spec:
  namespacePolicy:
    single:
      resource: afm-cos-s3-nss
    type: Single
---
apiVersion: noobaa.io/v1alpha1
kind: NooBaaAccount
metadata:
  name: afm-cos-s3-acc
  namespace: openshift-storage
spec:
  allow_bucket_creation: true
  default_resource: afm-cos-s3-nss
  nsfs_account_config:
    gid: 0In e
    new_buckets_path: /
    nsfs_only: true
    uid: 0
---
apiVersion: objectbucket.io/v1alpha1
kind: ObjectBucketClaim
metadata:
  name: afm-cos-s3-obc
  namespace: openshift-storage
spec:
  additionalConfig:
    bucketclass: afm-cos-s3-bc
    bucketName: afm-cos-s3-nss-bc
    storageClassName: openshift-storage.noobaa.io

```

---

5. Apply the yaml file from Example 3-8 on page 24 to create resources.

## Updating the bucket access policy

You need the noobaa admin account to update the bucket access policy.

1. Get AWS\_ACCESS\_KEY\_ID and AWS\_SECRET\_ACCESS\_KEY from the output:

```
noobaa status --show-secrets
```

2. Create the policy.json file with the content in Example 3-9.

*Example 3-9 Creating the policy.json file*

---

```
{
  "Version": "2012-10-17",
  "Statement": [
    {
      "Sid": "id-1",
      "Effect": "Allow",
      "Principal": "*",
      "Action": ["s3:*"],
      "Resource": ["arn:aws:s3:::*"]
    }
  ]
}
```

---

3. Add a bucket policy (Example 3-10).

Replace the bucket name with the bucket name that you defined in the previous steps.

*Example 3-10 Adding a bucket policy*

---

ObjectBucketClaim.Spec.bucketName field

BUCKET\_NAME=afm-cos-s3-nss-bc

S3\_ENDPOINT=https://\$(oc get route s3 -n openshift-storage -o json | jq -r  
'\$.status.ingress[0].host')

```
AWS_ACCESS_KEY_ID=$AWS_ACCESS_KEY_ID AWS_SECRET_ACCESS_KEY=$AWS_SECRET_ACCESS_KEY
aws --endpoint-url=$S3_ENDPOINT --no-verify-ssl s3api put-bucket-policy --bucket
$BUCKET_NAME --policy file://policy.json
```

---

4. Verify the bucket:

```
BUCKET_NAME=afm-cos-s3-nss-bc
```

- a. Get the noobaa account. The noobaa account name that was defined in the previous steps.

```
NoobaaAccount.metadata.name
```

```
NOOBAA_ACCOUNT_NAME=afm-cos-s3-acc
```

```
noobaa account status $NOOBAA_ACCOUNT_NAME --show-secrets
```

- b. Get the AWS\_ACCESS\_KEY\_ID and AWS\_SECRET\_ACCESS\_KEY from the output:

```
AWS_ACCESS_KEY_ID=$AWS_ACCESS_KEY_ID
```

```
AWS_SECRET_ACCESS_KEY=$AWS_SECRET_ACCESS_KEY aws --endpoint $S3_ENDPOINT
--no-verify-ssl s3api list-objects --bucket $BUCKET_NAME
```

### 3.2.3 Performance tuning

You defined a storage accelerated bucket that can be accessed by workloads such as `watsonx.data`. The following list describes the default parameters that can be changed or disabled to improve performance:

|                                     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
|-------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>afmFileLookupRefreshInterval</b> | Defines the frequency of revalidation that is triggered by a look-up operation on a file such as <code>ls</code> or <code>stat</code> , from the IBM Fusion HCI. AFM sends a message to the external object bucket to determine that the metadata of the file is modified since it was last revalidated. If so, the latest metadata information at home is reflected on the IBM Fusion HCI.                                                                                                                         |
| <b>afmFileOpenRefreshInterval</b>   | Defines the frequency of revalidations that are triggered by the read and write operations on a file from the IBM Fusion HCI. AFM sends a message to the external object bucket to determine if the metadata of the file was modified since it was last revalidated.                                                                                                                                                                                                                                                |
| <b>afmDirLookupRefreshInterval</b>  | Defines the frequency of revalidation that is triggered by a look-up operation such as <code>ls</code> or <code>stat</code> on a directory from the IBM Fusion HCI. AFM sends a message to the external object bucket to find out whether the metadata of that directory is modified since it was last revalidated. If so, the latest metadata information at the external object bucket is reflected on the IBM Fusion HCI.                                                                                        |
| <b>afmDirOpenRefreshInterval</b>    | Defines the frequency of revalidations that are triggered by the read and update operations on a directory from the IBM Fusion HCI. AFM sends a message to the external object bucket to find whether the metadata of that directory is modified since it was last revalidated.                                                                                                                                                                                                                                     |
| <b>afmObjectFastReaddir</b>         | Improves the objects download and readdir performance, when the <code>afmObjectFastReaddir</code> parameter value is set to <code>yes</code> at the fileset level. Extended attributes and ACLs are not fetched from a cloud object storage when this parameter is enabled. Also, deleted objects on a cloud object storage system are not reflected immediately on a cache when this parameter is enabled. You can use this option to pull the objects into the cache to run quick analytics on the target data.   |
| <b>afmParallelReadChunkSize</b>     | Defines the minimum chunk size of the read that needs to be distributed among the gateway nodes during parallel reads. A zero (0) value disables the parallel reads across multiple gateways. The parallel reads are routed through a single gateway node.                                                                                                                                                                                                                                                          |
| <b>afmObjectFastReaddir</b>         | Improves the objects download and readdir performance, when the <code>afmObjectFastReaddir</code> parameter value is set to <code>'yes'</code> at the fileset level. Extended attributes and ACLs are not fetched from a cloud object storage when this parameter is enabled. Also, deleted objects on a cloud object storage system are not reflected immediately on a cache when this parameter is enabled. You can use this option to pull the objects into the cache to run quick analytics on the target data. |

## Data not changing on the server

When you configure an accelerator or caching for an object bucket, there might not be any changes to the data on the object bucket.

Changes to the applications go to the object bucket through the accelerator. If there is no possibility of anybody making updates to the object bucket without the accelerator, then you can disable the refresh intervals.

For example, a Ceph object bucket exists, and the accelerator is placed on top of it. Applications in turn run on top of the accelerator.

If no changes are happening on the object bucket outside the accelerator path, then use the following configuration parameters:

```
mmchfileset $FILE_SYSTEM $fileset -p afmFileLookupRefreshInterval=disable
mmchfileset $FILE_SYSTEM $fileset -p afmFileOpenRefreshInterval=disable
mmchfileset $FILE_SYSTEM $fileset -p afmDirLookupRefreshInterval=disable
mmchfileset $FILE_SYSTEM $fileset -p afmDirOpenRefreshInterval=disable
```

## Data might change on the server

When the data can change on the server, then you can disable the file lookup interval without disabling the directory lookup. You can increase the revalidation time from the default of 60 seconds.

Enter the following commands to alter the frequency of revalidation:

```
mmchfileset $FILE_SYSTEM $fileset -p afmObjectFastReaddir=yes
mmchfileset $FILE_SYSTEM $fileset -p afmFileLookupRefreshInterval=disable
```

## Size of the download from each Gateway

AFM can read from multiple gateways at the same time. If the amount of data that is read is greater than a defined number, then parallel reads begin. For example, if 12 MB is the chunk size, then each gateway reads 12 MB and then pass the data to the main gateway to process the data.

```
mmchfileset $FILE_SYSTEM $fileset -p afmParallelReadChunkSize=12M
```

## *How it improves performance*

Parallel read data transfer improves the overall data read transfer performance of an AFM to cloud object storage fileset by using multiple gateway nodes.

## Multiple gateways

For filesets with mode LU or RO, you can use multiple gateways for better performance. Configure more gateways with the following commands:

```
mmchnode -gateway -N node
mmchfileset fs fileset-name -p afmGateway=all
```

To activate **afmGateway=all**, stop and restart the fileset by using the following commands:

```
mmafmctl perffs stop -j db2wh-db2u-perf-test-2-cos
mmchfileset perffs db2wh-db2u-perf-test-2-cos -p afmGateway=all
mmafmctl perffs start -j db2wh-db2u-perf-test-2-cos
```

For the other modes, such as IW and SW, use the **mmafmconfig** command:

```
mmafmconfig {add | update} MapName --export-map ExportServerMap
[--no-server-resolution]
```

The following command is a sample of the `mmafmconfig` command:

```
mmafmconfig add mymap --export-map  
169.46.118.100/fin37.ibm.com,10.242.33.16/fin38.ibm.com
```

The 2nd IP address in each pair is the gateway, and `fin37.ibm.com` and `fin38.ibm.com` are the addresses of the gateways.

### Populating the metadata cache

If you want to pre-fetch your metadata cache, you have several options.

The `mmchfileset` command retrieves the name entries the fastest from listv2 with minimum attributes. You can use the `ls` command to view the metadata. Use the following format for the command:

```
mmchfileset device filesetname -p afmObjectFastReaddir=yes
```

The `mmafmcscctl` command is not as fast as the `mmchfileset` command. The command bypasses the gateway and retrieves the names and full attributes:

```
mmafmcscctl device filesetName path download --metadata --outband
```

The `mmafmcscctl` is slower than the other commands. It uses the gateway and retrieves the names and full attributes:

```
mmafmcscctl device filesetame path download --metadata
```

## 3.3 Installing IBM watsonx.data on IBM Fusion HCI

Complete the following steps to install IBM watsonx.data.

1. Install IBM watsonx.data. For the procedure to install, see [Installing wastonx.data](#).
2. To configure IBM watsonx.data for IBM Fusion HCI storage, create a storage class with the appropriate settings for use with IBM watsonx.data. For the actual procedure to configure, see [Setting up IBM Storage Scale storage](#).
3. Create a watsonx.data instance. The watsonx.data operator is installed one time on the cluster and shared by many instances of watsonx.data on the cluster.

### Creating an accelerated bucket

As a prerequisite, create an AFM fileset and attach it to a bucket. Use the following procedure to create an accelerated bucket and connect to an existing externally managed object storage (Multicloud Gateway):

1. Log in to watsonx.data console.
2. From the navigation menu, select **Infrastructure Manager**.

3. To define and connect a bucket, click **Add component** and select **Add bucket**.
4. In the **Add bucket** window, provide the following details to connect to the accelerated bucket provided by MCG:

**Note:** Refer to the values that you set during the accelerated bucket creation for the Bucket name, endpoint, access key, and secret key.

- Bucket type. Select Ceph as the value for bucket type from list.
- Bucket name. Enter the name of your existing bucket.
- Endpoint. Enter the endpoint URL.
- Access key. Enter your access key.
- Secret key. Enter your secret key.
- Activation. Activate the bucket immediately or activate it later.
- Catalog type. Select the catalog type from the list.
- Catalog name. Enter the name of the catalog. The catalog is automatically associated with your bucket.

To add a bucket-catalog pair, see [Adding a bucket-catalog pair](#).



# Monitoring

The IBM Fusion HCI is like an Red Hat OpenShift cluster in a box. With the available default Red Hat OpenShift monitoring options, you can monitor the watsonx.data project or namespace. In addition, the IBM Fusion HCI has its own monitoring and logging capabilities to view different dashboards for storage, networking, and compute. For more information, see [Monitoring and logging](#).

However, for monitoring watsonx.data specifically, use the built-in Presto engine web interface for monitoring and managing queries. This web interface is accessible by an Red Hat OpenShift route to the coordinator pod. For more information, see [Exposing secure route to Presto server](#).

After the route has been exposed, you can open a web browser to the route's URL. The main page has a list of queries and includes information such as unique query ID, query text, query state, percentage completed, username, and source from which this query originated. The currently running queries are at the top of the page, followed by the most recently completed or failed queries.

Figure 4-1 shows an example of the main page.

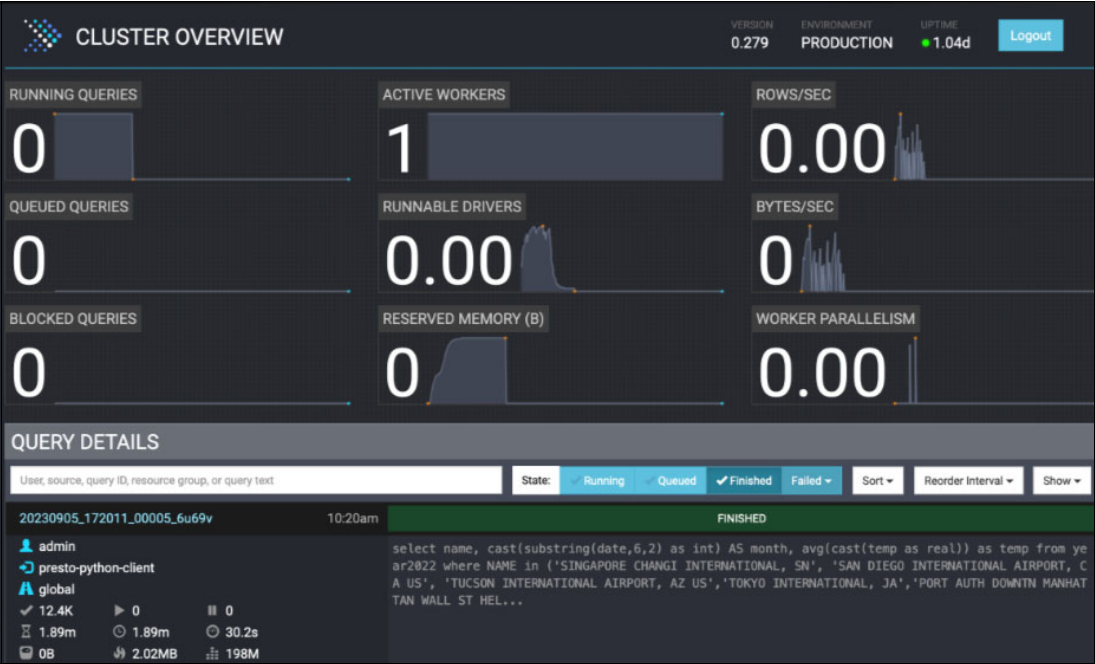


Figure 4-1 An example of the main page

For more detailed information about a query, click the query ID link. The query detail page has a summary section, graphical representation of various stages of the query and a list of tasks. Each task ID can be clicked to get more information about that task. For example, when you click the task ID, you see a page similar to Figure 4-2.

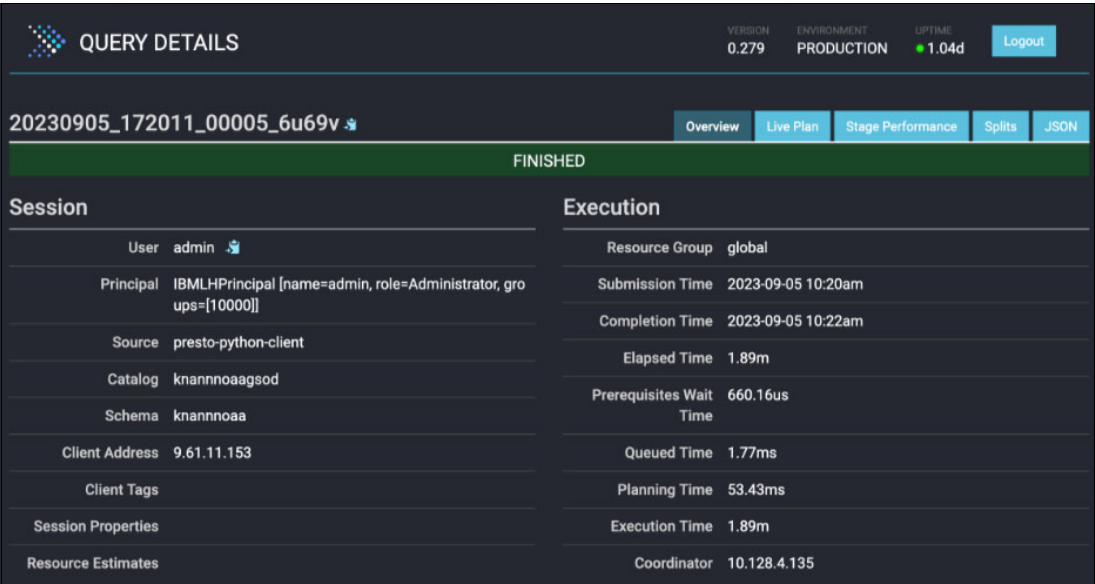


Figure 4-2 An example of the detailed view of an example task ID



## Backup and restore of IBM Cloud Pak for Data

This chapter describes backing up and restoring Cloud Pak for Data with IBM watsonx.data on IBM Fusion. It covers the non-disruptive backup of Cloud Pak for Data on a Fusion HCI and restore to an alternative Fusion HCI. The process includes setting up a backup location, creating backup policies on the Fusion hub system, and applying the policies on to the Cloud Pak for Data application that is deployed.

## 5.1 Considerations and requirements

The example includes two IBM Fusion HCI racks running Version 2.6.1 that are connected as hub and spoke. Supported versions of Cloud Pak for Data and watsonx.data service are installed on the hub system as the source for backup.

The following software and configurations are required when backing up your Cloud Pak for Data environment:

- ▶ IBM Fusion HCI 2.6.1 or later.
- ▶ Cloud Pak for Data 4.7.1 or later.
- ▶ Red Hat OpenShift Container Platform versions must be at the same major version on both source and target clusters. For example, IBM Fusion HCI 2.6.1 supports OpenShift Container Platform 4.10 and 4.12 and both source and target clusters must be at the same major version.
- ▶ Cloud Pak for Data and its services at the same release level.
- ▶ A supported version of Cloud Pak for Data is installed in private topology, and each Cloud Pak for Data instance includes the following components:
  - Two namespaces:
    - cpd-operator
    - cpd-instance
  - One shared cluster components: ibm-scheduler (optional scheduling service)
- ▶ With IBM Fusion 2.6.1, Backup & Restore should be configured across two Fusion HCI clusters with one of the clusters acting as the hub. The hub controls the backup and restore flow with one or more clusters that are connected to the hub as spokes. This setup allows for backups that are taken in one cluster to be restored in a different cluster.
- ▶ Cloud Pak for Data along with IBM watsonx.data can be installed on either the hub cluster or the spoke cluster.

## 5.2 Getting the prerequisites ready

Install more components before backing up the Cloud Pak for Data environment.

### 5.2.1 Installing the Cloud Pak Backup and Restore service

Complete the following steps:

1. Install the cpdbr-oadp service in the Cloud Pak for Data operators and Cloud Pak for Data shared namespaces of the cluster components, which include Cloud Pak for Data Scheduling service (if installed). Ignore the Scheduling service namespace if it is not installed. The cpdbr-oadp service must be installed on both the Hub and Spoke clusters.

To install the service, prepare your Hub and Spoke clusters to use cpd-cli. For more information, see [Cloud Pak for Data command-line interface CPD](#).

2. Install the cpdbr-oadp service in the following namespaces:
  - cpd-operator
  - cpd-scheduler

For more information, see [Installing the cpdbr service for IBM Fusion integration](#).

3. On the source cluster, install the cpdbr-oadp service by issuing the following command:

```
./cpd-cli oadp install \  
--component=cpdbr-tenant \  
--tenant-operator-namespace=<cpd-operator_ns> \  
--cpdbr-hooks-image-prefix=quay.io/cpdsre \  
--cpd-scheduler-namespace=cpd-scheduler \  
--log-level=debug \  
--verbose
```

4. After installation is done, verify that the cpdbr pod is deployed by running the following command:

```
oc get pods -A | grep cpdbr
```

The following line is an example of the expected output:

```
cpd-operator cpdbr-tenant-service-6dcc49464c-rr9jh
```

The installation of cpdbr-oadp also installs, generates, and applies the required recipes in the respective Cloud Pak for Data cpd-operator and cpd-scheduler namespaces (if installed). To verify, issue the command **oc get frcpe -A**, as shown in Example 5-1.

*Example 5-1 Example output*

---

```
oc get frcpe -A
```

| NAMESPACE              | NAME                 | AGE   |
|------------------------|----------------------|-------|
| cpd-operator           | ibmcpd-tenant        | 2m6s  |
| ibm-spectrum-fusion-ns | fusion-control-plane | 3d23h |
| ibm-spectrum-fusion-ns | fusion-cr-backup     | 10d   |

---

## 5.2.2 Installing the cpdbr service on the target cluster

Complete the following steps:

1. Install the cpdbr-tenant service on the target cluster, as shown in Example 5-2.

*Example 5-2 Installing cpdbr-tenant service on the target cluster*

---

```
$ cpd-cli oadp install --component=cpdbr-tenant --tenant-operator-namespace=cpd-operator  
processing request...  
cpd tenant operator namespace: cpd-operator  
clusterrole/cpdbr-tenant-service-clusterrole created  
clusterrolebinding/cpdbr-tenant-service-crb created  
role/cpdbr-tenant-service-role created in namespace kube-public  
rolebinding/cpdbr-tenant-service-rb created in namespace kube-public
```

---

2. After the installation is done, verify that ClusterRoleBinding was created, as shown in Example 5-3.

*Example 5-3 Verifying that ClusterRoleBinding was created*

```
$ oc get clusterrolebinding cpdbr-tenant-service-crb
```

| NAME                     | ROLE                                         | AGE |
|--------------------------|----------------------------------------------|-----|
| cpdbr-tenant-service-crb | ClusterRole/cpdbr-tenant-service-clusterrole | 37s |

```
$ oc get clusterrole cpdbr-tenant-service-clusterrole
```

| NAME                             | CREATED AT           |
|----------------------------------|----------------------|
| cpdbr-tenant-service-clusterrole | 2023-08-25T17:45:55Z |

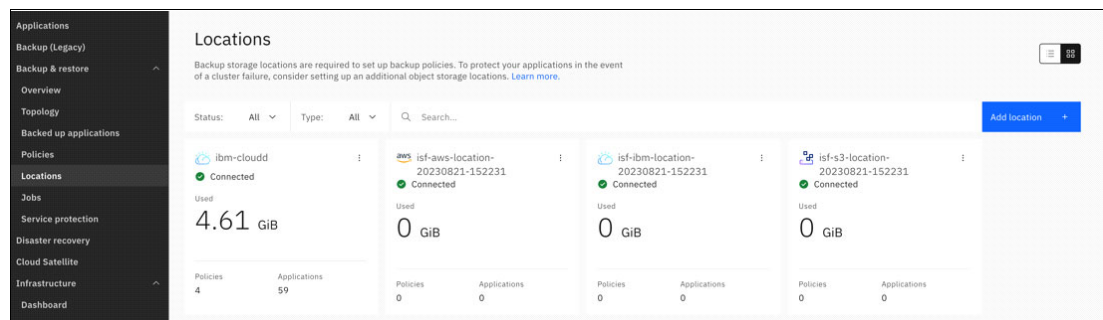
```
$ oc get clusterrolebinding | grep cpdbr
cpdbr-tenant-service-crb
```

## 5.2.3 Backup policies for Cloud Pak for Data applications

Before you create and apply backup policies to the Cloud Pak for Data applications, you must create an S3 compliant backup location in the Fusion UI of the hub system. The hub system is used to store the backups and is the source when the data is restored.

Complete the following steps:

1. To create a backup location, on IBM Fusion UI, select **Backup & restore** → **Locations** to add a backup location. In this setup, which is shown in Figure 5-1, we added an object storage from IBM Cloud as a backup location.



*Figure 5-1 Adding a backup location*

2. After creating the backup location, run the following command to list the backup locations that were created:

```
$ oc get fbsl -n ibm-spectrum-fusion-ns
```

| NAME       | PROVIDER           | PHASE     | STORAGETYPE |
|------------|--------------------|-----------|-------------|
| ibm-cloudd | isf-backup-restore | Connected | ibm         |

## 5.2.4 Creating and assigning a backup policy

Complete the following steps:

1. On IBM Fusion UI, select **Backup & restore** → **Policies** to add a backup policy, as shown in Figure 5-2. A Backup Policy specifies how frequently backups are taken, where backups are stored, and how long backups are retained.

The screenshot shows the IBM Storage Fusion UI. The left sidebar contains navigation options: Quick start, Events, Applications, Backup (Legacy), Backup & restore (selected), Overview, Topology, Backed up applications, Policies (selected), Locations, Jobs, Service protection, Disaster recovery, Cloud Satellite, Infrastructure, Dashboard, Nodes, Network, Services, and Settings. The main area displays the 'Policies' page. It includes a table of backup policies and a sidebar on the right showing details for the 'cpd-oper-policy'.

| Name                                           | Backup location   | Frequency                      | Time                                    | Retention |
|------------------------------------------------|-------------------|--------------------------------|-----------------------------------------|-----------|
| cpd-oper-policy                                | ibm-cloud         | Every day                      | 1:00 AM India Standard Time             | 30 Days   |
| ibm-ae1                                        | ibm-cloud         | Every day                      | 11:40 AM India Standard Time            | 2 Days    |
| ibm-test123                                    | ibm-cloud         | Every day                      | 12:00 AM Central European Standard Time | 2 Days    |
| isf-auto-inplace-backup-policy-20230821-152231 | In place snapshot | Every day                      | 9:38 AM Pacific Standard Time           | 1 Day     |
| spoke-ibm                                      | ibm-cloud         | Every M, Tu, W, Th, F, Sat, Su | 3:30 PM India Standard Time             | 2 Days    |
| spoke-inplace                                  | In place snapshot | Every 5 hours at minute 00     | America/Argentina/Buenos_Aires          | 1 Day     |
| test-inplace                                   | In place snapshot | Every 2 hours at minute 00     | Asia/Kolkata                            | 1 Day     |
| test123                                        | In place snapshot | Every day                      | 12:00 AM Central European Standard Time | 30 Days   |

The right sidebar shows details for the 'cpd-oper-policy':

- Name:** cpd-oper-policy
- Schedule:** Every day
- Frequency:** Every day
- Time:** 1:00 AM India Standard Time
- Retention:** 30 Days
- Location:** ibm-cloud
- Type:** IBM
- Location status:** Connected
- Used:** 4.61 GiB
- Applications (3):** apps.rackae1.mydomain.com, cpd-operator, ibm-cert-manager, ibm-licensing

Figure 5-2 Adding a backup policy

To list backup policies from the CLI, run the **oc get backuppolicies** command:

```
$oc get backuppolicies -n ibm-spectrum-fusion-ns
NAME          PROVIDER          BACKUPSTORAGELOCATION          SCHEDULE
RETENTION    RETENTIONUNIT
cpd-oper-policy    isf-backup-restore    ibm-cloud
00 1 * * *          30          days
```

- Assign the backup policy to Cloud Pak for Data applications. From the IBM Fusion UI, select **Backed up applications**, and then open the **Protect apps** menu and select the cluster where Cloud Pak for Data is deployed. Select the following applications:
  - cpd-operator
  - ibm-scheduler (if installed)

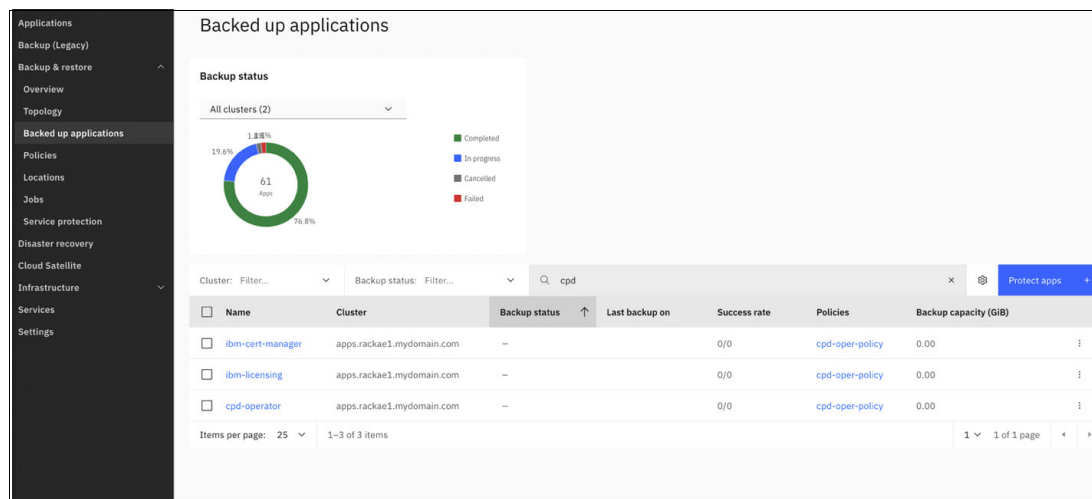


Figure 5-3 Verifying that the recipes are now associated to the corresponding policy assignments

- Examine the policies that are assigned by running the command in Example 5-4.

*Example 5-4 Checking the policies that are assigned*

```
$ oc get policyassignments.data-protection.isf.ibm.com -n ibm-spectrum-fusion-ns
| grep cpd-oper
NAME PROVIDER APPLICATION BACKUPPOLICY RECIPE RECIPENAMESPACE PHASE LASTBACKUPTIMESTAMP
CAPACITY
cpd-operator-cpd-oper-policy-apps isf-backup-restorecpd-operator cpd-oper-policyibmcpd-tenant
cpd-operator
Assigned <no value>
```

The recipes are not yet associated to the policy assignment. The recipes must be manually patched into the policies except for cpd-tenant, which is assigned automatically. If cpd-scheduler is assigned, the recipe must be patched. Example 5-5 shows patching cpd-scheduler.

*Example 5-5 Patching the recipes into the policy assignment*

```
$ oc -n ibm-spectrum-fusion-ns patch policyassignment
<cpd-scheduler-policy-assignment> --type merge -p
'{"spec":{"recipe":{"name":"ibmcpd-scheduler", "namespace":"cpd-scheduler",
"apiVersion":"spp-data-protection.isf.ibm.com/v1alpha1"}}}'
```

- Check the policy assignments again. The recipes are now attached to the backup policy assignment in Example 5-6.

*Example 5-6 Checking the policy assignments again*

```
$ oc get policyassignments.data-protection.isf.ibm.com -n ibm-spectrum-fusion-ns | grep cpd-oper
NAME PROVIDER APPLICATION BACKUPPOLICY RECIPE
RECIPENAMESPACE PHASE LASTBACKUPTIMESTAMP CAPAC
cpd-operator-cpd-oper-policy-apps isf-backup-restore cpd-operator cpd-oper-policy
ibmcpd-tenant cpd-operator Assigned <no value>
```

## 5.3 Backing up the source cluster

Complete the following steps:

1. Take a backup on the hub by selecting **Backed up applications** → **cpd-operator**. Then, select **Backup now** → **Backup**.

**Important:** The first backup of cpd-operator prepares follow-on backups to be valid for restore. Do *not* restore from the first backup. After the first backup is complete, take a second backup. The second backup, and all later backups, may be used for the restore.

2. After all the backups are finished, go to the Backed up applications page to confirm that the backups finished successfully.

## 5.4 Restoring to an alternative cluster

Before restoring to an alternative cluster, ensure that the target cluster is prepared for Cloud Pak for Data and watsonx.data installation. For more information, see [Preparing your cluster](#).

Complete the following steps:

1. When the alternative cluster is ready, edit guardian-configmap in the ibm-backup-restore project by increasing the restoreDatamoveTimeout parameter value to 240 minutes.
2. Then, select **Backup & restore** → **Topology** and verify that the hub and spoke are connected and in a healthy state.

The next step is to install the certificate manager and the IBM License Service. For more information, see [Installing shared cluster components for IBM Cloud Pak for Data](#).

3. Change the logging level from default INFO to DEBUG by running the following command:

```
oc patch cm -n ibm-backup-restore guardian-configmap
-p='{"data":{"logLevel":"DEBUG"}}'
```

4. Next, restore cpd-operator by selecting **Backed up applications** → **cpd-operator** → **Restore**.
5. Under Select a destination, click **Choose a different cluster to restore the application in** and then select the target cluster. Click **Next**.
6. In the next window, select the backup that you want to use and then click **Next**.
7. After the job is finished, verify the completion details by going to the Jobs section under Backup & restore.

8. Repeat these steps for “ibm-scheduler” (if installed). After the restore is finished for “cpd-operator”, and “ibm-scheduler”, confirm that the restore was successful by logging in to the Cloud Pak for Data user interface. Inspect the watsonx.data instance and all the previous data under the Instances page, as shown in Figure 5-4.

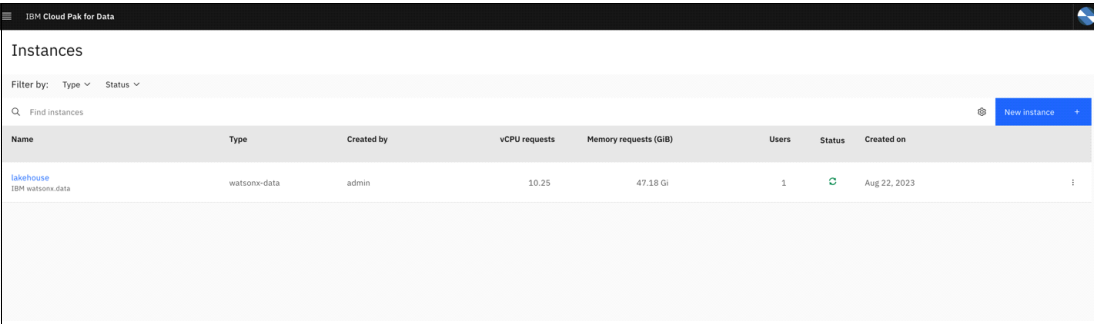


Figure 5-4 Verification of the restored watsonx.data instance in the Cloud Pak for Data user interface

Figure 5-5 shows that the watsonx.data test engine was restored successfully.

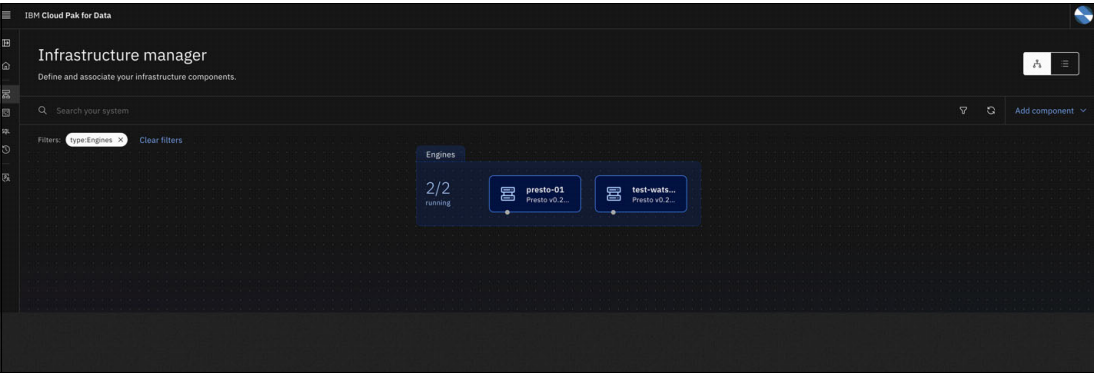


Figure 5-5 The watsonx.data test engine was restored successfully

Figure 5-6 shows the data that is restored from the “Data manager” view.

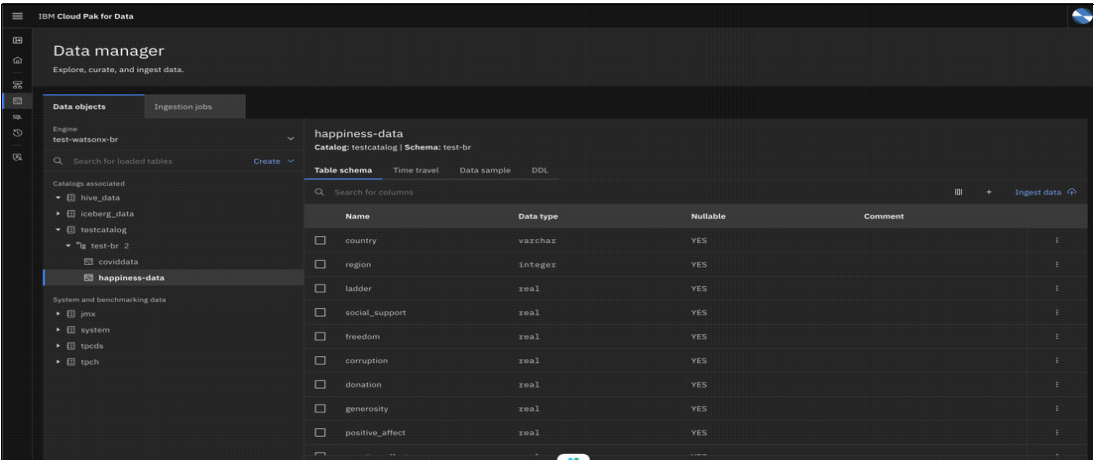


Figure 5-6 Data was restored from the “Data manager” view

9. Verify the restore by running `oc get pods -n cpd-instance` and ensuring that all pods are in a good state, as shown in Example 5-7.

*Example 5-7 Verifying the restore by running `oc get pods -n cpd-instance`*

---

| \$ oc get pods -n cpd-instance                   |       |           |             |       |
|--------------------------------------------------|-------|-----------|-------------|-------|
| NAME                                             | READY | STATUS    | RESTARTS    | AGE   |
| create-secrets-job-nwt8q                         | 0/1   | Completed | 0           | 8m35s |
| ibm-lh-lakehouse-hive-metastore-696f8fb6dd-8ss85 | 1/1   | Running   | 3 (16m ago) | 21m   |
| ibm-lh-lakehouse-minio-ff8f7b77f-h4n5r           | 1/1   | Running   | 0           | 21m   |
| ibm-lh-lakehouse-presto-01-presto-0              | 1/1   | Running   | 2 (16m ago) | 21m   |
| ibm-lh-lakehouse-presto543-presto-0              | 1/1   | Running   | 2 (16m ago) | 21m   |
| ibm-lh-postgres-edb-2                            | 1/1   | Running   | 0           | 26m   |
| ibm-lh-postgres-edb-3                            | 1/1   | Running   | 0           | 25m   |
| ibm-lh-postgres-edb-4                            | 1/1   | Running   | 0           | 24m   |
| ibm-lh-postgres-setup-job-8q6zf                  | 0/1   | Completed | 0           | 10m   |
| ibm-nginx-6995f698fd-9s9vv                       | 2/2   | Running   | 0           | 19m   |
| ibm-nginx-6995f698fd-sgvq6                       | 2/2   | Running   | 0           | 19m   |
| ibm-nginx-tester-55588dd7b-pnjvx                 | 2/2   | Running   | 0           | 21m   |
| lhconsole-api-85f77cc57d-k4wk9                   | 1/1   | Running   | 5 (18m ago) | 21m   |
| lhconsole-api-85f77cc57d-vmzkv                   | 1/1   | Running   | 5 (18m ago) | 21m   |
| lhconsole-nodeclient-6bb7475775-7x4js            | 1/1   | Running   | 0           | 21m   |
| lhconsole-ui-7c7dbb98d8-t8kpj                    | 1/1   | Running   | 0           | 21m   |
| usermgmt-6bf557c77c-2fs9g                        | 1/1   | Running   | 0           | 18m   |
| usermgmt-6bf557c77c-s5ql6                        | 1/1   | Running   | 0           | 18m   |
| usermgmt-ensure-tables-job-6vpp2                 | 0/1   | Completed | 0           | 7m20s |
| zen-audit-67944bcc74-v2445                       | 1/1   | Running   | 0           | 21m   |
| zen-core-5f7786c596-bmr9n                        | 2/2   | Running   | 3 (18m ago) | 19m   |
| zen-core-5f7786c596-hzb8v                        | 2/2   | Running   | 3 (18m ago) | 19m   |
| zen-core-api-58f7f7664d-2thq8                    | 2/2   | Running   | 0           | 19m   |
| zen-core-api-58f7f7664d-97fzz                    | 2/2   | Running   | 0           | 19m   |
| zen-core-create-tables-job-x7qb7                 | 0/1   | Completed | 0           | 6m50s |
| zen-core-pre-requisite-job-qh5sg                 | 0/1   | Completed | 0           | 5m7s  |
| zen-metastore-edb-2                              | 1/1   | Running   | 0           | 26m   |
| zen-metastore-edb-3                              | 1/1   | Running   | 0           | 25m   |
| zen-minio-0                                      | 1/1   | Running   | 0           | 21m   |
| zen-minio-1                                      | 1/1   | Running   | 0           | 21m   |
| zen-minio-2                                      | 1/1   | Running   | 0           | 21m   |
| zen-minio-create-buckets-job-ccmdv               | 0/1   | Completed | 0           | 8m44s |
| zen-pre-requisite-job-lxghr                      | 0/1   | Completed | 0           | 5m51s |
| zen-validate-metastore-edb-connection-job-m979g  | 0/1   | Completed | 0           | 7m44s |
| zen-watchdog-7f5dcd6789-zkcr1                    | 1/1   | Running   | 5 (13m ago) | 18m   |
| zen-watchdog-create-tables-job-kgz7q             | 0/1   | Completed | 0           | 6m36s |
| zen-watchdog-post-requisite-job-cgchv            | 0/1   | Completed | 0           | 3m33s |
| zen-watchdog-pre-requisite-job-x6pn5             | 0/1   | Completed | 0           | 3m51s |
| zen-watcher-6db76dd5c5-rw6jj                     | 2/2   | Running   | 0           | 18m   |

---

10.Run **oc get pvc -n cpd-instance** to ensure that each pvc is in a good state and bound, as shown in Example 5-8.

*Example 5-8 Ensuring that each pvc is in a good state and bound*

---

```
$ oc get pvc -n cpd-instance
```

| NAME                           | STATUS                   | VOLUME                                   |
|--------------------------------|--------------------------|------------------------------------------|
| export-zen-minio-0             | Bound                    | pvc-0b5e56d4-8dff-42fd-9dc9-bbf903365bfd |
| 10Gi RWO                       | ibm-storage-fusion-cp-sc | 21m                                      |
| export-zen-minio-1             | Bound                    | pvc-3729cd73-6950-4fb3-ace6-d86389b50f5d |
| 10Gi RWO                       | ibm-storage-fusion-cp-sc | 21m                                      |
| export-zen-minio-2             | Bound                    | pvc-b92377fb-a2d7-4973-983e-d04988fda54c |
| 10Gi RWO                       | ibm-storage-fusion-cp-sc | 21m                                      |
| ibm-lh-lakehouse-minio-pvc     | Bound                    | pvc-608134b5-0c6c-41b9-9537-ff5e0878851d |
| 488284Mi RWO                   | ibm-storage-fusion-cp-sc | 43m                                      |
| ibm-lh-postgres-edb-2          | Bound                    | pvc-adf13e39-14af-4b55-866f-27e3184a6157 |
| 9540Mi RWO                     | ibm-storage-fusion-cp-sc | 43m                                      |
| ibm-lh-postgres-edb-3          | Bound                    | pvc-cf179bc5-38e0-49f9-bc16-b535d31b0f00 |
| 9765625Ki RWO                  | ibm-storage-fusion-cp-sc | 26m                                      |
| ibm-lh-postgres-edb-4          | Bound                    | pvc-be3b1749-7f31-487b-b60e-aa4f733adb65 |
| 9765625Ki RWO                  | ibm-storage-fusion-cp-sc | 25m                                      |
| ibm-zen-objectstore-backup-pvc | Bound                    | pvc-d03edb68-67bb-4799-a386-5f3d223cd7be |
| 20Gi RWO                       | ibm-storage-fusion-cp-sc | 43m                                      |
| zen-metastore-edb-2            | Bound                    | pvc-88f1dd8e-3534-4ee3-82b1-3ef2b5904f25 |
| 10Gi RWO                       | ibm-storage-fusion-cp-sc | 43m                                      |
| zen-metastore-edb-3            | Bound                    | pvc-7d9498fc-3c17-456d-b45d-90576b1d8b0d |
| 10Gi RWO                       | ibm-storage-fusion-cp-sc | 26m                                      |

---

11.Run **oc get catalogsource -n cpd-operator** and **oc get pods -n cpd-operator** to ensure that the Cloud Pak for Data operators are in a good state, as shown in Example 5-9.

*Example 5-9 Ensuring that the Cloud Pak for Data operator is in a good state*

---

```
$ oc get catalogsource -n cpd-operator
```

| NAME                                                    | DISPLAY                                         |
|---------------------------------------------------------|-------------------------------------------------|
| cloud-native-postgresql-catalog                         |                                                 |
| ibm-cloud-native-postgresql-4.14.0+20230616.111503      | grpc IBM 39m                                    |
| cpd-platform                                            | ibm-cp-datacore-3.2.0+20230818.131833.371.420-8 |
| grpc IBM 39m                                            |                                                 |
| ibm-watsonx-data-catalog                                |                                                 |
| ibm-watsonx-data-1.0.2+20230816.142123.1192-linux-amd64 | grpc IBM 38m                                    |
| opencloud-operators                                     | ibm-cp-common-services-4.1.0                    |
| grpc IBM 37m                                            |                                                 |

```
$ oc get pods -n cpd-operator
```

| NAME                                                             | READY | STATUS |
|------------------------------------------------------------------|-------|--------|
| 28347d5b35b4a7e67ebbadc34bae6a27cf624ee1ec0388b16aa215aa76mjpbk  | 0/1   |        |
| Completed 0 37m                                                  |       |        |
| 457c18b305fb59f54375ecc9faa6f530db9ca6c2f2adbf6a4c2d831673shjq5  | 0/1   |        |
| Completed 0 36m                                                  |       |        |
| 82a84c3f4c0679f4c09e1261e671982f2405935dd8a66d50738940740dc1lj48 | 0/1   |        |
| Completed 0 31m                                                  |       |        |
| 8ee3408bcdd092e94f0be116278a44a1007c5fa49b13bcdbfbc1c6f002r7p57  | 0/1   |        |
| Completed 0 36m                                                  |       |        |

|                                                                 |     |         |
|-----------------------------------------------------------------|-----|---------|
| 94368157bcda899a2502d5cdf67291342961a91e61937aa778e80dc348skh7k | 0/1 |         |
| Completed 0 35m                                                 |     |         |
| b41a5640f98ac37869ac16c0eccdee0b225cc565114472ab5a50df351ergbsc | 0/1 |         |
| Completed 0 37m                                                 |     |         |
| cloud-native-postgresql-catalog-nqjn4                           | 1/1 | Running |
| 0 40m                                                           |     |         |
| cpd-platform-hp4w6                                              | 1/1 | Running |
| 0 39m                                                           |     |         |
| cpd-platform-operator-manager-6bc68dc8d-7xbjz                   | 1/1 | Running |
| 0 15m                                                           |     |         |
| cpdbr-tenant-service-6dcc49464c-zph7h                           | 1/1 | Running |
| 0 41m                                                           |     |         |
| create-postgres-license-config-m6ljg                            | 0/1 |         |
| Completed 0 32m                                                 |     |         |
| create-postgres-license-config-xdc4s                            | 0/1 |         |
| Completed 0 30m                                                 |     |         |
| e71db3df91177a0feccb558c266c053c60f349f28459e9ae7c2c55f685vzlzk | 0/1 |         |
| Completed 0 33m                                                 |     |         |
| ibm-common-service-operator-5f688bffdb-jzppd                    | 1/1 | Running |
| 0 15m                                                           |     |         |
| ibm-lakehouse-controller-manager-6c54bdbb6f-c6stq               | 1/1 | Running |
| 0 15m                                                           |     |         |
| ibm-namespace-scope-operator-66f4878bff-w9bt5                   | 1/1 | Running |
| 0 37m                                                           |     |         |
| ibm-watsonx-data-catalog-fvl6d                                  | 1/1 | Running |
| 0 38m                                                           |     |         |
| ibm-zen-operator-5646fffd6-bb95f                                | 1/1 | Running |
| 0 15m                                                           |     |         |
| meta-api-deploy-7bcbf6c896-nkw17                                | 1/1 | Running |
| 0 30m                                                           |     |         |
| opencloud-operators-8nw56                                       | 1/1 | Running |
| 0 37m                                                           |     |         |
| operand-deployment-lifecycle-manager-5f94f78-9rwln              | 1/1 | Running |
| 0 15m                                                           |     |         |
| postgresql-operator-controller-manager-1-18-5-6cb46bfd94-4v7m5  | 1/1 | Running |
| 0 15m                                                           |     |         |
| pre-zen-operand-config-job-lbj9l                                | 0/1 |         |
| Completed 0 29m                                                 |     |         |
| pre-zen-operand-config-job-n6x zr                               | 0/1 |         |
| Completed 0 30m                                                 |     |         |
| setup-job-ft4lg                                                 | 0/1 |         |
| Completed 0 15m                                                 |     |         |

---

To access the catalogs in IBM watsonx.data, the IBM watsonx.data service must be restarted by running the following command:

```
oc rollout restart sts,deploy -l 'component in
(ibm-lh-presto-coordinator,ibm-lh-presto,ibm-lh-hive-metastore)' -n cpd-instance
```

The command and its output are shown in Example 5-10.

*Example 5-10 Restarting IBM watsonx.data service*

---

```
$ oc rollout restart sts,deploy -l 'component in
(ibm-lh-presto-coordinator,ibm-lh-presto,ibm-lh-hive-metastore)' -n cpd-instance
Warning: would violate PodSecurity "restricted:v1.24": seccompProfile (pod or
container "ibm-lh-lakehouse-presto-01-presto" must set
securityContext.seccompProfile.type to "RuntimeDefault" or "Localhost")
statefulset.apps/ibm-lh-lakehouse-presto-01-presto restarted
Warning: would violate PodSecurity "restricted:v1.24": seccompProfile (pod or
container "ibm-lh-lakehouse-presto-01-presto-coordinator" must set
securityContext.seccompProfile.type to "RuntimeDefault" or "Localhost")
statefulset.apps/ibm-lh-lakehouse-presto-01-presto-coordinator restarted
Warning: would violate PodSecurity "restricted:v1.24": seccompProfile (pod or
container "ibm-lh-lakehouse-presto-01-presto-worker" must set
securityContext.seccompProfile.type to "RuntimeDefault" or "Localhost")
statefulset.apps/ibm-lh-lakehouse-presto-01-presto-worker restarted
Warning: would violate PodSecurity "restricted:v1.24": seccompProfile (pod or
container "ibm-lh-lakehouse-hive-metastore" must set
securityContext.seccompProfile.type to "RuntimeDefault" or "Localhost")
deployment.apps/ibm-lh-lakehouse-hive-metastore restarted
```

---

After the command finishes, verify that the watsonx.data service was successfully restarted by running the following command:

```
oc get deploy,sts -n cpd-instance
```

# Related publications

The publications listed in this section are considered particularly suitable for a more detailed discussion of the topics covered in this paper.

## IBM Redbooks

The following IBM Redbooks publications provide additional information about the topic in this document. Note that some publications referenced in this list might be available in softcopy only. For the current online list of Fusion Redbooks select [here](#).

- ▶ *IBM Storage Fusion HCI System: Metro Sync Disaster Recovery Use Case*, REDP-5708
- ▶ *IBM Storage Fusion Multicloud Object Gateway*, REDP-5718
- ▶ *IBM Storage Fusion Product Guide*, REDP-5688

You can search for, view, download or order these documents and other Redbooks, Redpapers, web docs, drafts, and additional materials, at the following website:

[ibm.com/redbooks](https://ibm.com/redbooks)

## Other publications

These publications are also relevant as further information sources:

## Online resources

- ▶ IBM Documentation for IBM Fusion 2.7.x  
<https://www.ibm.com/docs/en/sfhs/2.7.x>
- ▶ IBM Documentation for IBM watsonx.data  
<https://cloud.ibm.com/docs/watsonxdata>
- ▶ IBM Fusion  
<https://www.ibm.com/products/storage-fusion>
- ▶ IBM watsonx.data  
<https://www.ibm.com/products/watsonx-data>
- ▶ IBM watsonx.data together with IBM Storage Fusion HCI System ([video](#))

## Help from IBM

IBM Support and downloads

[ibm.com/support](https://ibm.com/support)

IBM Global Services

[ibm.com/services](https://ibm.com/services)





REDP-5720-00

ISBN 0738461458

Printed in U.S.A.

Get connected

