

Oracle® Cloud

Using Oracle AI Data Platform



G50054-33
September 2026



Oracle Cloud Using Oracle AI Data Platform,

G50054-33

Copyright © 2025, 2026, Oracle and/or its affiliates.

Primary Author: Adam Donald

This software and related documentation are provided under a license agreement containing restrictions on use and disclosure and are protected by intellectual property laws. Except as expressly permitted in your license agreement or allowed by law, you may not use, copy, reproduce, translate, broadcast, modify, license, transmit, distribute, exhibit, perform, publish, or display any part, in any form, or by any means. Reverse engineering, disassembly, or decompilation of this software, unless required by law for interoperability, is prohibited.

The information contained herein is subject to change without notice and is not warranted to be error-free. If you find any errors, please report them to us in writing.

If this is software, software documentation, data (as defined in the Federal Acquisition Regulation), or related documentation that is delivered to the U.S. Government or anyone licensing it on behalf of the U.S. Government, then the following notice is applicable:

U.S. GOVERNMENT END USERS: Oracle programs (including any operating system, integrated software, any programs embedded, installed, or activated on delivered hardware, and modifications of such programs) and Oracle computer documentation or other Oracle data delivered to or accessed by U.S. Government end users are "commercial computer software," "commercial computer software documentation," or "limited rights data" pursuant to the applicable Federal Acquisition Regulation and agency-specific supplemental regulations. As such, the use, reproduction, duplication, release, display, disclosure, modification, preparation of derivative works, and/or adaptation of i) Oracle programs (including any operating system, integrated software, any programs embedded, installed, or activated on delivered hardware, and modifications of such programs), ii) Oracle computer documentation and/or iii) other Oracle data, is subject to the rights and limitations specified in the license contained in the applicable contract. The terms governing the U.S. Government's use of Oracle cloud services are defined by the applicable contract for such services. No other rights are granted to the U.S. Government.

This software or hardware is developed for general use in a variety of information management applications. It is not developed or intended for use in any inherently dangerous applications, including applications that may create a risk of personal injury. If you use this software or hardware in dangerous applications, then you shall be responsible to take all appropriate fail-safe, backup, redundancy, and other measures to ensure its safe use. Oracle Corporation and its affiliates disclaim any liability for any damages caused by use of this software or hardware in dangerous applications.

Oracle®, Java, MySQL, and NetSuite are registered trademarks of Oracle and/or its affiliates. Other names may be trademarks of their respective owners.

Intel and Intel Inside are trademarks or registered trademarks of Intel Corporation. All SPARC trademarks are used under license and are trademarks or registered trademarks of SPARC International, Inc. AMD, Epyc, and the AMD logo are trademarks or registered trademarks of Advanced Micro Devices. UNIX is a registered trademark of The Open Group.

This software or hardware and documentation may provide access to or information about content, products, and services from third parties. Oracle Corporation and its affiliates are not responsible for and expressly disclaim all warranties of any kind with respect to third-party content, products, and services unless otherwise set forth in an applicable agreement between you and Oracle. Oracle Corporation and its affiliates will not be responsible for any loss, costs, or damages incurred due to your access to or use of third-party content, products, or services, except as set forth in an applicable agreement between you and Oracle.

For information about Oracle's commitment to accessibility, visit the Oracle Accessibility Program website at <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=docacc>.

Oracle customer access to and use of Oracle support services will be pursuant to the terms and conditions specified in their Oracle order for the applicable services.

Contents

About This Content

Part I Introduction to Oracle AI Data Platform Platform

1 Overview of Oracle AI Data Platform

What is Oracle AI Data Platform Used For?	1
Managed Integration with Open Source	1
Personas for Oracle AI Data Platform Users	2
Common Use Cases for Oracle AI Data Platform	2
About Oracle AI Data Platform Service Level Objective	3

2 Features of Oracle AI Data Platform

3 Get Started with Oracle AI Data Platform

Requirements for AI Features	1
Storage Encryption Options	2
Create Your First Oracle AI Data Platform	3
Access an Oracle AI Data Platform	5
Access Oracle AI Data Platform from a URL	6
Understanding the Workspace and UI Flow	6
Next Steps After Setup	6
Edit an Oracle AI Data Platform	7
Delete an Oracle AI Data Platform	7
Session Management	7

Part II Data Engineering

4 About Medallion Architecture

5 Ingest Data

Internal Sources	1
External Sources	3
Fusion Data in Oracle AI Data Platform with BICC	5
Create an Oracle Business Intelligence Cloud Connector Connection to Oracle AI Data Platform	5
Add Fusion Data Sources with a BICC Connection	6
Extract Fusion Data from a BICC Connection to a Notebook	6

6 Workspaces

About Workspaces	1
Create a Workspace	3
Create a Workspace with Private Network Access Enabled	3
Workspaces, External Catalogs, and Data Assets in a Private Network	5
Work with Files	5
Edit a Workspace	8
Delete a Workspace	8
Create a Folder	8
Delete a Folder	8
Move a Folder	8
Copy a Folder	8

7 Notebooks

Develop Code in Notebooks	1
Create a Notebook	1
Attach an Existing Cluster to a Notebook	2
Create a Cluster for a Notebook	2
Default Language	2
Widgets in Notebooks	2
Notebook Widget Examples	3
Widget Value Overrides	5
Widget Management	5
Manage Notebooks	6
Rename a Notebook	6
Delete a Notebook	6
Clone a Notebook	6
Browse Resources While Editing Notebook	6

Run Notebooks	7
Run Options for Notebook Cells	8
Run Code from a Notebook	9
Run Code from Another Notebook	10
Share Notebook Output with oidlUtils	10
Notebook Navigation	12
Create a Markdown Cell	12
Python Wheel Files	12
Update a Wheel Package	15
Notebook Output and Results	15
Download Notebook Output	16
Copy Notebook Output	16
Notebook Appearance	16
Manage Job Runs	18
Create a Manual Run Job from a Notebook	18
Create a Scheduled Job Run from a Notebook	18
Notebook Keyboard Shortcuts	19
Migrate Existing Apache Spark Code to Oracle AI Data Platform	19

8 Workflows

Configure Jobs	2
About Jobs	2
Create a Job	3
Change Job Location	4
Delete a Job	4
Schedule a Job Using a Calendar	4
Schedule a Job Using a Cron Expression	5
Run a Workflow on Demand	5
Run a Workflow Job from the Jobs Page	6
Change a Job Run Schedule	6
Pause or Activate a Job Run Schedule	6
Repair Failed Job Runs	6
Configure Tasks	7
About Tasks	7
Create a Python Task	8
Create a Notebook Task	9
Create a Nested Job Task	10
Create an If/Else Task	10
Create a Jar Task	11
Modify a Task	12
View Task Logs	12

Delete a Task	12
Parameterization	12
About Parameters	13
Pass Parameters Between Tasks and Notebook	15
Add Parameters to a Job	16
Add Parameters to a Python Task	16
Delete Parameters from a Job	17
Add Parameters to a Task	17
Delete Parameters from a Task	18
Run a Job with Different Parameter Values	18
Monitor Jobs	18
About Jobs Runs	19
View All Workflow Job Runs	20
View a Job's Run History	20
Monitor a Specific Job Run	20
Workflow Jobs: Run As	20
Edit Run As Identity for a Job	22

9 Compute

Manage Compute	1
About Compute Clusters	1
Create a Quickstart Cluster	3
Create a Custom Cluster	4
Create a Cluster Using an Imported Configuration	6
Create an NVIDIA GPU Cluster	6
Tune an NVIDIA GPU Cluster for Apache Spark	8
Export Compute Configuration	9
Import Compute Configuration	9
Modify a Cluster	9
Delete a Cluster	9
View Cluster Details	10
Maintenance Updates for Compute Clusters	10
Manage Libraries	10
Libraries	10
Install a Library from a Workspace or Volume	11
Install a Library from an Uploaded File	11
Uninstall a Library	12
Notebook Scoped Libraries	13
Monitor Compute	14
View Spark UI	14
View Driver and Worker Logs	14

View Metrics	14
View Event Logs	15
View Notebooks	16
AI Compute	17
Create an AI Cluster	21
Create an AI Cluster for an Agent	21
Attach an Existing AI Cluster to an Agent	22
Edit an AI Cluster	22
Detach an Agent from AI Compute	23
Delete an AI Cluster	23
Stop an AI Cluster	24
Start an AI Cluster	24
Restart an AI Cluster	25

10 Connect to Compute

Connections	1
Connect Oracle Analytics to Oracle AI Data Platform	2
Get an Oracle Analytics Connection Configuration File	2
Connect to Oracle Analytics	3
Download JDBC Driver	3
Configure DBeaver	4
Configure DBeaver with the Spark Simba JDBC Driver	4
Create a Database Connection in DBeaver	4
Connect DBeaver to Oracle AI Data Platform using JDBC	5
Download ODBC Driver	5

11 Streaming

About Streaming	1
Configuring Spark Structured Streaming using Workflows	4

Part III Integration

12 Oracle Cloud Infrastructure GoldenGate Integration

13 OCI Generative AI (Pretrained Foundation Models)

Part IV Machine Learning

14 Experiments

Create an Experiment	1
Edit an Experiment	2
View Experiment Run Details	2
Create an Experiment Run in a Notebook with Sample Code	2
Register a Model from Experiment Run Details	3
Delete Experiment	4

15 Models

Register a Model from an Existing Experiment Run	1
View Model Details	2
Edit Model Details	2
Delete a Model	2
Create a Model Deployment	3
Deploy a Model	3
Undeploy a Model	3
Send a Query to a Model Deployment Endpoint	4
Roll Forward a Model Deployment Version	4
Roll Back a Model Deployment Version	4
Edit a Model Deployment	5
Delete a Model Deployment	5
View Model Deployment Activity	5

16 Customer-managed MLflow Servers

Part V AI Agents

17 Agent Creation

Multi-agent Systems and Supervisor Patterns	1
About the Visual Flow Canvas	2
Create an Agent	4
Add Chat Trigger and Agent to Visual Builder Canvas	5
Configure a Supervisor Agent	6
Suggested Supervisor Instructions	7

Configure Supervisor Agent Memory and State Isolation	8
Models Parameter Tab	9
Add Guardrails to an Agent	10
Add Executor Agents and Tools to an Agent	11
Executor Agent Configuration	12
Checklist for Agents through Visual Builder	15
Agents Through Code	16
Build an Agent Through Code by Upload	18
Build an Agent Through Code by Creating New Code	19
Set an Entry File for Agents through Code	20
Set a Dependency File for Agents through Code	20
Test Agent Code	21
Agent Skills in Coding Experience	22
Add a New Skill	29
Add a New Executable Capability to an Existing Skill	29
Agent Skills Troubleshooting	30
Example: Complete Agent Skill	31
Agent Testing	31
Test your Agents in the Playground	34
Create an Agent Test Session	34
Resume an Agent Test Session	35
Delete an Agent Test Session	36

18 About Agent Tools

Custom Tool	3
Add a Custom Tool to an Agent	10
Remote MCP Server Tool	11
Connect an Agent to a Remote MCP Server from the Visual Builder	21
HTTP Request Tool	22
Add an HTTP Request Tool to an Agent	28
Prompt Tool	29
Add a Prompt Tool to an Agent	38
RAG Tool	38
Add a RAG Tool to an Agent	41
SQL Tool	42
Add a SQL Tool to an Agent	50

19 Knowledge Bases

Create a Knowledge Base	2
Edit a Knowledge Base	3

	Delete a Knowledge Base	3
	Add a Data Source to a Knowledge Base	4
	Ingest Data to a Knowledge Base	5
	View Ingestion Job Run Status	5
	Delete a Data Source	5
20	Agent Memory	
21	Session Variables	
	Create a Session Variable in the Agent Variables Tab	7
	Refer to Session Variables in Agent Flow Instructions	8
	View Agents and Tools using a Session Variable	8
	Assign Values to Session Variables from the Playground	8
22	Guardrails	
	Enable Specific Guardrails in a Node	4
23	Agent Deployment	
	Deploy an Agent	2
	Deploy an Agent with OAuth2	2
	Undeploy an Agent	3
	A2A Agent Deployment	3
	Edit a Draft Agent Card	11
	Edit a Published Agent Card	13
24	Invoke a Deployed Agent	
	Methods to Invoke Endpoint URIs	1
25	Monitor Agents	
	View Agent Sessions	2
	View Agent Metrics	2
26	Agent Operations	
	Move an Agent	1
	Copy an Agent	1

Part VI Data Governance

27 Manage with Master Catalog

Master Catalog	1
Use Cases for Master Catalog	1
Cross-Tenancy External Tables and Volumes	2
Standard Catalogs	3
Create a Standard Catalog	4
Edit a Standard Catalog Description	4
Rename a Standard Catalog	4
Delete a Standard Catalog	5
External Catalogs	5
Create an External Catalog	7
Create an External Catalog for Private Networks	7
Refresh External Data Catalogs	8
Refresh External Catalogs with SQL	9
Edit an External Catalog Description	9
Rename an External Catalog	10
Edit an External Catalog Configuration	10
Delete an External Catalog	10
How to Delete Data from Tables in an External Catalog	11
Schema	12
Create a Schema	12
View Schema Details	13
Edit a Schema Description	13
Delete a Schema	13
Tables	14
Schema Evolution	17
Create a Managed Table	18
Create an External Table	19
Edit a Table	20
View Table Details	20
Delete a Table	21
Volumes	21
Create a Managed Volume	21
Create an External Volume	22
Edit a Volume	23
View Volume Details	23

Delete a Volume	23
Synonyms	24
View Synonym Details	25

28 Lineage (Preview)

Entity and Column Lineage (Preview)	6
View Data Lineage (Preview)	9
View Lineage for Specific Data Columns (Preview)	10
View Details for a Lineage Artifact (Preview)	10
Export Impact Analysis (Preview)	10
Filter Lineage Flow Diagram (Preview)	11
Search for Artifacts in Lineage Flow Diagram (Preview)	11
Change Lineage Flow Depth (Preview)	11
Share a Lineage Flow Diagram (Preview)	12

29 Credential Store (Preview)

About Credential Store (Preview)	1
Create Credentials (Preview)	3
Use Stored Credentials in a Notebook (Preview)	3
Modify Credentials (Preview)	4
Share Credentials (Preview)	4
Delete Credentials (Preview)	4
View Credential Details (Preview)	4
View Credential Store History (Preview)	5

30 Audit Log

Search and Filter Audit Logs	2
------------------------------	---

31 Auto Populate Catalog

About Auto Populate	1
Create Metadata Extractor	2
Manually Review Extracted Metadata Entities	3
View Reviewed Entities	3
View Metadata Extractor Details	3
Delete Metadata Extractor	3

32 Share Data

Data Sharing	1
Create a Share	3
Modify a Share	3
Delete a Share	3
View Share Details	3
Add an Asset to a Share	3
Remove an Asset from a Share	4
Add Recipients to Data Sharing	4
Add an Existing Recipient to a Share	4
Manually Activate a Recipient	4
Resend Activation Link to a Recipient	5
Remove Recipient from a Share	5
Modify a Recipient	5
View Recipient Details	5
Delete a Recipient	5

Part VII DevOps

33 Git Integration (Preview)

Git Folders in Oracle AI Data Platform (Preview)	1
Create a Git Folder (Preview)	1
Edit Git Folder (Preview)	2
Delete Git Folder (Preview)	2
Modify Git Folder Settings (Preview)	2
View Git Folder History (Preview)	2
Create Git Branch (Preview)	2
Delete Git Branch (Preview)	3
Git Pull (Preview)	3
Git Push (Preview)	3
Merge Git Branches into Main (Preview)	4
Git Rebase (Preview)	4
Git Reset (Preview)	5

34 Bundles and CI/CD (Preview)

Create a Bundle (Preview)	2
What Gets Generated After Bundle Creation (Preview)	3
Publish Path Configuration (Preview)	3
How Bundle Artifacts are Published (Preview)	5

Deploy a Bundle (Preview)	5
How Bundles are Deployed to Published Paths (Preview)	5
Sync a Bundle (Preview)	6
Recommended Promotion Workflow for Bundles (Preview)	6
Bundle Variables (Preview)	7
Named Variables in Environments (Preview)	12
Bundle Variables and Credential Store (Preview)	13
Purge a Bundle (Preview)	15
Purge Behavior for Published Artifacts (Preview)	15
Existing Bundles and Publish Paths (Preview)	16
Troubleshooting Bundles (Preview)	16

35 Oracle AI Data Platform (AIDP) Developer Extension

Install Oracle AI Data Platform (AIDP) Developer Extension on Visual Studio Code	2
Install Oracle AI Data Platform (AIDP) Developer Extension Manually	2
Set Up Oracle AI Data Platform (AIDP) Developer Extension in Visual Studio Code	2

Part VIII Security

36 Permissions Model

About Permissions	1
Workspace Permissions	3
Create Workspace Permissions	3
Modify Workspace Permissions	4
Delete Workspace Permissions	4
Workspace Folder Permissions	4
Workspace File Permissions	5
Create File and Folder Permissions	7
Modify File and Folder Permissions	7
Delete File and Folder Permissions	7
Compute Cluster Permissions	7
Create Cluster Permissions	8
Modify Cluster Permissions	8
Delete Cluster Permissions	9
Job Permissions	9
Create Job Permissions	10
Modify Job Permissions	10
Delete Job Permissions	10
Notebook Permissions	10
Create Notebook Permissions	11

Modify Notebook Permissions	11
Delete Notebook Permissions	12
Master Catalog Permissions	12
Create Master Catalog Permissions	13
Modify Master Catalog Permissions	13
Delete Master Catalog Permissions	13
Standard Catalog Permissions	13
External Catalog Permissions	14
Create Catalog Permissions	15
Modify Catalog Permissions	15
Delete Catalog Permissions	16
Schema Permissions	16
Create Schema Permissions	17
Modify Schema Permissions	17
Delete Schema Permissions	18
Table Permissions	18
Create Table Permissions	19
Modify Table Permissions	19
Delete Table Permissions	19
Volume Permissions	19
Create Volume Permissions	20
Modify Volume Permissions	20
Delete Volume Permissions	20
Agent Permissions	21
Create Agent Permissions	22
Modify Agent Permissions	22
Delete Agent Permissions	22
AI Compute Cluster Permissions	22
Create AI Cluster Permissions	23
Modify AI Cluster Permissions	23
Delete AI Cluster Permissions	24
Knowledge Base Permissions	24
Create Knowledge Base Permissions	24
Modify Knowledge Base Permissions	25
Delete Knowledge Base Permissions	25
Git Folder Permissions (Preview)	25
Create Git Folder Permissions (Preview)	26
Modify Git Folder Permissions (Preview)	26
Delete Git Folder Permissions (Preview)	26
Bundle Permissions (Preview)	27
Create Bundle Permissions (Preview)	27
Modify Git Bundle Permissions (Preview)	28

Delete Bundle Permissions (Preview)	28
Credential Permissions (Preview)	28
Create Credential Permissions (Preview)	28
Modify Credential Permissions (Preview)	29
Delete Credential Permissions (Preview)	29

37 IAM Policies for Oracle AI Data Platform

38 Roles

About Roles	1
Map Active Directory Groups to IAM Groups	2
Create a Role	2
Modify a Role	2
Delete a Role	2
Assign Members to a Role	3
Remove Member from Role	3

39 Network Sources

Part IX Administration

40 Pricing

41 Administration Settings (Preview)

Git Linked Accounts (Preview)	1
Add Git Credentials by Using Personal Access Tokens (Preview)	1
Set Default Git Credentials (Preview)	1
Edit Git Credentials (Preview)	2
Delete Git Credentials (Preview)	2

42 AI Feature Management

Create an Oracle Autonomous AI Lakehouse Instance to Enable AI	2
Choose an Existing Oracle Autonomous AI Lakehouse Instance to Enable AI	2

43 Notifications

44 Limits

Part X Resources

45 Use APIs, SDK, and CLI

REST APIs	1
SDK and CLI	1

46 Reference

SQL Grammar	1
Catalog SQL Grammar	1
Schema SQL Grammar	6
Volume SQL Grammar	8
Table SQL Grammar	10
DML Queries	14
Aidputils for Agents and Tools Sample Code	14
Agent without Tools	15
SQL Tool Test	17
Prompt (LLM) Tool Test	17
Custom Code Tool - Hello World	18
Custom Code Tool - Developer Toolkit	19
Agent with Tool Registration in Oracle AI Data Platform	23
Observability: Logging, Tracing, and Metrics	25
Agent Instantiation and Usage with Aidputil Packages	26
Guardrails Configuration	30
Agent LangGraph Code Samples	33
Hello World with LangGraph	34
ReAct Agent with Prompt Tool LangGraph Sample Code	34
ReAct Agent with RAG Tool LangGraph Sample Code	37
ReAct Agent with SQL Tool LangGraph Sample Code	39
ReAct Agent with Custom Tools LangGraph Sample Code	41
Multi-agent System – Supervisor Pattern Sample Code	44

47 Known Issues

About This Content

This guide describes how to create, access, manage, and develop data resources in Oracle AI Data Platform.

Audience

This guide is intended for data scientists and developers who use AI Data Platform to:

- Create secure, single-pane of glass data platforms for data governance and management.
- Create coding solutions for managing, analyzing, and enriching data

Conventions

The following text conventions are used in this document:

Convention	Meaning
boldface	Boldface type indicates graphical user interface elements associated with an action, or terms defined in text or the glossary.
<i>italic</i>	Italic type indicates book titles, emphasis, or placeholder variables for which you supply particular values.
monospace	Monospace type indicates commands within a paragraph, URLs, code in examples, text that appears on the screen, or text that you enter.

Part I

Introduction to Oracle AI Data Platform

Oracle AI Data Platform simplifies cataloging, ingesting, and analyzing data for data professionals in an organization. The Oracle AI Data Platform service provides the platform and the framework to create data analytics pipelines.

Users are able to create an AI Data Platform, build data catalogs across data in their data estate and set role-based access control (RBAC) policies, and use notebooks powered by Spark to prepare, analyze, and enrich their data. The data catalog enables seamless usage and ingestion of metadata from supported external services. OCI and Oracle Analytics users can connect to AI Data Platform and analyze their data. All these capabilities are provided in a single pane of glass experience with enterprise security features across all capabilities.

Oracle AI Data Platform is an enterprise-grade, unified platform that simplifies the cataloging, preparation, and analysis of data across an organization's data estate. It brings together data, AI, analytics, and governance services within a single, cohesive experience enabling users to build and operationalize AI-powered applications securely and at scale. Oracle AI Data Platform unifies Oracle Autonomous AI Lakehouse, Oracle Analytics, OCI Object Storage, OCI Generative AI and Oracle Fusion Data Intelligence into a single, governed platform to build, deploy, and scale data and AI applications on your enterprise data.

Within this platform, your **Oracle AI Data Platform instance** provides a dedicated development environment for data engineers, scientists, and AI developers to design, orchestrate, and deploy data pipelines, models, set role-based access control (RBAC) policies, and use open source technologies such as Spark to prepare, analyze, and enrich their data.

Topics:

- [Overview of Oracle AI Data Platform](#)
- [Features of Oracle AI Data Platform](#)
- [Get Started with Oracle AI Data Platform](#)

1

Overview of Oracle AI Data Platform

This chapter provides information and procedures for new users getting started with Oracle AI Data Platform.

Topics:

- [What is Oracle AI Data Platform Used For?](#)
- [Managed Integration with Open Source](#)
- [Personas for Oracle AI Data Platform Users](#)
- [Common Use Cases for Oracle AI Data Platform](#)
- [#unique_21](#)

What is Oracle AI Data Platform Used For?

Oracle AI Data Platform provides an integrated environment for building, orchestrating, and operationalizing data and AI workflows.

Oracle AI Data Platform is designed for enterprises that need to:

- **Streamline Data Discovery and Governance:** AI Data Platform provides a centralized metadata repository (Master Catalog) that enhances searchability and governance of structured and unstructured data.
- **Enable Secure Data Collaboration:** Through RBAC-based access control, AI Data Platform allows different teams to work on shared datasets while maintaining strict security policies.
- **Accelerate Data Preparation and Processing:** With built-in notebooks and workflow orchestration, users can clean, transform, and enrich data efficiently.
- **Support Advanced Analytics, AI, and ML:** Users can build and deploy customizable AI agents to support their work with cataloged data. AI Data Platform also integrates with Apache Spark, allowing data scientists and analysts to run complex computations and model training directly within their data.
- **Ensure Seamless Integration Across Data Sources:** AI Data Platform supports external catalogs from Autonomous Database (ADB), Object Storage (OS), and third-party data sources, enabling users to query and analyze data without duplication.

Managed Integration with Open Source

Oracle AI Data Platform leverages and extends open-source technologies to provide a powerful yet managed experience.

Some key integrations include:

- **Apache Spark:** AI Data Platform's compute layer is powered by Spark, enabling scalable, distributed data processing.
- **Delta Lake Support:** AI Data Platform leverages Delta Lake to enhance data reliability, ACID transactions, and schema evolution.

- **Iceberg & Hudi Compatibility via Delta Uniform:** Through Delta Uniform, AI Data Platform extends support for Apache Iceberg and Apache Hudi, enabling interoperability across different storage formats. This ensures users can adopt a unified table format strategy while maintaining efficient query execution and data governance.
- **JDBC Integration for BI Tools:** AI Data Platform provides JDBC drivers, allowing seamless connectivity with external BI tools like Oracle Analytics Cloud (OAC) and third-party visualization platforms.

Personas for Oracle AI Data Platform Users

Oracle AI Data Platform serves a variety of users across different roles within an organization, each with unique needs and requirements.

Here's a general overview of the key personas who interact with AI Data Platforms:

- **Data Engineers** - Data engineers work with large-scale data pipelines, transforming raw data into usable formats for analysis. They rely on the robust capabilities of AI Data Platform to design and manage data workflows, ingest data from various sources, and ensure data quality. They are highly focused on automating processes, optimizing compute resources, and integrating different data systems seamlessly.
- **AI Engineers** - AI engineers build end-to-end flows for AI-powered agents. They build agents that can interface with supported tools, such as RAG, SQL, and prompts to assist users in streamlining data development and management.
- **Data Analysts** - Data analysts use AI Data Platform to discover, analyze, and generate insights from data. They require an intuitive interface and tools for querying and analyzing large datasets. AI Data Platform empowers them with interactive notebooks and seamless integration with business intelligence (BI) tools, helping them transform raw data into actionable insights for decision-makers.
- **Data Scientists** - Data scientists leverage the scalable compute capabilities of AI Data Platform for machine learning and advanced analytics tasks. They need access to diverse datasets, powerful processing tools, and the ability to run complex models. Spark-powered notebooks, AI/ML integration, and support for open-source libraries enable data scientists to build, test, and deploy models within the platform.
- **Data Stewards** - These users ensure that all data is handled in compliance with industry regulations and organizational policies. They focus on maintaining data privacy, auditing access, and monitoring data usage across the organization. AI Data Platform helps them manage metadata, enforce role-based access controls (RBAC), and ensure proper governance through cataloging, lineage tracking, and security policies.

Common Use Cases for Oracle AI Data Platform

Oracle AI Data Platform serves a variety of use cases across industries and business functions.

Medallion Architecture

- Implement a Medallion Architecture with bronze, silver, and gold layers.
- Use Delta Uniform and Iceberg for efficient data storage and query optimization.
- Enable zero-copy access to external data sources for seamless analytics.

ETL & Data Engineering

- Use Spark-based workflows and notebooks to process, transform, and enrich raw data.

- Automate data pipelines with low-code/no-code workflow orchestration.
- Handle large-scale batch processing and real-time data ingestion.

Machine Learning, AI and Data Science

- Train and deploy machine learning models using Spark-powered notebooks.
- Enable large-scale feature engineering and data transformation.
- Provide managed execution environments for Python and PySpark workloads.

Building AI Agents Leveraging Enterprise Data

- Create conversational AI agents to assist with retrieving and developing data.

Enterprise Data Catalog & Governance, Delta Sharing

- Centralized metadata management for structured and unstructured data.
- Role-based access control (RBAC) for secure data access and collaboration.
- Integration with external catalogs, including Oracle Autonomous AI Database and OCI Object Storage.
- AI Data Platform supports Delta Sharing, enabling secure, real-time, and governed data sharing across organizational boundaries.

Analytics, Business Intelligence & Reporting

- Connect OCI, Oracle Analytics Cloud (OAC), and third-party BI tools via JDBC like Tableau, Power BI.

Multi-Cloud & Hybrid Data Integration

- Enable federated query execution across multiple OCI services.
- Integrate with third-party cloud storage and databases for hybrid analytics.
- Maintain data sovereignty and compliance across multiple environments.

About Oracle AI Data Platform Service Level Objective

Oracle AI Data Platform is built on Oracle Cloud Infrastructure and leverages the Oracle PaaS Cloud services.

Oracle AI Data Platform availability is tied to dependent Oracle services such as Oracle Identity Cloud Service, Oracle Autonomous AI Lakehouse, and Oracle Cloud IaaS services, including compute and storage. Each of these services has published service level agreements under [Oracle Cloud SLA](#). See the Oracle PaaS and IaaS Public Cloud Services Pillar Document (PDF).

Oracle AI Data Platform has a service-level objective (SLO) of 99.95%.

Use the following definitions to calculate the Service Level for AI Data Platform.

Term	Definition
Service Level	Oracle measures the Service Level over the immediately preceding calendar month by: • Subtracting the sum of the Error Rate for each hour of that month (the Error Rate Sum) from 100. • Dividing the Error Rate Sum by the total number of hours in that month. • Multiplying the result by 100 to calculate the percentage.
Error Rate	The total number of Failed Service REST API Calls in a one-hour interval during the measured month, divided by the total number of Service REST API Calls during that interval.
Service REST API Call	An HTTP request that fulfills the applicable Cloud Service's REST API specification.
Failed Service REST API Call	A Cloud Service REST API Call processed by Your User that results in a 5xx (Server Error) status code.

Features of Oracle AI Data Platform

Oracle AI Data Platform is a modern data platform designed to simplify data ingestion, processing, and analytics at scale. It provides a seamless integration of compute, storage, and cataloging capabilities to enable efficient data management.

Key features of AI Data Platform include:

Workspace

A workspace in AI Data Platform acts as an isolated environment where users can manage and organize their data lake resources, including workflows, notebooks, and libraries. Workspaces enable efficient collaboration and governance by keeping resources grouped logically.

Compute

AI Data Platform provides scalable CPU and GPU compute resources for executing data processing and analytics workloads. Users can leverage Spark-based execution environments for high-performance processing, supporting batch and interactive workloads.

AI Compute

AI compute provides scalable resources for agents, both for testing and production. Once an agent is attached to an AI compute, agent developers can create sessions and start interacting with their agent. Multiple agents can be either attached (for testing) or deployed (for production) on a given AI compute.

Agents

Agents are end-to-end agentic applications. Agents are defined through a graph of steps represented by different types of nodes, such as agents and tools. Agents are defined through a no-code visual flow builder and through code via third-party libraries, such as LangGraph.

Notebook

AI Data Platform includes notebooks as an interactive development environment for writing and executing code. It supports Python and SparkSQL enabling users to transform, analyze, and visualize data directly within AI Data Platform.

Workflow

The workflow component allows users to define and orchestrate data pipelines made of notebooks, Python tasks, if-else, and other job tasks. Users can create, schedule, and monitor workflows for ETL, data transformations, and analytics automation.

Master Catalog

The Master Catalog serves as the central metadata repository for all structured and unstructured datasets within an AI Data Platform. It provides unified governance and data discovery, allowing users to search and manage datasets across different schemas and storage locations.

Catalog

A catalog in an AI Data Platform is a logical grouping of schemas, tables, volumes, and models, providing a structured way to organize datasets. Users can create multiple catalogs for different projects or teams to ensure effective data segmentation.

Knowledge Bases

Knowledge bases leverage Oracle AI Database 23ai Vector Search capability to store vector embeddings from documents stored in an AI Data Platform, such as PDF, DOCX, and HTML files. Through Oracle AI Database 23ai's vector search capabilities, knowledge bases empower AI agents to perform semantic searches and retrieve semantically relevant documents.

Schema

A schema defines the structure within a catalog, organizing tables and views under a common namespace. Schemas help in logically structuring data for different applications and analytics workloads.

Table

A table in an AI Data Platform represents structured datasets that can be queried and processed. Tables support various storage formats, including Delta Uniform, ensuring compatibility with multiple query engines.

View

A view is a virtual table in an AI Data Platform that provides a queryable representation of data stored in underlying tables. Views allow for simplified access to transformed datasets without requiring data duplication.

Volume

A volume is a storage abstraction in an AI Data Platform that provides a managed space for persisting raw, processed, and curated data. It supports efficient data access and integration with Object Storage.

Auto Populate

The Auto Populate feature simplifies metadata management by automatically detecting and registering new datasets located in OCI Object Storage. This reduces manual effort in keeping data catalogs up to date.

Role-Based Access Controls (RBAC)

AI Data Platform implements RBAC to enforce fine-grained access control across different resources. Users can define roles and permissions for workspaces, catalogs, datasets, agents, and knowledge bases to ensure secure collaboration.

Audit Log

Audit logs in AI Data Platform capture detailed records of user activities. These logs help monitor usage, ensure compliance, and investigate issues such as unauthorized access or configuration changes.

Three-Part Namespace

AI Data Platform adopts a three-part namespace (Catalog.Schema.Table) for accessing datasets, enabling a structured and consistent way to reference data across the platform. This standardization improves interoperability and ease of access.

3

Get Started with Oracle AI Data Platform

Oracle AI Data Platform helps you quickly build, orchestrate, and manage end-to-end data workflows across your organization. This guide walks you through the essential steps to create your first AI Data Platform and start working with data in a governed and scalable environment.

Topics:

- [Requirements for AI Features](#)
- [Storage Encryption Options](#)
- [Create Your First Oracle AI Data Platform](#)
- [Access an Oracle AI Data Platform](#)
- [Access Oracle AI Data Platform from a URL](#)
- [Understanding the Workspace and UI Flow](#)
- [Next Steps After Setup](#)
- [Edit an Oracle AI Data Platform](#)
- [Delete an Oracle AI Data Platform](#)
- [Session Management](#)
- [IAM Policies for Oracle AI Data Platform](#)
- [Pricing](#)

Requirements for AI Features

In order to use AI features in Oracle AI Data Platform, you need to provision an Oracle Autonomous AI Lakehouse instance.

As part of creating your AI Data Platform instance, an Oracle Autonomous AI Lakehouse must be created and configured. You can choose to create the Autonomous AI Lakehouse instance as part of creating your first AI Data Platform instance or choose to use an existing Autonomous AI Lakehouse instance. If you are using an existing Autonomous AI Lakehouse, you must also provide the password for the ADMIN user of that instance.

Note

Existing Autonomous AI Lakehouse instances must be at least 26ai or above. When you create a new Autonomous AI Lakehouse instance as part of creating your AI Data Platform instance, it is automatically configured as 26ai or above.

The Autonomous AI Lakehouse instance serves as the foundational data layer for AI workloads within AI Data Platform and is used as the default vector store for all AI-native capabilities. The Autonomous AI Lakehouse instance requires Oracle Autonomous AI Lakehouse version 26ai or later and is mandatory for enabling AI features in AI Data Platform, including:

- Knowledge Bases (KB),
- Agent Flows,
- and AI Compute.

Without an Autonomous AI Lakehouse instance, AI Data Platform continues to support structured data ingestion and transformation, but AI-native capabilities remain disabled.

Note

If your tenancy has reached one or more Autonomous AI Lakehouse quotas, you can't create a new Autonomous AI Lakehouse instance when creating your AI Data Platform instance. In that case, you must either select an existing Autonomous AI Lakehouse instance or contact your support team to request a quota increase.

You must have the following IAM policies enabled for AI features:

Allow Oracle AI Data Platform service to call LLMs hosted on the OCI GenAI service:

```
allow any-user to use generative-ai-family in tenancy where all
{ request.principal.type='aidataplatfom' }
```

Note

Depending on when your AI Data Platform was created, your principal type may be 'datalake' instead of 'aidataplatfom'. You can check your principal type by viewing the details of your AI Data Platform OCID.

Allow Oracle AI Data Platform service to provision and manage Oracle Autonomous AI Lakehouse resources to enable AI features:

```
allow group <aidpAdminGroup> to manage autonomous-databases in <aidp
compartment>
```

Storage Encryption Options

You can use Oracle-managed encryption keys, or you can choose your own customer-managed key (CMK) from OCI Vault when your organization requires direct control over key lifecycle and rotation.

Oracle AI Data Platform encrypts data at rest by default. When you configure your first AI Data Platform instance, you can choose to keep the default Oracle-managed encryption keys, or you can choose your own CMK from OCI Vault when your organization requires direct control over key lifecycle and rotation. Use Oracle-managed keys when you want the simplest setup and don't need to manage key material yourself. Use customer-managed keys when your security policy requires you to control the vault, key rotation process, access boundaries, or audit trail for the encryption key.

Customer-Managed Key Prerequisites (Preview)

Before you select a customer-managed key, make sure the following prerequisites are in place:

- You have permission to view OCI Vaults and keys in the target compartment. For more information configuring vaults and keys in OCI, see [Configuring IAM Policies](#).
- You have an OCI Vault available in the same region as your AI Data Platform deployment.
- You have a master encryption key in that vault.
- The key is enabled and available for use.
- You must choose an AES symmetric master key with a 256-bit length.

Default Behavior – Oracle Managed Keys

If you don't select a customer-managed key, AI Data Platform uses Oracle-managed keys automatically. This is the default option and requires no additional vault configuration. Choose this option when:

- You want the fastest provisioning path.
- Your team doesn't need to manage or rotate the encryption key directly.
- Your security policy allows Oracle-managed encryption for service data at rest.

Considerations Before Choosing Customer-Managed Keys

Customer-managed keys provide more control, but they also add operational responsibility.

- If the vault or key becomes unavailable, encryption-dependent operations can fail.
- If the key is disabled, deleted, or scheduled for deletion, access to protected data can be impacted.
- Key rotation, compartment design, and IAM policies become part of your operational runbook.
- Provisioning or update flows can take longer than the default Oracle-managed path.

Troubleshooting

If you don't see the customer-managed key option when setting up your AI Data Platform instance, check the following:

- Confirm that the feature is enabled in your tenancy.
- Confirm that you have permission to read vaults and keys.
- Confirm that the vault and key are in the same region as the AI Data Platform instance.

Limitations

You cannot change your encryption keys after creating an AI Data Platform instance. In order to change encryption strategy, you must create a new AI Data Platform instance and move the code to that.

Create Your First Oracle AI Data Platform

Before you create your first Oracle AI Data Platform, you need to ensure you have the correct IAM permissions.



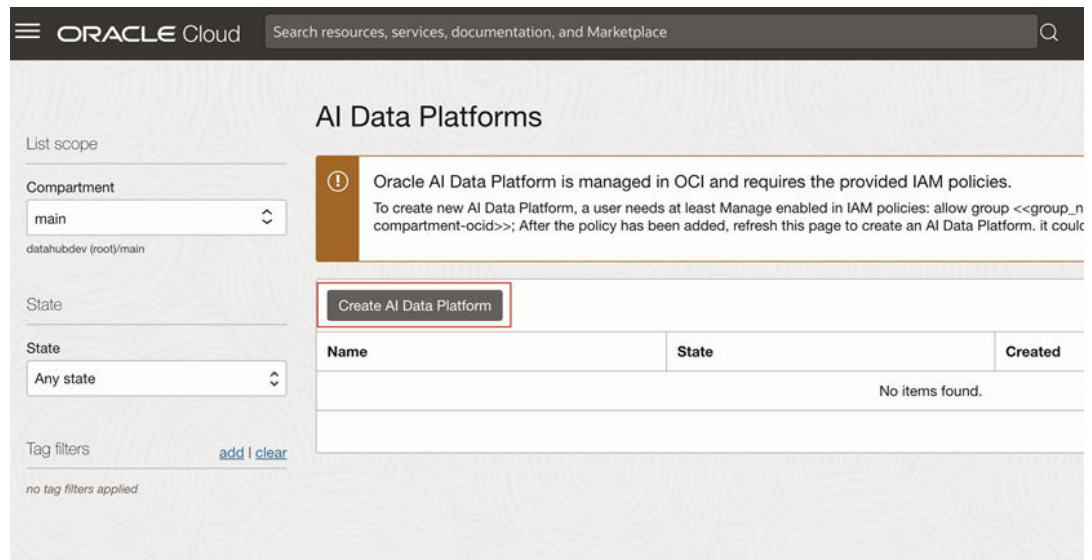
[LiveLabs Sprint](#)



[Video](#)

Before creating or using an AI Data Platform, check the following:

- You have the appropriate permissions for the compartment where you want to create the AI Data Platform. For more information, see [Managing Compartments](#).
 - If you already have tag-defaults set for your compartment you may face issues while creating an AI Data Platform instance in that scope. For more information, see [Known Issues](#).
- In the Home section of OCI, click **Analytics & AI**.
 - Click **AI Data Platform**.



- Click **Create AI Data Platform**. AI Data Platform performs a check to validate that all required IAM policies are in place.
- Provide a name and description for your AI Data Platform instance.
- Provide a name and description for your first AI Data Platform workspace.
- Select whether your AI Data Platform instance uses a public or private network configuration.
 - For private networks, you choose a compartment and VCN where Private Service Access (PSA) is already enabled. You then select a subnet and one or more network security groups.
 - If your data sources are in a different private network, you need to select that option and provide the VCN, subnet, and network security groups for that private network as well.
 - If your data sources require access through a SCAN proxy, require advanced DNS configuration, or your private network enforces a network source, expand **Advanced – optional**, select the appropriate options and provide the required settings.
- An Oracle Autonomous AI Lakehouse instance is required for AI features. Select whether to create a new instance or use an existing instance.
 - For **Create new**, provide a new password for the ADMIN user of the instance.
 - For **Choose existing**, select the compartment and instance and enter the password for the ADMIN user of the existing instance. Your existing instance must be 26ai or above.

8. To associate an Oracle Analytics Cloud instance with your AI Data Platform, select whether to use an existing Analytics Cloud instance or create a new one at the same time as your AI Data Platform instance.
 - a. For **Choose existing**, select the compartment, provide a domain name, and select the existing instance. Provide the vault secret OCID as the password.
 - b. For **Create new**, provide the domain name and vault secret OCID.
 - c. Select **None** to associate an Analytics Cloud instance later.
9. Select whether encryption keys are Oracle-managed (default setting) or customer-managed.
 - **Oracle-managed key**: This is the default option and requires no additional vault configuration.
 - **Customer-managed key**: Select if your organization requires direct control over key lifecycle and rotation. You need to provide the **Vault OCID** and **Key OCID** for your encryption key. Click **Validate** to confirm the provided OCIDs are valid.

Storage encryption

Choose a master encryption key for the bucket.

Key selection

☐ Oracle-managed key
Oracle manages the encryption key.

☒ Customer-managed key
Managed by you. You must have a valid encryption key. To create a key use Vault.

Vault OCID
Enter vault OCID

Key OCID
Enter key OCID

Validate

10. Select one of the following from **Add Policies**:
 - **Standard** (Recommended): Applies access settings broadly at the tenancy level.
 - **Advanced**: Allows you to configure fine-grained access at the compartment level.
11. Review the missing policies, if any, and apply them as necessary.
12. Expand **Optional Policies**. Review and add any additional policies your instance requires. For more information, see [IAM Policies for Oracle AI Data Platform](#).
13. Click **Advanced Options** to provide metadata tags for your AI Data Platform. Click **Add tag** to add multiple.
14. Click **Create AI Data Platform**.

Access an Oracle AI Data Platform

You can access any Oracle AI Data Platform that you or a role you belong to have been provided permissions for.

You require at least USE permissions from the AI Data Platform IAM policies to access an AI Data Platform. For more information, see [IAM Policies for Oracle AI Data Platform](#).

1. In the Home section of OCI, click **Analytics & AI**.
2. Click **AI Data Platform**.
3. Click the AI Data Platform you want to access.

Access Oracle AI Data Platform from a URL

You can access your Oracle AI Data Platform instance directly from a URL.

The URL for your Oracle AI Data Platform uses the following format, where *tenant_name* is the name of the OCI account you use to log in and *domain_name* is the domain you select:

```
https://<hash>.datalake.oci.oraclecloud.com/#?  
&tenant=<<tenant_name>>&domain=<<domain_name>>
```

You can also get a direct link to your instance by right-clicking your AI Data Platform instance name on the home page and clicking **Copy Link**.

You can share your AI Data Platform instance URL with other users. Those users need to have appropriate permissions to access the instance. For more information, see [About Permissions](#).

Understanding the Workspace and UI Flow

Once you've created or accessed your Oracle AI Data Platform, the user interface helps you navigate and manage all key components seamlessly inside the instance.

Left Navigation Overview

When inside an AI Data Platform, the left sidebar contains:

- Home - Brings you back to the dashboard
- Master Catalog - View all registered data assets across workspaces
- Workspace - Direct access to workspace details
- Select Workspace - Drop-down to switch between different workspaces
 - Inside a workspace, you'll see:
 - * Workflow - Design and schedule data workflows
 - * Compute - Manage Spark compute environments
- Data Sharing - Share data with other services supporting Delta Share protocol
- Auto populate catalog - Automatically ingest metadata from connected sources
- Notifications - Search and filter system messages you receive
- Roles - Configure access control (RBAC)
- Audit logs - Search and filter the history of objects in your AI Data Platform

Next Steps After Setup

Once you have created your Oracle AI Data Platform, you can prepare it for regular use.

When you have your AI Data Platform created and have accessed it for the first time, these are the next steps to ensure its ready for use:

- Review IAM Policies and Roles - Ensure correct access is set up for users and groups via the Roles tab. For more information, see [Roles](#).

- Set Up Compute - Configure Spark compute in the Compute tab to run your workflows efficiently. For more information, see [Compute](#).
- Ingest and Organize Data - Start by creating a Catalog to organize your data sources. For more information, see [Data Management](#).
- Explore Notebooks - Use notebooks for ad hoc exploration, transformations, or machine learning tasks. For more information, see [Notebooks](#).

Edit an Oracle AI Data Platform

You can edit the details of an Oracle AI Data Platform instance you own.

1. In the Home section of OCI, click **Analytics & AI**.
2. Click **AI Data Platform**.
3. Click **Actions** next to your AI Data Platform and click **Edit**.
4. Modify the details and click **Save**.

Delete an Oracle AI Data Platform

You can delete unused or redundant Oracle AI Data Platform instances you own.

1. In the Home section of OCI, click **Analytics & AI**.
2. Click **AI Data Platform**.
3. Click **Actions** next to your AI Data Platform and click **Delete**.
4. Select whether to delete:
 - All managed data in the Master Catalog
 - All files and folders in all workspaces
5. Enter **Delete** where prompted. Click **Delete**.

Session Management

Oracle AI Data Platform enforces session management controls to ensure a secure and reliable user experience.

The session lifetime for AI Data Platform is 8 hours. User sessions automatically expire after this session lifetime. This behavior aligns with the Oracle Cloud Infrastructure (OCI) security standards and reduces the risk of unauthorized access by ensuring that session tokens are not valid indefinitely.

After 8 hours, users must re-authenticate to continue working in the instance.

These controls are designed to balance usability with strong security and governance practices across the platform.

Part II

Data Engineering

The section explains the methods for developing your data in Oracle AI Data Platform.

Data engineers focus on building and maintaining the systems that data analysts use to access and manipulate data. They use big data technologies like Apache Spark and programming languages including Python and SQL to process and manage data located in object storage, databases, and data warehouses. They are responsible for the initial stages of the data analytics and data science workflow, such as collecting, storing, and transforming data. Their work ensures that the data is accessible and is of high quality so that other data scientists and analysts can use it for their work. Data Engineers also use CI/CD principles for data pipelines and code to manage version control and promote collaboration with data scientists, analysts, and other stakeholders.

Topics:

- [Notebooks](#)
- [Workflows](#)
- [Compute](#)
- [Ingest Data](#)
- [Integration](#)
- [Streaming](#)

4

About Medallion Architecture

Medallion architecture is a layered design pattern for organizing and processing data in a data lakehouse environment.



In a medallion architecture, data is divided into pipelines for three progressive stages: bronze, silver, and gold. The main benefits of a medallion architecture:

- Trusted, high-quality data
- Traceability and auditability
- Modular, efficient pipelines

Oracle AI Data Platform provides you with the tools to establish your bronze, silver, and gold stages and automate cleaning, validating, and enriching your data through each medallion layer.

The bronze layer stores raw data ingested from your source systems, preserving the data in its original form. It acts as your single source of truth and supports traceability and auditability. Data is pulled directly into AI Data Platform, preserving the original data structure and supporting fast, scalable retrieval for downstream processing. Data is then ingested into the bronze layer and is ready for processing.

Data from the bronze layer is processed and sent to the silver layer. Workflows, agents, and other automation help you clean, remove duplicates, and transform your data so it is analytics-ready. Data in the silver layer is often joined across sources to provide standardized, consistent information.

In the silver layer, data analysts manage and develop the cleaned data to prepare it for consumption in data dashboards, advanced analytics, and machine learning models. Data prepared this way is stored in the gold layer, and data in the gold layer is treated as the most trustworthy, and best suited for training models and developing data visualizations, dashboards, and more.

5

Ingest Data

This chapter explains how Oracle AI Data Platform can ingest data from different internal and external sources.

Topics:

- [Internal Sources](#)
- [External Sources](#)

AI Data Platform enables seamless ingestion of data from both external and internal sources using Spark-based notebooks. Whether you're pulling data from cloud services, on-premise databases, or Oracle-native platforms, AI Data Platform provides flexible, code-driven ingestion methods that support your data engineering workflows at scale.

With AI Data Platform ingestion connectors, you can:

- Use drag and drop features to generate notebook code for fast connection setup.
- Ingest batch or near real-time data from a wide variety of systems.
- Leverage Spark and JDBC-based patterns to read, write, and process data efficiently.
- Register external sources as catalogs for direct querying without duplication.

Explore the following sections to learn more:

- **External Sources** - Ingest data from MySQL, PostgreSQL, Kafka, and more.
- **Internal Sources** - Connect with Oracle-native systems like Oracle Autonomous AI Lakehouse, Fusion via BICC, and Oracle Database/Oracle AI Database.

Internal Sources

Oracle AI Data Platform supports ingestion from internal Oracle sources using built-in ingestion connectors. These connectors enable users to seamlessly extract data using Spark-based notebooks and integrate it into their workflows and data pipelines.

Ingestion connectors abstract the complexities of connection setup, providing optimized access patterns for both batch and near real-time ingestion from Oracle-native services.

AI Data Platform provides sample code templates in the [Oracle AI Data Platform Samples](#) Git repository to support ingesting data from several internal sources using Spark in notebooks.

Table 5-1 Internal Sources

Source	Access Type	Integration Method	Description	External Catalog Support	Sample Code Available
Fusion	Extract Only	Preconfigured Spark Templates	Extracts data from Fusion SaaS applications via BICC into AI Data Platform tables or volumes.	No	Yes
Oracle Siebel	Read only	Preconfigured Spark Templates	Reads from Oracle Siebel database into AI Data Platform tables or volumes	No	Yes
Oracle PeopleSoft	Read only	Preconfigured Spark Templates	Reads from Oracle PeopleSoft database into AI Data Platform tables or volumes	No	Yes
REST Endpoints	Read Only	JDBC via Spark Notebook	Reads from APIs for ingesting semi-structured data like JSON.	No	Yes
MySQL HeatWave	Read Only	JDBC via Spark Notebook	Move data between AI Data Platform and MySQL HeatWave using JDBC.	No	Yes
Oracle Autonomous AI Lakehouse	Read/Write + Zero-Copy	JDBC or External Catalog	Ingest from or register Oracle Autonomous AI Lakehouse as an external catalog for querying data directly without duplication.	Yes	Yes
Oracle Autonomous AI Transaction Processing	Read/Write + Zero-Copy	JDBC or External Catalog	Ingest from or register as an external catalog for querying data directly without duplication.	Yes	Yes
Oracle Database	Read/Write	JDBC or External Catalog	Supports data ingestion from on-prem or OCI databases.	Yes	Yes

Table 5-1 (Cont.) Internal Sources

Source	Access Type	Integration Method	Description	External Catalog Support	Sample Code Available
Exadata	Read/Write	JDBC or External Catalog	Access Exadata systems for high-performance reads and writes using JDBC.	Yes	Yes
NetSuite	Read Only	JDBC via Spark Notebook	Reads from NetSuite into AI Data Platform tables or volumes	NO	Yes

Table 5-2 Spark SQL to , Oracle Autonomous AI Database, Exadata Data Type mapping

Spark SQL Type	Oracle AI Database, Oracle Autonomous AI Database, Exadata Data Type
ByteType	NUMBER(38,10)
ShortType	NUMBER(38,10)
IntegerType (INT)	NUMBER(38,10)
LongType	NUMBER(38,10)
FloatType	FLOAT(126)
DoubleType	NUMBER(38,10)
DecimalType(p,s)	NUMBER(p,s)
StringType	VARCHAR2(4000 CHAR)
BinaryType	BLOB
BooleanType	VARCHAR2(4000 CHAR)
DateType	DATE
TimestampType	TIMESTAMP(6)
ArrayType	VARCHAR2(4000 CHAR)
MapType	Not supported
StructType	VARCHAR2(4000 CHAR)
CalendarIntervalType	Supported if converted to String/VARCHAR2

External Sources

Oracle AI Data Platform supports ingestion of data from a wide range of sources using Spark-based notebook connectors. These connectors enable users to ingest and process data directly from external sources in a flexible, code-driven manner.

AI Data Platform provides sample code templates in the [Oracle AI Data Platform Samples](#) Git repository to support ingesting data from several external systems using Spark in notebooks. These templates are pre-built and customizable, allowing users to quickly connect, read, and write data from various commonly used systems.

Table 5-3 External Ingestion Sources

Source	Access Type	Integration Method	Description	External Catalog Support	Sample Code Available
MySQL	Read/Write	JDBC via Spark Notebook	Ingest and export data between AI Data Platform and MySQL databases using JDBC connectors.	Yes	Yes
Azure SQL	Read/Write	JDBC via Spark Notebook	Ingest and export data between AI Data Platform and MySQL databases using JDBC connectors.	Yes	Yes
Snowflake	Read/Write	JDBC via Spark Notebook	Ingest and export data between AI Data Platform and MySQL databases using JDBC connectors.	Yes	Yes
PostgreSQL	Read/Write	JDBC via Spark Notebook	Supports bidirectional data movement with PostgreSQL via JDBC.	No	Yes
MS SQL Server	Read/Write	JDBC via Spark Notebook	Connect and transfer data from Microsoft SQL Server using Spark and JDBC.	No	Yes
Kafka	Read	Kafka Consumer in Spark Notebook	Stream ingestion from Kafka topics	No	Yes
Hive	Read/Write	JDBC via Spark Notebook	Ingest and export data between AI Data Platform and Hive databases using JDBC connectors.	No	Yes

Fusion Data in Oracle AI Data Platform with BICC

You can connect your Oracle AI Data Platform directly to raw Fusion data using the Business Intelligence Cloud Connector (BICC).

You can use Oracle Business Intelligence Cloud Connector (BICC) to extract business intelligence and other data in bulk and load it into designated external storage areas. To learn more about creating an extract using BICC, see [Creating a Business Intelligence Cloud Extract](#).

To create a connection with BICC to your AI Data Platform, you need the following prerequisites:

Fusion Applications Prerequisites

You need to ensure you have the required permissions for your Fusion Applications instance. The user logging into the Fusion Applications instance must have:

- Administrator permissions for the instance
- `ORA_ASM_APPLICATION_IMPLEMENTATION_ADMIN_ABSTRACT` role or a role that includes it

Oracle AI Data Platform Prerequisites

In the OCI compartment where your AI Data Platform instance resides, you need to have:

- An Object Storage Bucket
- Object Storage bucket name
- Object Storage namespace
- Object Storage host name
- OCID value of your AI Data Platform instance tenancy
- OCID value of the user with the API key created to access the Object Storage Bucket

Create an Oracle Business Intelligence Cloud Connector Connection to Oracle AI Data Platform

To use the Oracle Business Intelligence Cloud Connector (BICC) to access your Fusion data directly from your Oracle AI Data Platform.

1. Open any browser and enter `https://<saas cloud host name>:<saas cloud port number>/biacm` to sign into your Fusion Applications instance.
2. Click **Configure External Storage**.
3. Click the OCI Object Storage Connection tab, then click **Add**.
4. Enter the required OCI parameters. See the required OCI parameters in [Oracle AI Data Platform Prerequisites](#).
5. Click **Generate API Signing Key**, then click **Export Public Key**.
6. Open [Oracle Cloud](#) in another window and log in as the user with access to the Object Storage Bucket.
7. In the top-right, click **Profile**, then click **User Details**.

8. Click **API Key** and add the exported public API key from BICC.
9. Return to the BICC window. Click the Console tab, then click **Test the Connection**.
10. Click **Save**.

Add Fusion Data Sources with a BICC Connection

You can use your BICC connection to your Fusion Applications instance to connect data sources to your Oracle AI Data Platform.

1. Open any browser and enter `https://<saas cloud host name>:<saas cloud port number>/biacm` to sign into your Fusion Applications instance.
2. Click **Manage Jobs**, then click **Add** to create a new job.
3. Select the offerings and required public virtual objects you want to add to your AI Data Platform. Click **Save**.
4. Click **Manage Job Schedules**, then click **Add** to create a new schedule.
5. Set the schedule to run immediately or on a ongoing basis. Click **Save**.
6. Once the job has run, check the Object Storage Bucket in OCI to confirm the data is available. Data is exported as zipped CSV files.

Extract Fusion Data from a BICC Connection to a Notebook

Once you have connected Fusion Data to your Oracle AI Data Platform from a BICC connection, you can extract that data from a notebook.

1. From your home page, navigate to a notebook.
2. In your notebook, enter Spark code to extract the data. For example:

```
spark.read.format("aidataplatfrom") \
  .option("type", "FUSION_BICC") \
  .option("fusion.service.url", "https://<saas cloud host name>:<saas cloud port number>") \
  .option("user.name", "john.smith") \
  .option("password", "***password**") \
  .option("schema", "Financial") \
  .option("fusion.external.storage", "FA4IDL") \
  .option("datastore",
    "CrmAnalyticsAM.GeographiesAnalyticsAM.Geography") \
  .load().show()
```

3. Read the data to a data frame and save the data frame as a delta table in a catalog.

6

Workspaces

Workspaces are containers for organizing your notebooks and workflows.

Topics:

- [About Workspaces](#)
- [Create a Workspace](#)
- [Create a Workspace with Private Network Access Enabled](#)
- [Workspaces, External Catalogs, and Data Assets in a Private Network](#)
- [Edit a Workspace](#)
- [Delete a Workspace](#)
- [Create a Folder](#)
- [Delete a Folder](#)
- [Move a Folder](#)
- [Copy a Folder](#)

About Workspaces

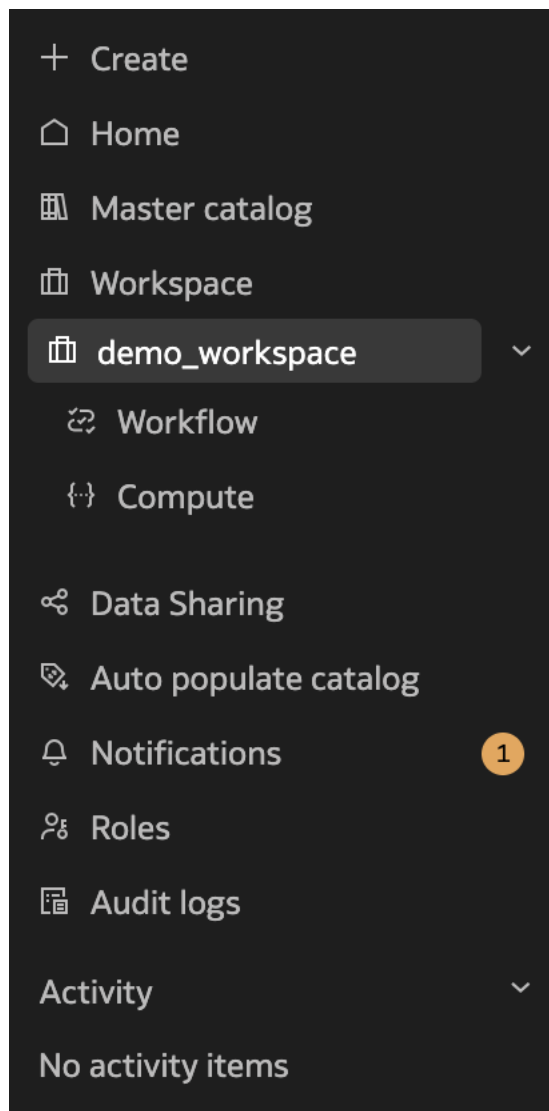
Workspaces are isolated logical containers that organizes your Oracle AI Data Platform notebooks, files, folders, workflows and compute clusters.



A workspace establishes a mapping between the file-like resources you need (e.g. notebooks, files, folders) with corresponding workflow resources (e.g. jobs, job runs) and corresponding compute clusters.

One AI Data Platform instance can have multiple workspaces. A default workspace is created while creating your AI Data Platform instance. You can create additional workspaces within the same AI Data Platform instance depending on your business needs. For example, you can create different workspaces for different teams or business units (e.g. Sales, Marketing, Engineering).

From OCI console, clicking on your AI Data Platform instance takes you to the AI Data Platform home page. At the left-most navigation pane of the home page, you need to select a specific workspace before accessing the workspace files, folder, notebooks, workflows and compute clusters.



Adding Resources to Workspaces

Within a workspace, new files, folders, notebooks and jobs can be created from scratch using the  button.

New files, folders or notebooks can also be uploaded from the local machine. Multi-file upload is coming soon.

Folders

You can use folders in your workspaces to help organize files and other folders containing code, libraries, configurations, metadata or data.

When you create a new workspace, it automatically begins with a Shared folder. You can create additional folders to help organize your files.

You can manage permissions to your folders and assign who can view and modify the content of each folder. For more information, see [Workspace Folder Permissions](#).

When you delete a folder, you will delete all data and metadata contained in that folder.

Editing and renaming workspace resources

You can edit the contents of and rename the files, folders, notebooks or jobs within a workspace.

When you click on a workspace resource, depending on the resource type, the corresponding editor will open where they can be edited and renamed inline. For example, notebooks will open in a notebook editor.

You can rename folder names and descriptions inline at the top of the workspace details page.

File metadata behavior for edit and overwrite operations

When you edit or overwrite an existing file, AI Data Platform preserves the file's creation metadata and refreshes its update metadata. The following behavior applies to both Edit File and Overwrite File operations:

Metadata field	Behavior
createdBy	Remains unchanged and continues to identify the user who originally created the file.
createdTime	Remains unchanged and continues to show when the file was originally created.
updatedBy	Changes to the user who performed the latest edit or overwrite operation.
updatedAt	Changes to the timestamp of the latest edit or overwrite operation.

Create a Workspace

You create a workspace as a container for your files, folders, notebooks, workflows and compute clusters.

To create a workspace, you need the `CREATE_WORKSPACE` permission or the `AI_DATA_PLATFORM_ADMIN` role. The `CREATE_WORKSPACE` permission can be granted directly from the Workspace Listing page.

1. On the Home page, click **Workspace**.
2. Click  **Create Workspace**.
3. Fill in the required details and select the default catalog you want your workspace to pull data from.
4. Click **Create**.

Create a Workspace with Private Network Access Enabled

You can create a workspace that can access data assets in a private network. By default, a workspace is not enabled for accessing data assets in a private network and can only access data assets that are not in a private network. A workspace can only access data assets residing in one private network at a time.

In order to create a workspace that can access existing data assets in an existing private network, you need to have the IAM policies to inspect the VCN of the data asset, the subnet of the data asset, and, if network-security-group is used, to inspect the network-security-group of

the data asset within their respective compartments. For more information, see [IAM Security Policies](#).

1. On the Home page, click **Workspace**.
2. Click **Create Workspace**.
3. Select **Enable private network**.

Network Configuration
If you want to access data sources in that private network, enable private network.

☒ **Enable private network**

In compartment
Select a compartment **Browse**

Choose a VCN
Choose

In compartment
Select a compartment **Browse**

Choose a subnet
Choose

Advanced options
Pick the network security group that has appropriate inbound and outbound traffic rules.

In compartment
Select a compartment **Browse**

Network security group
Add

Add

Are any of your data sources that behind a SCAN proxy?
☒ **Yes**

SCAN details (optional) (Provide the SCAN details if clusters in this workspace will access data sources that are behind a SCAN proxy)

DNS name
Add dns

Port
Add port

Add

4. Select the VCN from the corresponding OCI compartment.
5. Select the subnet from the corresponding OCI compartment.
6. Provide the details of the network security group. Click **Add** if there are multiple.
7. Optional: Select **My data sources require access through a SCAN proxy** if any of your data sources are behind a SCAN proxy. Enter the SCAN proxy settings in the fields provided.

This is a prerequisite if clusters in this workspace need to access Oracle databases that are behind SCAN proxies.
8. Optional: Select **My data sources require advanced DNS configuration** if you have enabled network source and need inbound traffic to come from in-region Service Gateways in your VCN, if any data source is in another region, if your data source is in AWS S3, or if you need to access any oraclecloud.com service behind a private endpoint. Enter the DNS settings in the fields provided.
9. Click **Create**.

Workspaces, External Catalogs, and Data Assets in a Private Network

In order to create an external catalog for a data asset that is in a private network, you will first need to create a workspace enabled for the same private network.

During the creation of an external catalog, you need to pick the workspace that is enabled for the private network of the data asset. For more information, see [Create an External Catalog for Private Networks](#).

Clusters in a workspace can read-write to this external catalog only if the workspace is already enabled for the private network of the data asset.

In order to access data assets in a private network, the VCN of that private network needs to have security rules allowing ingress traffic from all source IP inside that VCN to all destination ports. If that security rule is not already configured, you need to create a new Ingress security rule. For more information, see [Security Rules](#).

If your data asset is in another VCN within the same region, the two VCN needs to be configured with [Local VCN peering \(within region\)](#) using Local Peering Groups or using [Local VCN Peering through an upgraded Dynamic Routing Gateway \(DRG\)](#).

If your data asset is in another VCN in a different region, the two VCN needs to be configured with [Remote VCN peering \(across regions\)](#) using Remote Peering Connections (RPCs) or using [Remote VCN Peering through an Upgraded Dynamic Routing Gateways \(DRG\)](#).

Work with Files

You can store files in volumes in Oracle AI Data Platform and users can organize these files in folders within a volume. AI Data Platform provides you multiple ways to access data stored in volumes and workspaces.

AI Data Platform supports multiple methods for accessing data stored in volumes:

- **POSIX-style paths:** Allow users to provide access to data relative to the driver root (/). Users can read/write data to volumes or workspace folders.
- **URI-style paths:** Allow users to provide access to data using a URI scheme. For example, if you want to read data in OCI Object Storage, you should provide a valid URI scheme to read/write that data.

Here are some examples:

Source	Access Pattern	Example
Volume	POSIX	Example 1 <pre>df_csv = spark.read.csv("/ Volumes/ <<catalog_name>>/ <<schema_name>>/ <<volume_name>>/ <<file_name>>.csv", header=True, inferSchema=True, sep=",")</pre> Example 2 <pre>import pandas as pd df_panda_csv=pd.read_csv v("/Volumes/ <<catalog_name>>/ <<schema_name>>/ <<volume_name>>/ <<file_name>>.csv", header=0, sep=",")</pre>

Example 3

Source	Access Pattern	Example
Workspace	POSIX	Example 1
		<pre>df_csv = spark.read.csv("/ Workspace/ <<folder_path>>/ <<file_name>>.csv", header=True, inferSchema=True, sep=",") df_csv.show()</pre>
		Example 2 <pre>import pandas as pd df_panda_csv=pd.read_csv v("/Workspace/ <<folder_path>>/ <<file_name>>.csv", header=0, sep=",") df_panda_csv.head()</pre>
		Example 3 <pre>import osos.listdir("/ Workspace/ <<folder_path>>/")</pre>
OCI Object Storage	URI	
		<pre>spark.read.format("json").load("file:/// Workspace/ <<folder_path>>/ <<file_name>>.json").show()</pre>
OCI Object Storage	URI	
		<pre>df_csv = spark.read.csv("oci:// <<bucket_name>>@<<namesp ace>>/<<folder/file>>", header=True, inferSchema=True, sep=",")</pre>

Edit a Workspace

Workspace attributes can be changed after the workspace is created.

Most changes to the network configuration of a workspace first require all existing compute clusters to be manually stopped.

1. On the Home page, click **Workspaces**.
2. Next to the workspace you want to edit, click ... **Actions** and click **Edit**.
3. Make your changes to the workspace and click **Edit**.

Delete a Workspace

Deleting your workspace removes it and all your contained jobs, tasks, and compute clusters.

1. On the Home page, click **Workspaces**.
2. Next to the workspace you want to delete, click ... **Actions** and click **Delete**.
3. On the confirmation page, click **Delete**.

Create a Folder

You can create folders in workspaces where you have appropriate permissions.

1. In the workspace you want to create a folder, click ⊕ **Create** then **Folder**.
2. Provide the folder name and a folder description. Click **Create**.

Delete a Folder

You can delete folders and their contained data and metadata.

1. Navigate to the folder you want to delete in your workspace.
2. Click ... **Actions** then click **Delete**.
3. Click **Delete**.

Move a Folder

You can move folders and their contained data and metadata to another location in your workspace.

1. Navigate to the folder you want to move to a new location in your workspace.
2. Click ... **Actions** then click **Move**.
3. Select a new location for your folder and its contents.
4. Click **Move**.

Copy a Folder

You can create a duplicate of a folder and its contents in your workspace.

1. Navigate to the folder you want to copy to another location in your workspace.
2. Click ... **Actions** then click **Copy**.
3. Select a location to copy the folder and its contents to.
4. Click **Copy**.

7

Notebooks

This chapter provides information on using and managing notebooks in your workspace.

Topics:

- [Develop Code in Notebooks](#)
- [Manage Notebooks](#)
- [Run Notebooks](#)
- [Python Wheel Files](#)
- [Notebook Navigation](#)
- [Notebook Output and Results](#)
- [Notebook Appearance](#)
- [Manage Job Runs](#)
- [Migrate Existing Apache Spark Code to Oracle AI Data Platform](#)

Develop Code in Notebooks

Data engineers and data scientists can use notebooks in their Oracle AI Data Platform as a common tool for interactively developing code and exploring data.

Oracle AI Data Platform currently supports Python, SQL, and Scala languages in notebooks. Notebooks can be scheduled or configured to run as part of a workflow. To run notebooks, you need to attach a compute cluster.

Your AI Data Platform comes with integrated managed notebooks for an intuitive developer experience.

You can use the sample code in the [Oracle AI Data Platform Samples](#) Git repository for examples of code you can use with your notebook.

Auto-save

Notebooks are automatically saved every two minutes.

Importing and Exporting Notebooks

You can currently import a notebook file (*.ipynb) from your local machine to your workspace.

Exporting notebooks is not currently supported.

Create a Notebook

You can create a notebook in any workspace you have administrator permissions.

1. On the Home page, navigate to your workspace.
2. Click **Create** and click **Notebook**.

3. Fill in the name and description, then click **Create**.

Attach an Existing Cluster to a Notebook

Notebooks require an attached cluster to provide compute power for developed code.

1. On the Home page, navigate to your workspace and open your notebook.
2. Click **Actions** then click **Attach an existing cluster**.
3. Click on the cluster you want to use from the list.

Your notebook will show **Cluster: (ClusterName) running** when it has been successfully attached. This can take up to several minutes.

Create a Cluster for a Notebook

You can create a new cluster directly from the notebook interface and attach it immediately.

For more information, see [About Compute Clusters](#).

1. On the Home page, navigate to your workspace and open your notebook.
2. Click **Actions** then click **Create cluster**.
3. Select **Runtime version**.
4. Select the driver options for your cluster.
5. Select the worker options for your cluster. These options apply to all cluster workers.
6. Select whether the number of workers is static or scales automatically.
 - If **Static amount**, specify the number of workers.
 - If **Autoscale**, specify the minimum and maximum number of workers the cluster can scale to.
7. For **Run duration**, select whether the cluster will stop running after a set duration of inactivity. If **Idle timeout** is selected, specify the idle time, in minutes, before the cluster will time out.
8. Click **Create**.

Default Language

You can use notebooks to develop and run Apache Spark code in Python, SQL, or Scala.

The default language for notebooks is Python. You can change the default language for the whole notebook or for individual cell(s) to SQL, Scala, or Markdown or raw text. You can combine Python, SQL, and Scala code in different cells within the same notebook.

Notebooks have syntax highlighting for Python, SQL, and Scala. New notebook cells will be created based on the default language of the notebook.

Widgets in Notebooks

Widgets let you define input values in a notebook.

You can provide values when running the notebook interactively. Jobs, workflow task, or parent notebooks can also provide values at run time.

You use widgets when the same notebook must run with different values, such as a processing date, environment, region, or dataset name. Defining inputs with widgets lets you reuse the notebook without editing its code for each run.

Notebooks can include the following widget types

Widget Type	Description
text	Accepts any string value.
dropdown	Users select one value from a defined list.
combobox	Users select a defined value or enter a custom value.
multiselect	Users select one or more values from a defined list.

Widget names must be unique within a notebook. All widget values are returned as strings. A multiselect widget returns selected values as a comma-separated string.

Note

You need at least USE permissions for the notebook to change widget values for interactive notebook runs. You need MANAGE permissions to create or remove widgets.

Notebook Widget Examples

The following examples demonstrate different use cases for widgets in your notebooks.

When you run the notebook interactively, the widgets display in the widget panel. You can change their values before running the notebook.

Text Widget

This example creates a single text widget takes the provided string value, in this case the processing date:

```
aidutils.widgets.text(  
    name="date",  
    defaultValue="2026-01-01",  
    label="Processing date"  
)
```

Dropdown Widget

This example creates a single dropdown widget that allows users to select one value from a defined list:

```
aidutils.widgets.dropdown(  
    name="environment",  
    defaultValue="dev",  
    choices=["dev", "test", "prod"],  
    label="Environment"  
)
```

Combobox Widget

This example creates a single combobox widget, which provides defined values for a user to select from a list, and a field to enter a custom value:

```
aidutils.widgets.combobox(  
    name="campaign",  
    defaultValue="renewal_outreach",  
    choices=["renewal_outreach", "win_back", "loyalty_offer"],  
    label="Campaign"  
)
```

Multiselect Widget

This example creates a single multiselect widget, which provides defined values for a user to select multiple options from:

```
aidutils.widgets.multiselect(  
    name="regions",  
    defaultValue="North America",  
    choices=["North America", "Europe", "Asia Pacific", "Latin America"],  
    label="Regions"  
)  
  
regions = aidutils.widgets.get("regions")
```

Multiple Widgets

You can use `aidutils.widgets` to create multiple widgets in your notebook. This example creates a text widget that has the processing date as a string and a dropdown widget for selecting an environment from a defined list, then outputs the results.

```
aidutils.widgets.text(  
    name="date",  
    defaultValue="2026-01-01",  
    label="Processing date"  
)  
  
aidutils.widgets.dropdown(  
    name="environment",  
    defaultValue="dev",  
    choices=["dev", "test", "prod"],  
    label="Environment"  
)  
  
date = aidutils.widgets.get("date")  
environment = aidutils.widgets.get("environment")  
  
print(date)  
print(environment)
```

Widget Value Overrides

A notebook uses the default value when no other value is provided. To use a different value for a workflow job run, you add a parameter to the notebook task that overrides the original value.

Note

The parameter key must match the widget name.

For example, a notebook that defines an environment widget with a default value of dev can receive the following task parameter:

```
environment = prod
```

When the task runs, the notebook uses prod instead of dev.

You can also pass values to a notebook called from another notebook. In this example, you use the parameters argument in `aidutils.notebook.run()`:

```
aidutils.notebook.run(  
    "process_data",  
    parameters={  
        "date": "2026-03-01",  
        "environment": "test"  
    }  
)
```

When the called notebook contains widgets named date and environment, the values passed by the parent notebook override the widget defaults for that execution. For more information about configuring workflow parameters, see [Parameterization](#).

Widget Management

You can remove one or more widgets from a notebook with `remove()` and `removeall()` commands from the `aidutils` library.

Note

You need MANAGE permissions to create or remove widgets.

Use `remove()` to remove one widget from your notebook:

```
aidutils.widgets.remove("environment")
```

Use `removeAll()` to remove every widget from a notebook:

```
aidutils.widgets.removeAll()
```

Manage Notebooks

You can rename and delete notebooks your own. You can also clone notebooks to copy their contents and work with its code in a new notebook.

You rename and delete your notebooks from their action menu in your workspace. You clone your notebook by opening it and selecting the **Clone** option from the File menu.

Rename a Notebook

If the name of your notebook is no longer helpful or relevant, you can change it at any time.

1. On the Home page, navigate to your workspace.
2. Next to the notebook you want to rename, click **Actions** then **Rename**.
3. Enter a new name and click **Save**.
4. Optional: You can also change the name of an open notebook by clicking the name and entering a new one.

Delete a Notebook

You can delete notebooks that you have administrator permissions for.

1. Navigate to your workspace.
2. Next to the notebook you want to delete, click **Actions** then **Delete**.
3. Click **Delete**.

Clone a Notebook

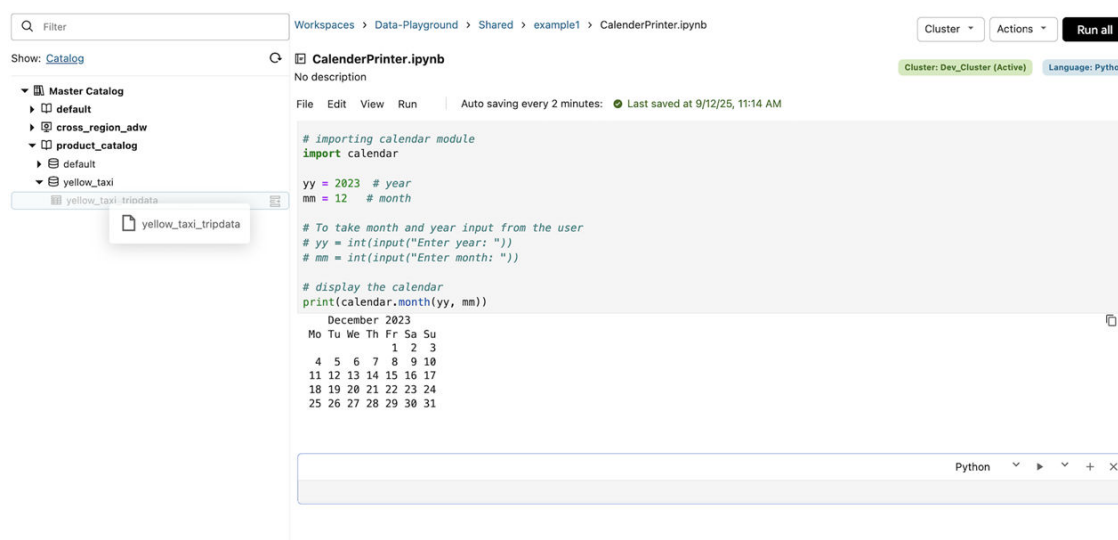
You can clone an existing notebook to create a copy of the content of that notebook you can modify while retaining the original.

1. Open the notebook you want to clone.
2. In the notebook toolbar, click **File** then click **Clone**.
3. Enter a new name for the cloned notebook.
4. Click **Browse** to select the workspace folder to save the cloned notebook into. If no folder is selected, the cloned notebook is created in the same folder as the notebook you are cloning.
5. Select whether to include or exclude outputs. Outputs are included by default. Clear the selection to exclude outputs.
6. Click **Clone**. The cloned notebook is created in the workspace folder you specified.

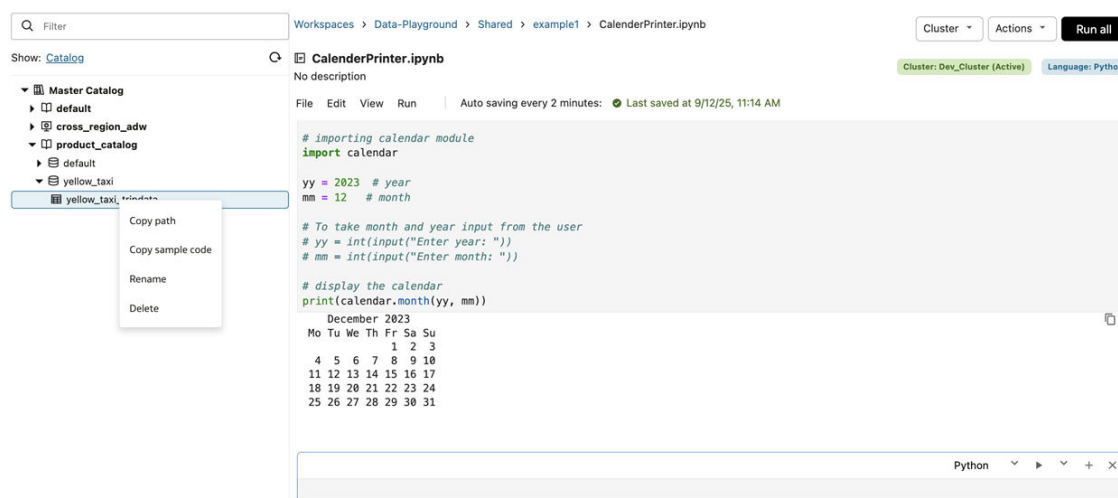
Browse Resources While Editing Notebook

When you are in a notebook, you can browse the Catalog or workspace objects on the left side without leaving your notebook.

If you drag and drop any object from the left hand pane to the notebook, the object name or the full path is copied and pasted to the notebook cell (depending on the context).



You also have a button and context menu options available for each catalog or workspace object in the left hand pane. The context menu at the left navigation has options to copy sample code, copy name, or copy path and so that you can paste to your notebook cell.



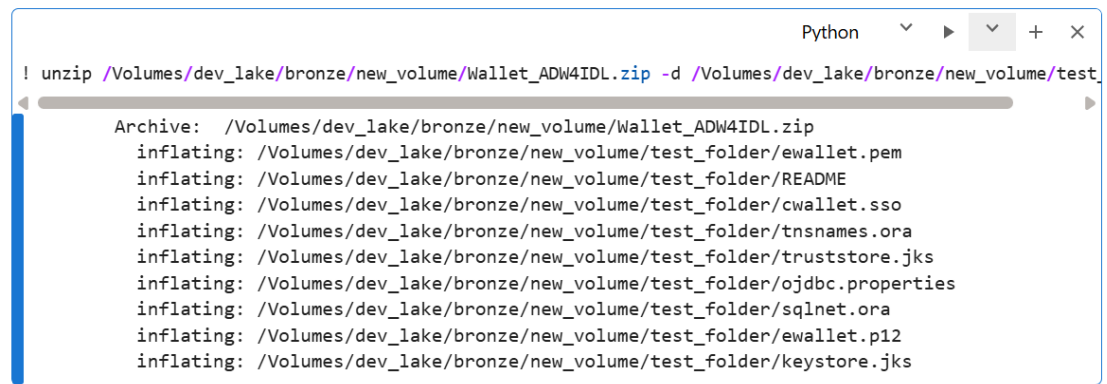
Run Notebooks

You can run code in notebooks you own or from notebooks that are shared with you.

Code can be run from a notebook using three methods: running on demand, running as a one-off manual run, or creating a scheduled notebook job. Jobs run on demand are run only once.

Running Terminal Commands Within a Notebook

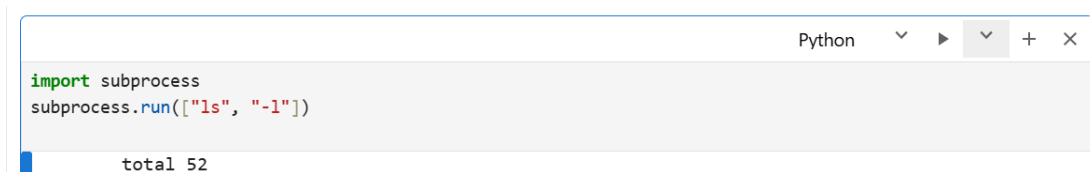
You can run basic terminal commands or shell commands within a notebook by prefixing with an '!'. For example, you can use the **unzip** command to extract from ZIP files in the workspace.



```
Python  ▾ ▶ ▾ + ×  
! unzip /Volumes/dev_lake/bronze/new_volume/Wallet_ADW4IDL.zip -d /Volumes/dev_lake/bronze/new_volume/test_

Archive:  /Volumes/dev_lake/bronze/new_volume/Wallet_ADW4IDL.zip
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/ewallet.pem
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/README
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/cwallet.sso
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/tnsnames.ora
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/truststore.jks
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/ojdbc.properties
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/sqlnet.ora
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/ewallet.p12
  inflating: /Volumes/dev_lake/bronze/new_volume/test_folder/keystore.jks
```

You can also use the **subprocess** module in Python for shell script execution.



```
Python  ▾ ▶ ▾ + ×  
import subprocess  
subprocess.run(["ls", "-l"])
```

total 52

You can also use native Python modules like **zipfile** for tasks like unzipping files as an alternative to shell commands.

Limitations

Currently, Oracle AI Data Platform does not have native support for pip install, CI/CD, Git, or version control systems.

Run Options for Notebook Cells

The Run menu in notebooks provides you with options for running cells in a notebook.

You can find all options for running cells in your notebook from the **Run** menu at the top of your notebook.

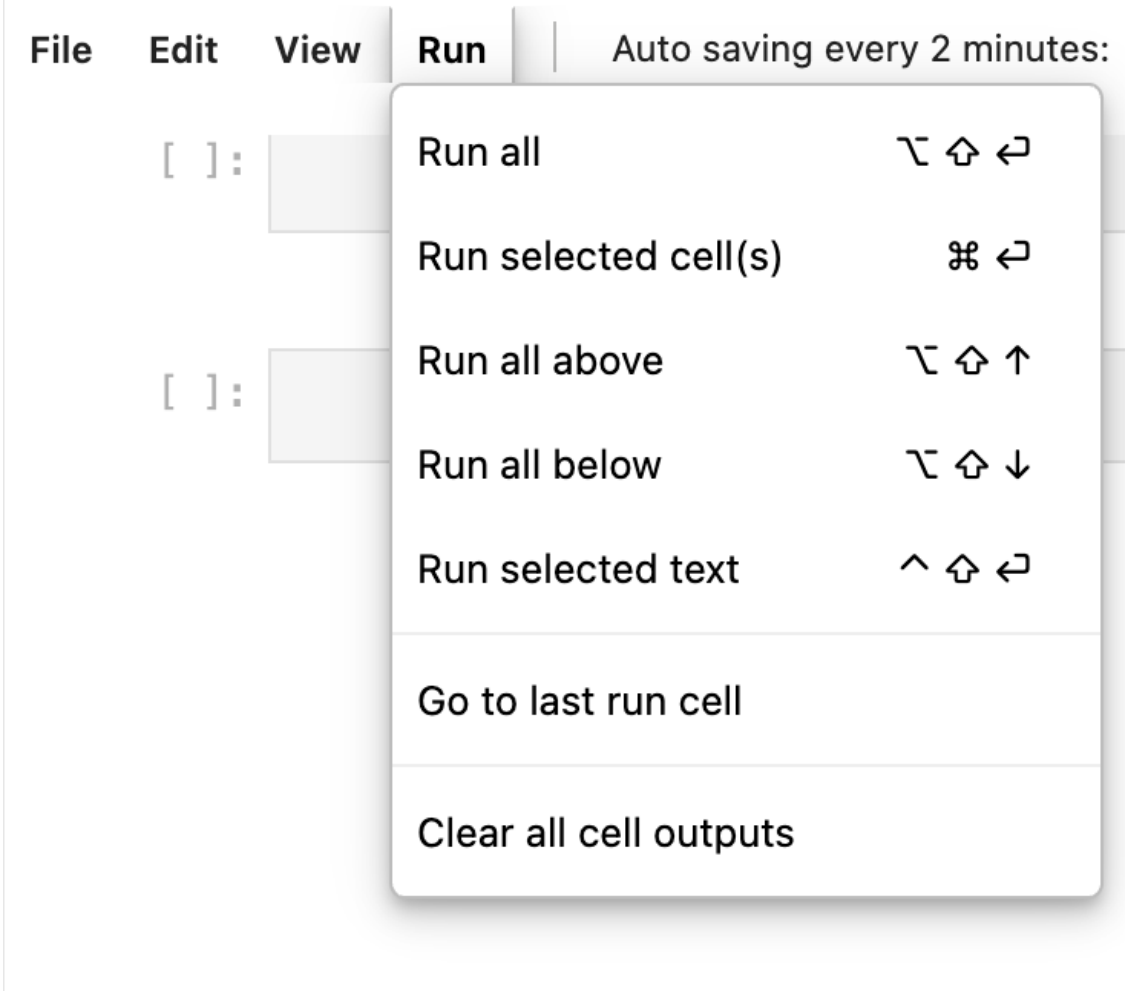


Table 7-1 Run Options for Notebook Cells

Option	Description
Run all	Runs all cells in the notebook sequentially.
Run selected cell(s)	Runs the currently selected cell or cells.
Run all above	Runs the currently selected cell and any cells that appear above the currently selected cell in the notebook.
Run all below	Runs the currently selected cell and any cells that appear below the currently selected cell in the notebook.
Run selected text	Runs the selected segment of code in a cell.
Go to last run cell	Navigates you to the most recently run cell in the notebook.
Clear all cell outputs	Removes outputs from all cells in the notebook.

Run Code from a Notebook

You can choose to run all code developed in a notebook at once, or one cell at a time.

Keyboard shortcuts for running code in a notebook are:

- **MacOS:** Cmd + Return
- **Windows:** Ctrl + Enter

You can run code in a single cell by clicking the ▶ **Play** button, or run the whole notebook by clicking **Run all**.

1. On the Home page, click **Workspace**.
2. Navigate to your notebook.
3. Click **Run all**.
4. Check the status of your notebook job run by clicking **Workflow** then **Job Runs**.

Run Code from Another Notebook

You can use `%run` magic command in a notebook to include code from another notebook.

In the following example, you bring in code from a notebook named **called-notebook.ipynb**, to a notebook **caller-notebook.ipynb**.

1. Install the **nbconvert** Python library.
2. Use the `%run` command in a cell, as in the following example:

```
%run /Workspace/folder1/called-notebook.ipynb
```

After following these steps, the notebook named **called-notebook.ipynb** is immediately run using your user principal (i.e. **caller-notebook.ipynb**) and using the attached cluster of **caller-notebook.ipynb**. All the functions and variables defined in **called-notebook.ipynb** immediately become available in the notebook named **caller-notebook.ipynb**.

Share Notebook Output with oidlUtils

You can capture and share content output by your notebook tasks by using utilities available in `oidlUtils`.

`oidlUtils` is a set of utilities available to all users of Oracle AI Data Platform. When sharing content between notebooks, `oidlUtils` can be called to pass arguments to one notebook and pass output back to the caller notebook, and it can be called in a job task in a notebook to pass output back to the parent task, meaning the task calling on the notebook. Used this way, you can capture and use structured output returned by your notebook tasks.

These `oidlUtils` modules are available for use with notebooks:

Module	Description	Example
notebook	Orchestrate notebook task runs and return a single structured result (commonly a JSON string) to the caller.	<pre>oidlUtils.notebook.run(notebook_path: str, timeout_seconds: int = 0, (Optional) parameters: dict (Optional)) -> str</pre>

Module	Description	Example
notebook	Allow a notebook to exit a task run and return a single string result (commonly a JSON payload) to the caller notebook or job/task output API.	<code>oidlUtils.notebook.exit(value: str)</code>

Example 1: Notebook to Notebook Sharing

In this example you use Notebook A to invoke Notebook B. Notebook B returns a result payload that Notebook A is set up to capture and use.

Notebook A

```
result = oidlUtils.notebook.run("NotebookB", 0)
print("Output from Notebook B:", result)
```

Notebook B

```
import json

payload = {
    "status": "SUCCESS",
    "rows_processed": 1234,
    "output_table": "sales_gold",
    "run_id": "run_2026_02_11"
}

json_payload = json.dumps(payload)
oidlUtils.notebook.exit(str(json_payload))
```

Output from Notebook B

```
{"status": "SUCCESS", "rows_processed": 1234, "output_table": "sales_gold",
"run_id": "run_2026_02_11"}
```

Example 2: Passing Output through Job Tasks

In this example, you return a JSON file when your notebook is run as a task: `oidlUtils.notebook.exit(json.dumps(payload))`.

```
import json

payload = {
    "status": "SUCCESS",
    "output_table": "sales_gold",
    "rows_processed": 1234
}

oidlUtils.notebook.exit(json.dumps(payload))
```

Next you run the job with your notebook task and get the task output through the API call:
endpoint = f"https://<workspace-url>/jobs/runs/get-output?run_id={task_run_id}"
and response = requests.get(endpoint, headers=headers).json().

```
import requests
```

```
task_run_id = "<task_run_id>"
```

```
endpoint = f"https://<workspace-url>/jobs/runs/get-output?  
run_id={task_run_id}"  
response = requests.get(endpoint, headers=headers).json()
```

Last, you capture the output returned by the notebook with `job_result = response['notebook_output']['result']`.

```
job_result = response["notebook_output"]["result"]  
payload = json.loads(job_result)  
  
print(payload["output_table"]) # Output : sales_gold  
print(payload["rows_processed"]) # Output : 1234
```

Notebook Navigation

You can create and maintain a table of contents that can be used to organize and navigate your notebook.

You can click the table of contents icon in the top left of your notebook to display a notebook outline. The table of contents is automatically generated based on markdown headings you can create, enabling easy organization and navigation.

You can add formatted text, headings, lists, and documentation as markdown to organize and explain notebook content for yourself and other users.

Create a Markdown Cell

You can create markdown cells to provide headings in your notebook table of contents for ease of organization and navigation.

1. From the Home page, navigate to your notebook.
2. From the cell type drop-down list, select **Markdown**.
3. Add your markdown to the cell.
 - Markdown cells can include formatted text, headings, lists, and other documentation.
 - To create a heading, start a line with # followed by a space. Heading 1 uses one #, heading 2 uses ##. Add additional # for each additional heading level.

Python Wheel Files

You can use Python wheel files to quickly and easily add Python code to notebooks in your workspace.

Python wheel files (.WHL files) allow you to package and distribute your Python code in a single file. This makes installing existing Python code to your notebooks a simpler process.

You need to have Python wheel and setuptools packages installed. You can use pip to install these packages from your notebook. For example:

```
%pip install wheel setuptools
```

Your Python wheel files should follow a format similar to this:

```
/Workspace/<project>/
  setup.py
  customer_churn/
    __init__.py
    analytics.py
    predictor.py
    utils.py
  data/
    samples_customers.csv
```

In your wheel file, you should:

- Keep setup.py at the project root
- Keep importable package code inside the package folder. In the example given, the package folder would be customer_churn/
- Put reusable logic libraries in .PY files, not in .IPYNB files.

A setup.py file looks like this:

```
from setuptools import setup, find_packages

setup(
    name="customer_churn",
    version="1.0.0",
    author="admin",
    author_email="admin@corporation.com",
    description="A package for customer churn analysis and prediction",
    packages=find_packages(),
    classifiers=[
        "Development Status :: 4 - Beta",
        "Programming Language :: Python :: 3",
        "License :: OSI Approved :: MIT License",
    ],
    python_requires=">=3.6",
    install_requires=[
        "pyspark>=3.0.0",
    ],
    package_data={
        'customer_churn': [
            'data/*.csv',
            '*.ipynb', # Include notebook files
        ],
    },
    include_package_data=True,
)
```

A `customer_churn/__init__.py` file looks like this:

```
"""
Customer Churn Analysis Package
Provides tools for calculating and predicting customer churn
"""

__version__ = "1.0.0"
__author__ = "admin"

from .analytics import ChurnAnalyzer
from .utils import load_customer_data, calculate_churn_rate

__all__ = [
    'ChurnAnalyzer',
    'load_customer_data',
    'calculate_churn_rate'
]
```

Every import in `__init__.py` must point to a real `.PY` file in the package. For example, both `utils.py` and `analytics.py` should be inside the `customer_churn` folder.

Build and Install a Wheel

To build a wheel, you run these commands from your notebook:

```
%%bash
cd <project_folder_location> python setup.py bdist_wheel
ls -lh dist      #wheel file is created inside dist/ folder
```

For example, for our customer churn model, our build looks like this:

```
%%bash
cd /Workspace/<projectName>/customer_churn_whl
python setup.py bdist_wheel
ls -lh dist
```

To install our customer churn wheel file, run the following commands:

```
pip install /Workspace/<project>/dist/customer_churn-1.0.0-py3-none-any.whl
```

Once installed, you can verify the installation with the following:

```
pip show customer_churn
```

Finally, you can test the installation with the following:

```
import customer_churn
from customer_churn import ChurnAnalyzer, load_customer_data

print("Package imported successfully!")
print(f"Version: {customer_churn.__version__}")
```

Update a Wheel Package

You can update a Python wheel package by editing the source files directly in your workspace project.

① Note

Do not edit the installed files in site-packages, and do not modify the .WHL file directly.

1. Navigate to the files in your project source. For example, `Workspace/<projectName>/customer_churn`.
2. Update `setup.py` to change package meta data or packaging information, such as:
 - Package name
 - Version
 - Dependencies
 - Packaged data files
3. Update `__init__.py` to change what the package exports at the time of import, such as:
 - Public classes or functions
 - Package version
 - Top-level imports
4. After making your edits, rebuild the wheel:

```
%%bash
cd /Workspace/<projectName>/customer_churn_wheel
python setup.py bdist_wheel
ls -lh dist
```

5. Uninstall the old wheel:

```
pip uninstall /Workspace/<projectName>/dist/<OldWheelName>.whl
```

6. Install the new wheel:

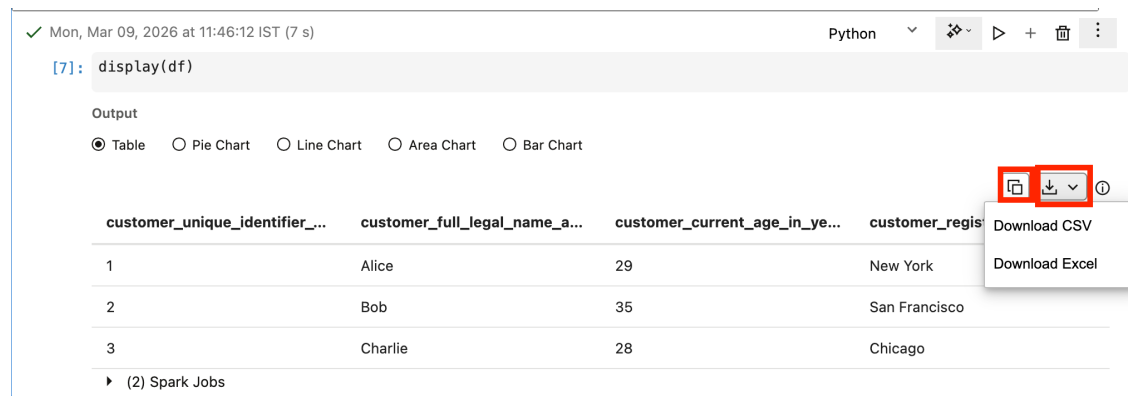
```
pip install /Workspace/<projectName>/dist/<NewWheelName>.whl
```

Notebook Output and Results

You can see notebook outputs and results in a new cell that appears right after the cell with code.

While a cell is in progress, you can cancel the execution of the cell. If a notebook is run as a workflow job, the output is not visible in the same notebook. In that case the output is visible in the output area of the corresponding workflow job run.

You can download the output from the output cell as a CSV or Excel file. You can also copy the contents of the output cell directly to your clipboard.



Download Notebook Output

You can download the resulting output of a notebook cell directly from the results panel.

1. In the top-right of the cell you want to download output from, click **Download**.
2. From the menu, select the format you want to download output as.

Copy Notebook Output

You can copy the resulting output of a notebook cell directly from the results panel.

1. In the top-right of the cell you want to download output from, click **Copy**.
2. The contents of the output cell are copied directly to your clipboard. You can paste it elsewhere in your notebook or to an external location.

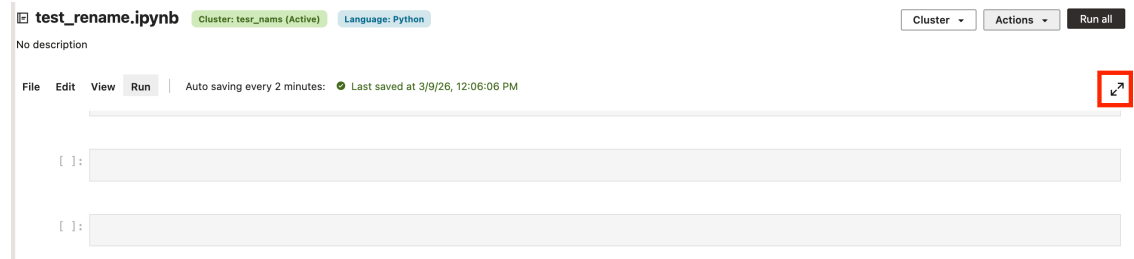
Notebook Appearance

You can alter the screen space available for working in notebooks by minimizing the left navigation panel or expanding the notebook view.

You can minimize the Oracle AI Data Platform left navigation panel while working on your notebook to increase the usable screenspace by clicking **Minimize panel** on the bottom-right of the navigation panel.

 Create Home Master catalog Workspace test_nams  Workflow Compute Agent flows Data sharing Auto-populate catalog Notifications  Roles Audit logs

You can also expand the notebook by clicking **Expand** at the top right of the notebook, expanding the space available and making larger cells and outputs easier to read.



Cell Run Numbers

Each notebook cell displays a run number that indicates the order in which cells were run. This number updates every time the cell is run. You can run cells in any order, so the run numbers may not match the physical order of cells in your notebook.

Manage Job Runs

You can create job runs to manage how and when code is run from your notebooks.

Manual job runs can be run again or later set up to run on a schedule. Scheduled job runs are automatically triggered based on the schedule you set. Unless a schedule is configured, manual jobs are only run once.

Create a Manual Run Job from a Notebook

You can create an unscheduled job that you can run manually from code you've developed in your notebook.

1. On the Home page, click **Workspace**.
2. Navigate to your notebook.
3. Click **Actions**, then click **Schedule**.
4. Provide a name and description for the job.
5. Click **Browse** and select the location to store your job. Click **Select**.
6. Select a compute cluster from the **Cluster** dropdown.
7. For **Schedule**, select **Manual Run**.
8. Click **Create**.

Create a Scheduled Job Run from a Notebook

You can create a scheduled job that runs automatically from code you've developed in your notebook.

1. On the Home page, click **Workspace**.
2. Navigate to your notebook.
3. Click **Actions**, then click **Schedule**.

4. Provide a name and description for the job.
5. Click **Browse** and select the location to store your job. Click **Select**.
6. Select a compute cluster from the **Cluster** dropdown.
7. For **Schedule**, select **Schedule**.
8. Select a **Schedule Status**.
 - Select **Active** if you want the schedule to be enabled immediately.
 - Select **Paused** if you want to manually enable the scheduled run at a later time.
9. Provide a time zone for the schedule to be based on.
10. Select the **Schedule Type**.
 - For **Calendar**, you must specify the frequency and which hours or days the schedule will repeat on.
 - For **Cron Expression**, you must provide the schedule in the form of a cron expression.
11. Check the listed run time at the bottom to confirm your schedule is correct. Click **Create**.

Notebook Keyboard Shortcuts

You can use keyboard shortcuts to simplify using commands in your notebook.

Windows	macOS	Action
Ctrl + Enter	Cmd + Return	Execute cell
Shift + Enter	Shift + Return	Execute cell and advance to next cell
Ctrl + S	Cmd + S	Save notebook
Ctrl + N	Ctrl + N	New notebook
Ctrl + Z	Cmd + Z	Undo
Ctrl + Y	Cmd + Y	Redo
Ctrl + C	Cmd + C	Copy
Ctrl + X	Cmd + X	Cut
Ctrl + V	Cmd + V	Paste
Ctrl + Alt + F	Ctrl + Option + F	Find and Replace
Ctrl + Shift + A	Ctrl + Shift + A	Insert cells above
Ctrl + Shift + B	Ctrl + Shift + B	Insert cells below
Ctrl + Alt + Up	Ctrl + Option + Up	Move cell up
Ctrl + Alt + Down	Ctrl + Option + Down	Move cell down
Ctrl + D	Ctrl + D	Delete cell
Alt + Shift + Enter	Option + Shift + Return	Run All
Alt + Shift + Up	Option + Shift + Up	Run all above cells

Migrate Existing Apache Spark Code to Oracle AI Data Platform

You can adapt your Apache Spark code to migrate it for use in Oracle AI Data Platform notebooks.

If you are migrating existing Spark code from other platforms, you can use the following guidelines to adapt your code for use in notebooks.

Table 7-2 Apache Spark to AI Data Platform Migration Guidelines

Guideline	Details
Remove SparkSession creation commands	AI Data Platform automatically creates a SparkContext for each compute cluster. We recommend removing the session creation commands or replacing them with SparkSession.builder().getOrCreate().
Remove session termination commands, like sys.exit() or spark.stop()	All purpose compute clusters are shared clusters, so if any users stop the SparkSession, by using sys.exit() or spark.stop() for example, the cluster needs to be restarted for everyone. To avoid disruption, we recommend avoiding those commands in the notebooks.

8

Workflows

Workflows in Oracle AI Data Platform provide a powerful and flexible way to automate data processing tasks. With workflows, users can define and orchestrate complex data pipelines that can run on-demand and based on a pre-defined schedule. Workflows can be composed of multiple tasks, each performing a specific action, and can include advanced features such as dependencies, triggers, and error handling.



[Video](#)

Topics:

- [Configure Jobs](#)
- [Configure Tasks](#)
- [Parameterization](#)
- [Monitor Jobs](#)

Key Features of AI Data Platform Workflows

- Automation: Automate complex data tasks and processes.
- Orchestration: Define the sequence and dependencies of tasks in a pipeline.
- Scheduling: Run workflows on a schedule or trigger based on specific events.
- Monitoring: Track workflow status, logs, and execution history.
- Parameterization: Pass parameters to customize the behavior of workflows and tasks.

Core Concepts

- Job: A collection of tasks executed in sequence or parallel to complete a data processing job.
- Task: The individual steps that make up a workflow. Tasks can include actions like running Python code, executing a notebook, if-else task or running another job task.
- Job Run: An instance of a job execution. A job can be triggered multiple times, each time representing a new job run.
- Trigger: Defines the conditions under which a workflow is executed, such as on a schedule, or if it's manually triggered.
- Dependencies: Define the order of task execution or specify conditions under which certain tasks run.
- Parameters: Values passed to workflows or tasks to customize their execution. Parameters can be defined at the job, task, or runtime level.

Benefits/Use Cases of Using Workflows

- Streamlined Automation - Simplify the execution of recurring data tasks by automating them through workflows.
- Parallel Processing - Speed up data processing by running tasks in parallel.

- Customizable Execution - Modify workflows at runtime with parameters to meet specific needs.
- Improved Efficiency - Reduce manual interventions and errors, enabling smoother operations.

Workflows in an AI Data Platform enable a wide range of use cases, including automated ETL pipelines, data integration from multiple sources, and advanced analytics. Users can automate data quality monitoring, machine learning model training, and deployment. These capabilities drive efficiency and scalability for modern data-driven workflows.

Best Practices

- Task Modularization - Break down workflows into reusable tasks to simplify management and improve maintainability.
- Efficient Resource Allocation - Optimize workflows for better performance by running tasks in parallel when appropriate.
- Error Handling - Use retries, error notifications, and fallback mechanisms to ensure workflows run reliably.
- Compute Assignment - Assign specific compute resources to each task based on workload size, optimizing performance and cost.

By following these best practices, you can design workflows that are scalable, reliable, and efficient, ensuring optimal performance and easier management in AI Data Platform.

Configure Jobs

This section covers configuring jobs and job runs in your Oracle AI Data Platform.

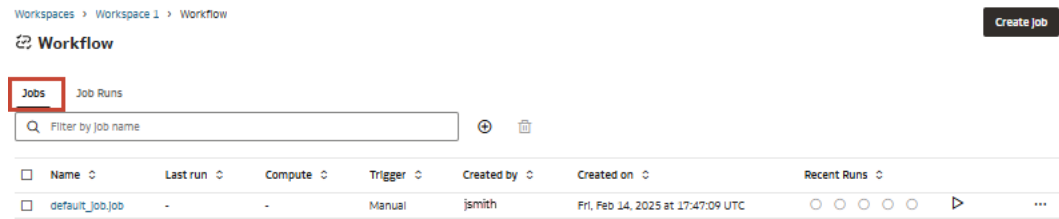
Topics:

- [About Jobs](#)
- [Create a Job](#)
- [Change Job Location](#)
- [Delete a Job](#)
- [Schedule a Job Using a Calendar](#)
- [Schedule a Job Using a Cron Expression](#)
- [Run a Workflow on Demand](#)
- [Run a Workflow Job from the Jobs Page](#)
- [Change a Job Run Schedule](#)
- [Pause or Activate a Job Run Schedule](#)

About Jobs

You create workflows for your data by constructing jobs.

You can track and manage all jobs in your Oracle AI Data Platform from the **Jobs** tab in the Workflow home page. From the **Jobs Run** tab you can see job history and the status of runs currently in progress.



Jobs are way of organizing and orchestrating groups of tasks as part of workflow. You can use workflows for common data processing, such as ETL workflows, Python scripting, running notebooks, and machine learning.

Jobs can vary in complexity. One job may have only a single task that runs a notebook, while another has more than a hundred tasks and nested jobs carrying out complex tasks with multiple conditions and dependencies.

Running a job starts the tasks inside with the sequences and conditions you specified. Jobs can even be nested inside of other jobs, embedding their sequence of tasks as a single node.

Every job run generates a job file that is stored in a user-specified location during job creation. This job file acts as a reference for managing and tracking job executions. You can run, schedule, and view the JSON representation of the job directly from the file in the workspace, ensuring transparency and reproducibility. This approach allows seamless job monitoring, versioning, and integration within automated workflows.

Jobs can configured to run on a calendar schedule, scheduled to run based on a cron expression, or run immediately.

Create a Job

You need to create a job to begin using workflows in Oracle AI Data Platform.

1. Navigate to your workspace and click **Workflow**.
2. On the **Jobs** tab, click **Create Job**. You can also click **Create Job** on the top right.

The screenshot shows the 'Create Job' form. It has the following fields and controls:

- Job Name:** A text input field containing 'job_1'. A 'Required' label is to the right.
- Description:** A text input field with the placeholder text 'Add Description'.
- Job Location:** A text input field containing 'jobs' and a 'Browse' button to its right.
- Max Concurrent Runs:** A text input field containing the number '1'.
- Buttons:** 'Cancel' and 'Create' buttons are located at the bottom right of the form.

3. Provide a name and description for your job.
4. Click **Browse** and select the location to save the job in your AI Data Platform. Click **Select**.
5. Enter a number for **Max Concurrent Runs**.
6. Click **Create**.

Change Job Location

You can change the location of a job after creation.

1. On the Home page, click **Workflow**.
2. Click the job you want to change the location for.
3. Click the **Details** tab.
4. Next to Location, click **Browse**.
5. Pick a new location for the job and click **Select**.

Delete a Job

You can delete jobs that you no longer need.

1. Navigate to your workspace and click **Workflows**.
2. Next to the job you want to delete click **Options** then click **Delete**.
3. Click **Delete**.

Schedule a Job Using a Calendar

Workflow jobs can be scheduled to run on an automated basis.

1. On the Home page, click **Workflow**.
2. Click the job you want to make a schedule for.
3. Click the Details tab.
4. Next to schedule, click **Add**.
5. Choose whether your schedule begins as **Active** or **Paused**.
 - Select **Active** if you want the schedule in effect immediately.
 - Select **Paused** if you want to activate the schedule at a later time.
6. Select the **Time Zone** the schedule uses as a basis.
7. From **Schedule Type**, select **Calendar**.
8. Select whether the schedule will be run hourly, daily, weekly, or monthly. You need to provide additional information for **Hourly**, **Weekly**, and **Monthly** options.
 - For **Hourly**, select the hours on which the schedule repeats.
 - For **Weekly**, select the days of the week the schedule repeats.
 - For **Monthly**, select the days of the month the schedule repeats.
9. Enter the time of day the schedule runs. The time of day is in 24 hour format, beginning at 00:00 and ending at 23:59.
10. Check the listed run time at the bottom to confirm your schedule is correct. Click **Create**.

Schedule a Job Using a Cron Expression

Workflow jobs can be scheduled to run on an automated basis using cron expressions to determine date and times.

1. On the Home page, click **Workflow**.
2. Click the job you want to make a schedule for.
3. Click the Details tab.
4. Next to schedule, click **Add**.
5. Choose whether your schedule begins as **Active** or **Paused**.
 - Select **Active** if you want the schedule in effect immediately.
 - Select **Paused** if you want to activate the schedule at a later time.
6. Select the **Time Zone** the schedule uses as a basis.
7. From **Schedule Type**, select **Cron Expression**.
8. Enter the cron expression.
9. Check the listed run time at the bottom to confirm your schedule is correct. Click **Create**.

Run a Workflow on Demand

You can choose to run a workflow job immediately.

1. On the Home page, click **Workflow**.
2. Click on the job you want to run.
3. Click **Run Now**.

Run a Workflow Job from the Jobs Page

You can quickly run listed jobs directly from the Jobs page.

1. On the home page, click **Workflow**.
2. In the **Jobs** tab, next to the listed job you want to run, click **Run Now**.

☐	Name	Last run	Compute	Trigger	Created by	Created on	Recent Runs
☐	default_job.job	-	-	Manual	awdonald	Fri, Feb 14, 2025 at 17:47:09 UTC	☐ ☐ ☐ ☐ ☐ ☒

Change a Job Run Schedule

You can edit the schedule of a job run after creation to modify time, timezone, or frequency of that schedule.

1. On the Home page, click **Workflow**.
2. Click the job you want to make a schedule for.
3. Click the Details tab.
4. Next your schedule, click **Edit**.
5. Modify the schedule as needed. Check the listed run time at the bottom to confirm your schedule is correct.
6. Click **Save**.

Pause or Activate a Job Run Schedule

After setting a schedule for a job run you can choose to pause it then reactivate when needed.

1. On the Home page, click **Workflow**.
2. Click the job you want to make a schedule for.
3. Click the **Details** tab.
4. Next to your schedule click to **Pause** or **Activate** your schedule.

The option displayed is dependent on the current status of the schedule. If active, **Pause** is displayed. If paused, **Activate** is displayed.

Repair Failed Job Runs

You can attempt to repair a run that has failed by reviewing the timeline and details and rerunning the job with optional parameters to assist your diagnostic.

1. Navigate to your workflow.
2. Click **Job Runs**.
3. Next to the failed job run, click **Actions** and click **Repair run**.
4. Review the task details to determine possible causes of failure.
5. Click **Repair run** and select which tasks to rerun.
6. Add optional parameters that will only apply to this repair run.

- If you select **Key/Value** as the **Parameter** type, click  **Add** and enter parameters.
 - If you select **JSON**, enter parameter values in the space provided.
7. Click **Run repair**.

Configure Tasks

This section covers the creation and configuration of tasks.

Topics:

- [About Tasks](#)
- [Create a Python Task](#)
- [Create a Notebook Task](#)
- [Create a Nested Job Task](#)
- [Create an If/Else Task](#)
- [Create a Jar Task](#)
- [Modify a Task](#)
- [View Task Logs](#)
- [Delete a Task](#)
- [Repair Failed Job Runs](#)

About Tasks

Tasks are short, functional blocks of code you can piece together into a flow as part of job or promote to jobs themselves.

Tasks are the primary building blocks of all workflows in Oracle AI Data Platform. The type of task determines the type of code it uses. As part of a job, you connect tasks to determine their sequence and priority when the job is run.

Task type	Description
Notebook task	A task that has been saved to a notebook you can access
Python task	A task using a snippet of Python programming language
If/else condition	A task that uses if/else conditions
Nested job task	A task that uses an existing job and its tasks as a nested task
Jar task	A task that can run Scala or Java code compiled into Java Archive (JAR) files.

When you have more than one task, you can create sets of task dependencies where the success or failure of one task can trigger subsequent tasks in sequence. You can only create dependencies in jobs that have more than one task. See [Create a Notebook Task](#).

Tasks can be run in parallel with one another. You can do this by making two or more tasks dependent on the success or failure of another task in the same workflow, causing them to run at the same time.

Tasks can fail due to transient issues, such as network disruptions, resource unavailability, or temporary service failures. In these cases, your AI Data Platform automatically retries the task based on retry policies you configure when the task is created. As part of these policies, you define:

- **Retry Count:** The maximum number of retry attempts.
- **Retry Interval:** The wait time between retries.

In addition to standard task retries, AI Data Platform also supports **Retry on timeout**. If a task exceeds its execution time limit due to resource constraints or slow processing and you want to retry only for these scenarios, you can choose to automatically trigger a retry. These retry policies enhance workflow resilience, ensuring tasks have a higher chance of successful execution without manual intervention.

Git Providers as Task Source (Preview)

You can select **Git provider** as a source for Notebook and Python tasks. After selecting **Git Provider** as your source, you need to provide the relative file location, the Git provider (GitHub, Bitbucket, GitLab), your credentials, and the repository URL and branch. To be able to select your credentials, you need to have linked a Git account to your AI Data Platform. For more information, see [Git Linked Accounts \(Preview\)](#).

Note

When using Git Provider as your source, **File Location** is the relative path from the root of the repository. For example, `/GitDemo/01_GitSourceFile.ipynb`.

When and How to Use Compute Logs

You should check your compute logs if your task fails with resource or system related errors, such as out-of-memory errors or CPU use exceeding limits.

Review Spark logs if you see long wait times, unexpected retries, or job performance bottlenecks. These logs provide insight into the driver and worker nodes of the compute cluster backing your task and can help identify the source of possible problems.

For guidance on how to check your logs, see [Monitor a Specific Job Run](#).

You must have appropriate compute-level RBAC permission to view metadata and logs for the compute instance associated with the job. Contact your administrator to obtain these permissions if you are unable to view compute logs. For more information, see [About Permissions](#).

Create a Python Task

You can create a task as part of workflow job that uses Python scripting.

1. On the Home page, click **Workflow**.
2. Click on the job you want to make a task for.
3. Click **Add task**.
4. Enter a task name.
5. For **Task type** select **Python**.
6. Choose a **Source**:

- For **Workspace**, click **Browse** and navigate to the Python script you want to add as a task. Click **Select**.
- For **Git provider**:
 - a. Provide the relative **File location**.
 - b. Select your **Git provider** and **Git credential** from the drop-down menus.
 - c. Provide your **Git repository URL** and **Git branch**.
- 7. Select a compute cluster for the Python task, if one is not already attached.
- 8. Select the number of retries a task should attempt on failure. If you select more than 0, you must also specify how much time the job run should wait between retries and if retries should be attempted on timeout.

9. From **Depends on**, select any tasks you want to make this task dependent on. Select the conditional response to that dependency from the **Run if** dropdown.
10. Add additional parameters by providing their **Key** and **Value**. Click **Add parameter** to provide multiple parameters.

Create a Notebook Task

You create tasks using notebooks you have built in your Oracle AI Data Platform.

1. On the Home page, click **Workflow**.
2. Click on the job you want to make a task for.
3. Click **Add task**.
4. Enter a task name.
5. For **Task type** select **Notebook**.
6. Choose a **Source**:
 - For **Workspace**, click **Browse** and navigate to the Python script you want to add as a task. Click **Select**.
 - For **Git provider**:
 - a. Provide the relative **File location**.
 - b. Select your **Git provider** and **Git credential** from the drop-down menus.
 - c. Provide your **Git repository URL** and **Git branch**.
7. Select a compute cluster for the notebook task, if one is not already attached.
8. Select the number of retries a task should attempt on failure. If you select more than 0, you must also specify how much time the job run should wait between retries and if retries should be attempted on timeout.

9. From **Depends on**, select any tasks you want to make this task dependent on. Select the conditional response to that dependency from the **Run if** dropdown.
10. Add additional parameters by providing their **Key** and **Value**. Click **Add parameter** to provide multiple parameters.

Create a Nested Job Task

You can use another workflow job and its contained tasks as a nested task within another workflow.

1. On the Home page, click **Workflow**.
2. Click on the job you want to make a task for.
3. Click **Add task**.
4. Enter a task name.
5. For **Task type** select **Nested job task**.
6. From the **Jobs** drop-down list, select an existing job you want to make a task from.
7. Select the number of retries a task should attempt on failure. If you select more than 0, you must also specify how much time the job run should wait between retries and if retries should be attempted on timeout.

Retries

Retry 1 time and wait 0 seconds between retries

☐ Retry on timeout

8. From **Depends on**, select any tasks you want to make this task dependent on. Select the conditional response to that dependency from the **Run if** dropdown.
9. Add additional parameters by providing their **Key** and **Value**. Click **Add parameter** to provide multiple parameters.

Create an If/Else Task

You can create a task that uses if/else conditions based on catalog data to determine if the task triggers.

1. On the Home page, click **Workflow**.
2. Click on the job you want to make a task for.
3. Click **Add task**.
4. Enter a task name.
5. For **Task type** select **If/Else**.
6. Enter the conditions that determine if the task triggers. Click **Add** to set multiple conditions.
7. Enter the condition expression.
8. Select the number of retries a task should attempt on failure. If you select more than 0, you must also specify how much time the job run should wait between retries and if retries should be attempted on timeout.

9. From **Depends on**, select any tasks you want to make this task dependent on. Select the conditional response to that dependency from the **Run if** dropdown.
10. Add additional parameters by providing their **Key** and **Value**. Click **Add parameter** to provide multiple parameters.

Create a Jar Task

You can create tasks that run Scala or Java code compiled into Java Archive (JAR) files.

Note

Dependent library files must use JDK, Scala, or Spark versions compatible with the Oracle AI Data Platform cluster runtime at time of creation to avoid unexpected behavior.

1. On the Home page, click **Workflow**.
2. Click on the job you want to make a task for.
3. Click **Add task**.
4. Enter a task name.
5. For **Task type** select **JAR task**.
6. For **Main class name** specify the full name of the class that contains the main method you want to execute. For example, `ProcessTransaction`. This class must be included in one of the files added as a dependent library.
7. For **Dependent libraries**, click **Add**.
8. Select a source for your dependent library files. At least one library that contains the main class method specified above must be included.
 - For **Workspace** or **Volume**, browse your AI Data Platform workspace or volume, select the file or files you want to add as a library and click **Add**.
 - For **Upload file to workspace**, browse your local machine for the file or files to upload as a library and click **Add**.
9. In **Command line arguments**, provide whitespace-delimited arguments that you want to pass to the main class.
10. Select the number of retries a task should attempt on failure. If you select more than 0, you must also specify how much time the job run should wait between retries and if retries should be attempted on timeout.

11. From **Depends on**, select any tasks you want to make this task dependent on. Select the conditional response to that dependency from the **Run if** dropdown.

12. Add additional parameters by providing their **Key** and **Value**. Click **Add parameter** to provide multiple parameters.

Modify a Task

You can change existing attributes of a task, such as name, type, and parameters to alter how it functions in your job.

1. On the Home page, click **Workflow**.
2. Click on the job you want to configure tasks for.
3. In the **Tasks** tab, click the task you want to edit.
4. In the **Task Details** pane on the right, modify the task attributes as needed. Changes are saved automatically.

View Task Logs

You can view the run logs of individual tasks in a job.

1. On the Home page, click **Workflows**.
2. Click **Job Runs**.
3. Click the job you want to see the task logs for.
4. Click the task nodes to see the logs for that task.

Delete a Task

You can delete a task by removing the task node from a job.

1. On the Home page, click **Workflow**.
2. Click on the job you want to delete tasks from.
3. On the task node, click **Actions** then click **Remove node**.
4. Click **Delete**.

Parameterization

This section covers parameters and how they are used in jobs and tasks.

Topics:

- [About Parameters](#)
- [Pass Parameters Between Tasks and Notebook](#)
- [Add Parameters to a Job](#)
- [Delete Parameters from a Job](#)
- [Add Parameters to a Task](#)
- [Add Parameters to a Python Task](#)
- [Delete Parameters from a Task](#)
- [Run a Job with Different Parameter Values](#)

About Parameters

You can customize job runs by passing parameters that alter the behavior of the job, task, or job run.



Parameters can be provided at three different levels of a workflow: job level, task level, and job run level. Parameters take the following precedence in the case of conflicts: Job Run > Task > Job.

1. **Job Run Level Parameters:** Job run level parameters specified in the job run payload take precedence over task-level parameters defined in the job configuration. This means that any parameters specified during job run override the task-specific defaults.
2. **Task Level Parameters:** If job run parameters for a specific task are not provided, the task uses the parameters defined at the task level in the job configuration. Job-level parameters with the same name override task-level parameters for that specific task. So if you want a task specific parameter, you should name it differently than the job parameter.
3. **Job Level Parameters:** If neither task-level nor runtime parameters are provided, the default values set at the job level apply. Job-level parameters are considered defaults that can be used if no more specific parameters are available.

Note

Job parameters are immutable in task contexts. This means if there is a job with parameter *JobParamA* with resolved value *JobParamRuntimeValueA* then TaskA execution cannot change the value of *JobParamA*. The value of *JobParamA* remains *JobParamRuntimeValueA* for all tasks and the entire job run. As a result, if you wanted to share information between tasks you could use intermediate storage or output parameters to achieve that.

When task names, task value keys, or job parameter names contain special characters (such as !@\$%), you must surround those identifiers with backticks (` `). Only alphanumeric and underscore characters are usable without surrounding the identifier in backticks.

For example:

```
{
  "VariableWithSpecialChars": "{{job.parameters.`param$@`}}"
}
```

System Parameters are templated parameters whose value is provided by the system as part of workflow runs and subsequent task runs. You don't have to provide any value, default or otherwise, for these templated parameters. Oracle AI Data Platform has a fixed list of valid templated parameters / dynamic value references that are supported in workflow. System parameters are entered by surrounding them with two curly brackets. For, example `{{job.id}}`.

Table 8-1 Supported System Parameters

Parameter	Description
{{hub.id}}	The unique identifier assigned to the hub
{{hub.region}}	The region of the hub
{{workspace.id}}	The unique identifier assigned to the workspace
{{workspace.url}}	The URL of the workspace
{{job.id}}	The unique identifier assigned to the job
{{job.name}}	The name of the job at the time of the job run
{{job.run_id}}	The unique identifier assigned to the job run
{{job.repair_count}}	The number of repair attempts on the current job run
{{job.start_time.[argument]}}	A value based on the time (in UTC timezone) that the job run started. The return value is based on the argument option. See Options for date and time values.
{{job.parameters.[name]}}	The value of the job-level parameter with the key [name]
{{job.trigger.type}}	The trigger type of the job run. The possible values are Manual and Scheduled.
{{job.trigger.file_arrival.location}}	If a file arrival trigger is configured for this job, the value of the storage location
{{job.trigger.time.[argument]}}	A value based on the time (in UTC timezone) that the job run was triggered, rounded down to the closest minute for jobs with a cron schedule. The return value is based on the argument option. See Options for date and time values.
{{task.name}}	The name of the current task
{{task.run_id}}	The unique identifier of the current task run
{{task.execution_count}}	The number of times the current task was run (including retries and repairs)
{{task.notebook_path}}	The notebook path of the current notebook task
{{tasks.[task_name].run_id}}	The unique identifier assigned to the task run for [task_name]
{{tasks.[task_name].result_state}}	The result state of task [task_name]. The possible values are success, failed, excluded, canceled, skipped, timed out, upstream_canceled, and upstream_failed.
{{tasks.[task_name].error_code}}	The error code for task [task_name] if an error occurred running the task. Examples of possible values are RunExecutionError, ResourceNotFound, and UnauthorizedError. For successful tasks, this evaluates to an empty string.
{{tasks.[task_name].execution_count}}	The number of times the task [task_name] was run (including retries and repairs)
{{tasks.[task_name].notebook_path}}	The path to the notebook for the notebook task [task_name]
{{tasks.[task_name].values.[value_name]}}	The task value with the key [value_name] that was set by task [task_name]

Table 8-2 Options for Date and Time

Argument	Description
iso_weekday	Returns a digit from 1 to 7, representing the day of the week in the time stamp
is_weekday	Returns <code>true</code> if the timestamp is on a weekday
iso_date	Returns the date in ISO format
iso_datetime	Returns the date and time in ISO format
year	Returns the year part of the timestamp
month	Returns the month part of the timestamp
day	Returns the day part of the timestamp
hour	Returns the hour part of the timestamp
minute	Returns the minute part of the timestamp
second	Returns the second part of the timestamp
timestamp_ms	Returns the timestamp in milliseconds

Pass Parameters Between Tasks and Notebook

You can pass parameters from a task to a notebook and vice versa. This enables dynamic workflow behavior, allowing notebooks to adjust their processing based on runtime values.

The `oidlUtils.parameters` package provides the necessary functionality to handle these parameter operations. The `oidlUtils` package is a utility library in Oracle AI Data Platform that simplifies tasks such as parameter management, task value passing, and other workflow operations. It is commonly used in notebooks and tasks to get and set parameters across workflow stages.

Task key values are either strings or JSONs. For example, to use a string as your task key value, your parameter call would look like this:

```
oidlUtils.parameters.setTaskValue(key="payload", value="abc", "defaultValue")
```

To use a JSON as your task key value, your parameter call would look like this:

```
oidlUtils.parameters.setTaskValue(key="payload", value=json.dumps(payload),
"defaultValue")
```

Example Workflow: Passing Parameters

In this scenario, we have two notebooks in a workflow. Notebook 1 receives parameters from a task, processes them, and sets output parameters that are passed to Notebook 2 in the next task.

Notebook 1: Get and Set Parameters

```
# Get parameter if already set in the task
param_key = "param1"
param_value = oidlUtils.parameters.getParameter(param_key, "defaultValue")
print(param_value)
print("Param {} value is {}".format(param_key, param_value))
```

```
# Set parameter value in the task
output_parameter_key = "output_parameter"
output_param_value = oidlUtils.parameters.getParameter(output_parameter_key,
"defaultValue2")
print("Param {} value is {}".format(output_parameter_key, output_param_value))
oidlUtils.parameters.setTaskValue(output_parameter_key, "1234")
```

The first notebook retrieves a parameter (param1) passed from the task and then sets a new parameter (output_param_2), which will be used in the next task.

Notebook 2: Read the Output Parameter

```
output_param_2= "output_parameter"
param_value = oidlUtils.parameters.getTaskValue("GetSetParameter",
output_param_2, "defaultValue")
print("Param {} value is {}".format(output_param_2, param_value))
```

The second notebook receives the output_param_2 from Notebook 1 via the workflow task and processes it.

1. Task 1: Notebook 1

- In the first task, parameters can be passed to Notebook 1 from the job or task itself.
- Notebook 1 processes the parameters as input parameters param1 and sets new output parameters (e.g., output_param_2).

2. Task 2: Notebook 2

- In the second task, Notebook 2 receives the output parameter from Task 1 by referencing it directly in the notebook code as shown above by passing the name of first task "GetSetParameter" as set in the workflow.
- The value of output_param_2 is passed into Notebook 2 where it can be used for further processing.

This approach makes it easy to pass values dynamically between tasks and notebooks, allowing your workflows to be more flexible and adaptive.

Add Parameters to a Job

You can provide your jobs different parameters to track.

1. On the Home page, click **Workflow**.
2. Click the job you want add parameters to and click the **Details** tab.
3. Under Job Parameters, provide the key and value to track. To add multiple parameters, click  **Add Parameter**.

Changes you made are saved automatically.

Add Parameters to a Python Task

You can pass parameters to a Python task as command-line arguments. Use this approach when your Python script needs fixed input values, job parameters, or values produced by an earlier task.

You can provide parameters as fixed values or dynamic values:

Fixed Values

Fixed values are literal values that you enter directly in the parameter list. For example, the following parameter list uses fixed values:

```
["123", "26-03-2026", "26-03-2026", "1"]
```

Dynamic Values

Dynamic values are runtime values that use workflow parameter references enclosed in double curly braces. The following parameter list uses dynamic values:

```
[
    "{{job.parameters.repair_request_id}}",
    "{{job.parameters.need_by_date}}",
    "{{tasks.PredictSupplierPartDelivery.values.predicted_ready_date}}",
    "{{job.parameters.lead_time_buffer_days}}"
]
```

In the dynamic example:

- `repair_request_id` and `need_by_date` are job-level parameters.
- `predicted_ready_date` is a task value produced by the `PredictSupplierPartDelivery` task.
- `lead_time_buffer_days` is a job-level parameter.

Your Python script can read parameter values in the same order that they are passed to the task.

```
first_parameter = sys.argv[0]
second_parameter = sys.argv[1]
third_parameter = sys.argv[2]
fourth_parameter = sys.argv[3]
```

Delete Parameters from a Job

You update your jobs to remove parameters that are no longer needed.

1. On the Home page, click **Workflow**.
2. Click the job you want remove parameters from and click the **Details** tab.
3. Under Job Parameters, click  **Delete** next to the parameter you want to remove.

Changes you made are saved automatically.

Add Parameters to a Task

You can add parameters to tasks to alter their behavior.

1. On the Home page, click **Workflow**.
2. Click the job that contains the tasks you want add parameters to and click the **Tasks** tab.
3. Click the task you want to add parameters to.

4. Under **Parameters**, enter the key and value for your parameter. To add multiple parameters, click  **Add Parameter**.

Changes you made are saved automatically.

Delete Parameters from a Task

You can remove parameters you don't need from your tasks.

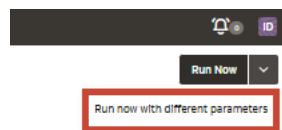
1. On the Home page, click **Workflow**.
2. Click the job that contains the tasks you want delete parameters from and click the **Tasks** tab.
3. Click the task you want to delete parameters from.
4. Click  **Delete** next to each parameter you want removed.

Changes you made are saved automatically.

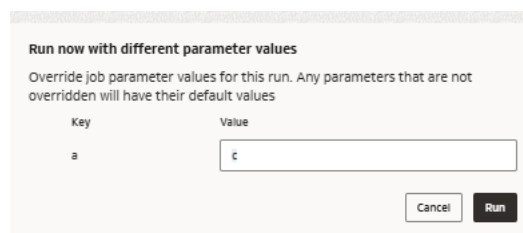
Run a Job with Different Parameter Values

You can choose to run a job immediately with modified parameter values.

1. On the Home page, click **Workflow**.
2. Click on the job you want to run.
3. Click the down arrow next to **Run Now**. Click **Run now with different parameters**.



4. Enter new values. These parameter values only apply to this job run.

A screenshot of a dialog box titled 'Run now with different parameter values'. It contains the text 'Override job parameter values for this run. Any parameters that are not overridden will have their default values'. Below this is a table with two columns: 'Key' and 'Value'. The 'Key' column has the value 'a'. The 'Value' column has a text input field containing the value 'c'. At the bottom right of the dialog are 'Cancel' and 'Run' buttons.

Key	Value
a	c

5. Click **Run**.

Monitor Jobs

This section covers how to track and search job runs in your Oracle AI Data Platform.

Topics:

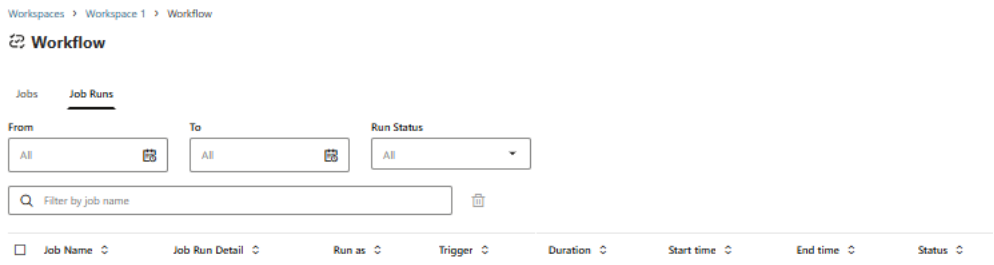
- [About Jobs Runs](#)
- [View All Workflow Job Runs](#)

- [View a Job's Run History](#)
- [Monitor a Specific Job Run](#)

About Jobs Runs

You can monitor the status of previous job runs, job runs in progress, or job runs that have been interrupted from the Job Runs page.

You can find a history of job runs in Job Runs tab of your job.



All jobs that have been run on your Oracle AI Data Platform are listed in the Job Runs page. The job run information is split into columns tracking different aspects of each job run. You can sort those columns in ascending or descending order by clicking the column header. You can filter the job runs displayed by using the dropdown menus, using the Search bar to filter by name, or combining multiple filter options.

Column name	Description
Job Name	Name of the job
Job Run Detail	Click View to see the details page of the job
Run as	User or role that triggered the run
Trigger	Trigger type. Either Manual or Scheduled
Duration	Length of time taken to run the job
Start time	Start time
End time	End time
Status	Current job status

The **Status** displays the current state of a job, using one of the values from the following list.

Status	Description
Success	Job run was successful
Failed	Job run failed
Canceled	Job run was canceled
Skipped	Job run was canceled because a previous run of the same job was already active
Timed Out	Job run exceeded the configured time limit and was stopped

View All Workflow Job Runs

You can see the recent history of your run workflow jobs from the **Job Runs** page.

1. On the Home page, click **Workflows**.
2. Click **Job Runs**.
3. Optional: Filter the results by selecting date ranges, run status, job name or any combination of filters to find the workflow job runs you need.

View a Job's Run History

You can see records of all previous runs for a job and filter the results for greater detail.

1. On the Home page, click **Workflow**.
2. Click the job you want to view the run history for.
3. Click the **Runs** tab.
4. Optional: Filter the results by selecting date ranges or run status from the dropdowns.

Monitor a Specific Job Run

You can track the status and history of a specific job run from the Job Runs page.

1. From the Job Runs page, click on **View Job Run Detail** for the job run you want to inspect.
2. Click the Details tab to review run-level metadata such as job parameters, compute configuration, schedule, and start/end time.
3. Use the buttons at the bottom of the Details page to see more specific information about the job run.
 - Click **View Details** to inspect compute configurations, like driver and worker shape.
 - Click **Spark UI** to inspect Spark stages, tasks, and resource usage for the run.
 - Click **Logs** to view driver and worker logs and review errors, warnings, and other runtime messages.
 - Click **Metrics** to monitor additional metrics related to compute, like CPU and memory usage.

To learn more about monitoring and troubleshooting activity at the compute level, see [Monitor Compute](#).

Workflow Jobs: Run As

You can run a workflow job in Oracle AI Data Platform using the Run As command to determine which user identity is used.

The identity you select when you use Run As governs two things during execution: whose permissions apply, and which credential is used to connect to downstream systems (such as Git repositories or other integrated services).

By default, a workflow could run under the personal identity of whoever built it which creates risk. If that person leaves the organization or their access changes, the workflow can silently break.

Run As solves this by letting you execute workflows under a secure, governed identity instead. This means:

- Scheduled and automated workflows aren't tied to any one person's login.
- Workflows keep running even if the original creator is no longer with the organization.
- Execution identity is explicit, auditable, and controlled by permissions.

Types of Identity

When configuring a job using Run As, you can choose from two types of identity: your own identity and a service account.

Your own identity: AI Data Platform checks whether you have a saved credential in the Credential Store. If one exists, it's used automatically. If not, execution falls back to your standard organizational (IAM) login.

A service account: Service accounts stored in the Credential Store can also be selected as the Run As identity but only if you've been granted at least USE permission on that service account's credential.

Situation	Result
You have no credential saved in the Credential Store	Job runs using your organizational login. You'll see a warning that the scheduled job may fail if your login session expires (typically after 24 hours).
You have a credential saved in the Credential Store	Job runs using that credential and continues to work reliably, with no expiration.
You add a credential to the Credential Store after creating a job	The next run automatically switches to using the new credential. If you remove it later, the job falls back to your organizational login.
You switch Run As to a service account you have access to	Scheduled runs use the service account's credential, with no expiration.

Ownership in Run and Edit Behavior

Run As also protects against a workflow running under the an incorrect identity when someone other than the original owner interacts with it.

When someone else runs your job: If you grant another user at least USE permission on a job you created, and they run it, Run As switches to them. The job now executes under their login or their saved credential.

When someone else edits your job: If you grant another user at least MANAGE permission and they make changes to the job, Run As switches to them — including for jobs that already have a schedule set. From that point forward, scheduled runs use their credential.

When a Run As is set to a service account: If you've set Run As to a service account, and another user with MANAGE permission later edits the job, the outcome depends on their access to that service account:

- If they also have at least USE permission on the service account, Run As stays on the service account.
- If they don't have at least USE permission on the service account, Run As switches to them personally, following the same behavior as above.

By using Run As identity, service accounts can be leveraged as a common credential for multiple users accessing or sharing a workflow job.

Git Repository Access

For jobs that sync a Git repository, AI Data Platform always uses the Git credentials belonging to the current Run As identity. The Git credentials of another user are never used. If a job runs as a specific user, only the Git credentials of that user are used to pull the repository. If a job runs as a service account, only that service account's Git credentials are used.

Note

If needed, you can use AI Data Platform REST APIs to set credentials for your service account. See [REST API for Oracle AI Data Platform](#).

Reliability

AI Data Platform is confirmed to resolve the correct Run As identity for every job run without adding measurable delay to the job start time.

Scheduled jobs using saved credentials or service account run without failing due to login expiration.

Clear warnings are issued within seconds in circumstances where your credentials are missing and AI Data Platform falls back to your organizational login.

Run As Monitoring

You can monitor the identity of users for job runs by viewing the job run history. For more information, see [View a Job's Run History](#).

Edit Run As Identity for a Job

You can change the identity a job is run as going forward after the job is created from the Details tab.

1. On the Home page, click **Workflow**.
2. Click the job you want to change the run identity for.
3. Click the Details tab.
4. Under **Run As**, click **Edit**.
5. Select your identity or the service account you want to run the job as. Select the service account you want to use from the dropdown menu.
6. Click **Save**.
7. Optional: Click **Run now** to run the job with the selected identity immediately.

9

Compute

This chapter covers the use of computing resources in Oracle AI Data Platform workspaces.

Topics:

AI Data Platform compute clusters are computing resources available in the workspace. You can use compute clusters to run workloads, such as data engineering, data science, and data analytics workloads.

You can use existing compute or, if you have the necessary permissions, create new compute to support your workloads. You can view the compute you have access to using the Compute section of the workspace.

- [Manage Compute](#)
- [Connect to Compute](#)
- [Manage Libraries](#)
- [Monitor Compute](#)

Manage Compute

This section covers the basic functions of creating, changing, or removing compute clusters in your Oracle AI Data Platform.

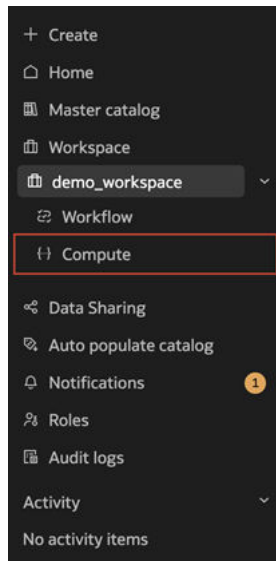
Topics:

- [About Compute Clusters](#)
- [Create a Quickstart Cluster](#)
- [Create a Custom Cluster](#)
- [Create an NVIDIA GPU Cluster](#)
- [Tune an NVIDIA GPU Cluster for Apache Spark](#)
- [Modify a Cluster](#)
- [Delete a Cluster](#)
- [View Cluster Details](#)
- [Maintenance Updates for Compute Clusters](#)

About Compute Clusters

All-purpose compute clusters provide you the compute resources to process your workloads in an Oracle AI Data Platform instance.

You manage your compute clusters from the Compute page in your AI Data Platform.



Types of Compute

Two types of compute exist in your AI Data Platform: all-purpose compute clusters and Default Master Catalog Compute Cluster.

You can only create all-purpose compute clusters in your AI Data Platform. All-purpose compute clusters are suitable for a versatile range of workloads and can be attached to your notebooks and used in workflows. Unless otherwise specified, any references to 'compute cluster' or 'cluster' in documentation refer to all-purpose compute clusters.

When you create a new all purpose compute cluster, you can choose either the Quickstart or Custom configuration. The Quickstart configuration is optimized to provide fast startup, while Custom configuration allows you to fine-tune your all purpose compute cluster to suit the specific workloads you need it to process. In both Quickstart and Custom configuration options, you can view cost projections and modify idle timeout options.

Note

Installing custom libraries to a Quickstart configured all purpose compute cluster automatically changes it to the Custom configuration. This can impact startup performance.

Default Master Catalog Compute Cluster is present in all AI Data Platform instances. This cluster is responsible for essential AI Data Platform functions, like search crawls, refreshing catalog objects, creating, editing, and deleting objects, and testing connections.

Cluster Runtime

All-purpose compute clusters can be created with an Apache Spark 3.5 runtime. The runtime environment is compatible with:

- Spark 3.5.0
- Delta 3.2.0 (pre-included)
- Python 3.11
- Scala 2.12

- Hadoop 3.3.4
- Java 17

Compute Configuration Export and Import

Once you have created a compute cluster, you can export the settings you configured for that cluster to your workspace or a volume in a catalog to use again for other clusters. You can import your saved configurations to other created clusters or import them as part of the creation process.

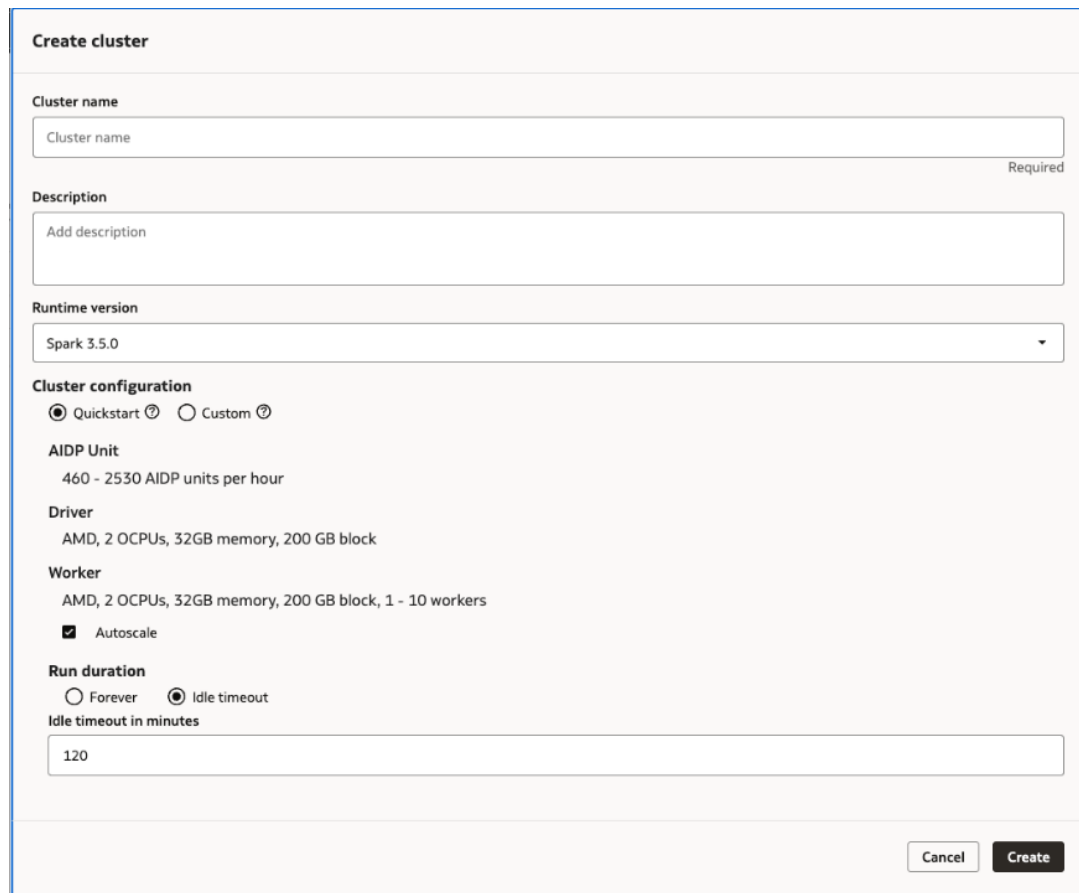
Your configurations are available to other users with access to workspace or volume you saved the configuration to, making it easy to share configurations across teams.

Create a Quickstart Cluster

You can choose to create an all-purpose compute cluster with preconfigured settings to process data and AI workloads in your Oracle AI Data Platform.

The quickstart configuration is an Apache Spark cluster with 1 driver and up to 10 workers, each with AMD 2 OCPU and 32 GB memory. Autoscaling is enabled by default for quickstart configuration. You can set your clusters to be constantly active or you can set an interval of inactivity after which the cluster will automatically stop (idle timeout). Stopped clusters will resume when called on by an attached workflow or notebook. You can edit your cluster at any time after creation.

1. Click **Create** in the left navigation panel then click **Compute**. You can also navigate to your workspace and click **Compute**, then click  **Create Cluster**.



Create cluster

Cluster name
Cluster name Required

Description
Add description

Runtime version
Spark 3.5.0

Cluster configuration
☒ Quickstart  ☐ Custom 

AIDP Unit
460 - 2530 AIDP units per hour

Driver
AMD, 2 OCPUs, 32GB memory, 200 GB block

Worker
AMD, 2 OCPUs, 32GB memory, 200 GB block, 1 - 10 workers

☒ Autoscale

Run duration
☐ Forever ☒ Idle timeout

Idle timeout in minutes
120

Cancel Create

2. Provide a name and description to identify your cluster.
3. Select **Runtime version**.
4. Select **Quickstart** as your cluster configuration.
5. Select whether the number of workers is static or scales automatically. Autoscaling is enabled by default for quickstart configuration.
6. For **Run duration**, select whether the cluster will stop running after a set duration of inactivity. If **Idle timeout** is selected, specify the idle time, in minutes, before the cluster will time out.
7. Click **Create**.

Create a Custom Cluster

You can create an all-purpose compute cluster with configuration settings of your own choosing to process data and AI workloads in Oracle AI Data Platform.

Custom clusters are intended for advanced users who want to leverage the full range of configuration options to suit their needs. You should select the driver and worker options that best fit the workloads you are going to process. You can set your clusters to be constantly active or you can set an interval of inactivity after which the cluster will automatically stop (idle timeout). Stopped clusters will resume when called on by an attached workflow or notebook. You can edit your cluster at any time after creation.

1. Click **Create** in the left navigation panel then click **Compute**. You can also navigate to your workspace and click **Compute**, then click  **Create Cluster**.

Create cluster

Cluster name

Cluster name

Required

Description

Add description

Runtime version

Spark 3.5.0

Cluster configuration

☐ Quickstart [Ⓢ]
☒ Custom [Ⓢ]

AIDP Unit

0 - 0 AIDP units per hour

Driver

Driver shape

Select Driver Shape

OCPUs

Select number of OCPUs

Memory (GB)

Select Memory

Block volume (GB)

Worker

Worker shape

Select Worker Shape

OCPUs

Select number of OCPUs

Memory (GB)

Select Memory

Block volume (GB)

Set number of workers

☐ Static amount
☒ Autoscale

Minimum number of workers

1

Maximum number of workers

10

Run duration

☐ Forever
☒ Idle timeout

Idle timeout in minutes

120

Advanced options

- Adding Spark configurations and Init script will result in slower startup

Cancel

Create

2. For **Create method**, select **Configure a new cluster**.
3. Provide a name and description to identify your cluster.
4. Select **Runtime version**.
5. Select **Custom** as your cluster configuration.
6. Select the driver options for your cluster.
7. Select the worker options for your cluster. These options apply to all cluster workers.
8. Select whether the number of workers is static or scales automatically.
 - If **Static amount**, specify the number of workers.
 - If **Autoscale**, specify the minimum and maximum number of workers the cluster can scale to.
9. For **Run duration**, select whether the cluster will stop running after a set duration of inactivity. If **Idle timeout** is selected, specify the idle time, in minutes, before the cluster will time out.
10. Click **Create**.

Create a Cluster Using an Imported Configuration

You can create an all-purpose compute cluster with configuration settings exported from an existing compute cluster in your Oracle AI Data Platform.

Importing settings from an exported YAML file allows you to quickly replicate clusters with the same settings. You can import YAML files from AI Data Platform workspaces or volumes shared with you.

You can edit your cluster at any time after creation.

1. Click **Create** in the left navigation panel then click **Compute**. You can also navigate to your workspace and click **Compute**, then click  **Create Cluster**.
2. For **Create method**, select **Use an exported configuration**.
3. From the **Configuration file** drop-down menu, select the exported configuration to use for your new cluster. Runtime, library, and environment settings are prefilled by the imported file.
4. Provide a name and description to identify your cluster.
5. Modify the remaining settings as needed.
6. Click **Create**.

Create an NVIDIA GPU Cluster

You can choose to use an NVIDIA GPU in an All Purpose Compute Cluster to accelerate any workload in your unified AI and data pipeline.

NVIDIA GPU shapes use the following configurations:

Table 9-1 NVIDIA GPU Shapes

GPU Count	OCPU	Block storage (GB)	GPU memory (GB)	CPU memory (GB)
1	15	1500	24	240
2	30	3000	48	480

Note

When you use NVIDIA GPU shapes, both the Driver and Worker shape must be an NVIDIA GPU. Mixing CPU and GPU shapes for the same cluster is currently not supported.

1. Click **Create** in the left navigation panel then click **Compute**. You can also navigate to your workspace and click **Compute**, then click  **Create Cluster**.

Create cluster

Cluster name
Cluster name Required

Description
Add description

Runtime version
Spark 3.5.0

Cluster configuration
☐ Quickstart [Ⓢ] ☒ Custom [Ⓢ]

AIDP Unit
0 - 0 AIDP units per hour

Driver
Driver shape
Select Driver Shape

OCPUs **Memory (GB)**
 Select number of OCPUs Select Memory

Block volume (GB)

Worker
Worker shape
Select Worker Shape

OCPUs **Memory (GB)**
 Select number of OCPUs Select Memory

Block volume (GB)

Set number of workers
☐ Static amount ☒ Autoscale

Minimum number of workers **Maximum number of workers**
 1 10

Run duration
☐ Forever ☒ Idle timeout
 Idle timeout in minutes
 120

✓ **Advanced options** - Adding Spark configurations and Init script will result in slower startup

Cancel Create

2. Provide a name and description to identify your cluster.
3. Select **Runtime version**.
4. Select **Custom** as your cluster configuration.
5. For your cluster driver options:
 - Select **NVIDIA GPU** as the Driver Shape.
 - Select 1 or 2 as the GPU count.
6. For your cluster worker options:
 - Select **NVIDIA GPU** as the Worker Shape.
 - Select 1 or 2 as the GPU count.
7. Select whether the number of workers is static or scales automatically.
 - If **Static amount**, specify the number of workers.
 - If **Autoscale**, specify the minimum and maximum number of workers the cluster can scale to.
8. For **Run duration**, select whether the cluster will stop running after a set duration of inactivity. If **Idle timeout** is selected, specify the idle time, in minutes, before the cluster will time out.
9. Click **Create**.

Tune an NVIDIA GPU Cluster for Apache Spark

You can tune your NVIDIA GPU clusters to optimize their performance by using recommendations from the GPU provider and by installing optional libraries.

Tuning GPU clusters can help optimize the performance of those clusters when called on by jobs in your Oracle AI Data Platform.

For NVIDIA GPU-based clusters, you can follow NVIDIA's [Tuning Guide](#) for recommendations and steps you can take to optimize performance.

You also have the option of installing NVIDIA cuDF and cuML plug-ins for Apache Spark libraries to assist with optimization:

- [cuDF plugin for Apache Spark](#) library is an accelerator for Apache Spark and provides a set of plugins that leverage GPUs to accelerate processing.
- [cuML plugin for Apache Spark](#) library enables GPU-accelerated, distributed machine learning on Apache Spark and provides several PySpark ML compatible algorithms powered by the cuML library.

The cuDF plugin library is commonly used first for feature engineering and data cleaning, and then cross validation is performed at scale using the cuML plugin library. You can use these libraries for use cases like fraud detection (time series), web clickstream, and A/B experimentation.

Table 9-2 Recommended Spark Configurations

Setting	Value	Note
spark.task.resource.gpu.amount	0.0625	1 / spark.executor.cores
spark.rapids.sql.concurrentGpuTasks	3	GPU memory / 8GB, maximum of 4
spark.rapids.shuffle.multiThreaded.writer.threads	32	CPU cores / GPU count per worker
spark.rapids.shuffle.multiThreaded.reader.threads	32	CPU cores / GPU count per worker
spark.shuffle.manager	com.nvidia.spark.rapids.spark350.RapidsShuffleManager	-
spark.rapids.shuffle.mode	MULTITHREADED	-
spark.plugins	com.nvidia.spark.SQLPlugin	-
spark.executor.resource.gpu.amount	1	-
spark.sql.files.maxPartitionBytes	2 GB	Optional, recommended for large datasets
spark.rapids.sql.batchSizeBytes	2 GB	Optional, recommended for large datasets
spark.rapids.memory.host.spillStorageSize	32 G	Optional, recommended for large datasets
spark.rapids.memory.pinnedPool.size	8 G	Optional, recommended for large datasets
spark.sql.adaptive.coalescePartitions.minPartitionSize	32 MB	Optional, recommended for large datasets
spark.sql.adaptive.advisoryPartitionSizeInBytes	160 MB	Optional, recommended for large datasets

Table 9-2 (Cont.) Recommended Spark Configurations

Setting	Value	Note
spark.rapids.filecache.enabled	True	Optional, recommended if workloads will be reusing datasets

Export Compute Configuration

You can export your compute cluster configuration and dependencies to a location in your workspace so it can be imported to other compute clusters.

1. Navigate to your workspace and click **Compute**.
2. Next to the compute cluster you want to export settings for, click ... **Actions** then click **Export compute configuration**. You can also click the cluster name, click **Actions** in the top-right, then click **Export compute configuration**.
3. Verify the dependencies, such as libraries and environment variables. Deselect any that should not be exported.
4. Choose a save location in your workspace or in a volume.
5. Click **Export**.

Import Compute Configuration

You can import existing compute cluster configurations and their dependencies to an all-purpose cluster using YAML configuration files in your workspace.

1. Navigate to your workspace and click **Compute**.
2. Next to the compute cluster you want to import settings for, click ... **Actions** then click **Import compute configuration**. You can also click the cluster name, click **Actions** in the top-right, then click **Import compute configuration**.
3. Search your workspace or volume for the configuration YAML files you want to upload and select one or more.
4. Review the configuration files you selected then click **Import**.

Modify a Cluster

You can change settings or add additional parameters for your clusters.

1. Navigate to your workspace and click **Compute**.
2. Next to the compute cluster you want to modify, click ... **Actions** then click **Edit**.
3. Modify the attributes of your compute cluster or add additional parameters as needed.
4. Click **Save**.

Delete a Cluster

You can delete compute clusters that are unused or no longer needed.

1. Navigate to your workspace and click **Compute**.

2. Next to the cluster you want to delete, click ... **Actions** and click **Delete**.
3. Click **Delete**.

View Cluster Details

You can review the shape and settings of a cluster at any time.

1. Navigate to your workspace and click **Compute**.
2. Click the name of the cluster you want to view details for.
3. Click the **Details** tab.

Maintenance Updates for Compute Clusters

Oracle AI Data Platform compute automatically applies maintenance updates without user intervention.

The maintenance updates cover any necessary security patches or bug fixes for operating system and AI Data Platform internal components. AI Data Platform verifies there are no running clusters before applying these monthly maintenance updates.

Manage Libraries

This section covers how to use and manage libraries connected to your compute clusters.

Topics:

- [Libraries](#)
- [Notebook Scoped Libraries](#)
- [Install a Library from a Workspace or Volume](#)
- [Install a Library from an Uploaded File](#)
- [Uninstall a Library](#)
- [Notebook Scoped Libraries](#)

Libraries

You can add cluster scoped libraries to make third-party or custom code available to your compute clusters while running notebooks or workflow jobs.

Cluster scoped libraries can be installed to extend the out of the box capabilities of compute clusters and applies to all notebooks and workflow jobs using that cluster. For example, visualization options, connectivity options (e.g. JDBC JARs), extractions (e.g. extracting text from PDF) or transformations.

The option for installing cluster scoped libraries is available in the Library tab of your cluster after the cluster status changes to Active. Your library file should be a .jar file or a Wheel (*.whl) file or a requirements.txt file.

You can also add initialization scripts during the creation of a cluster or by modifying an existing cluster. For more information, see [Modify a Cluster](#).

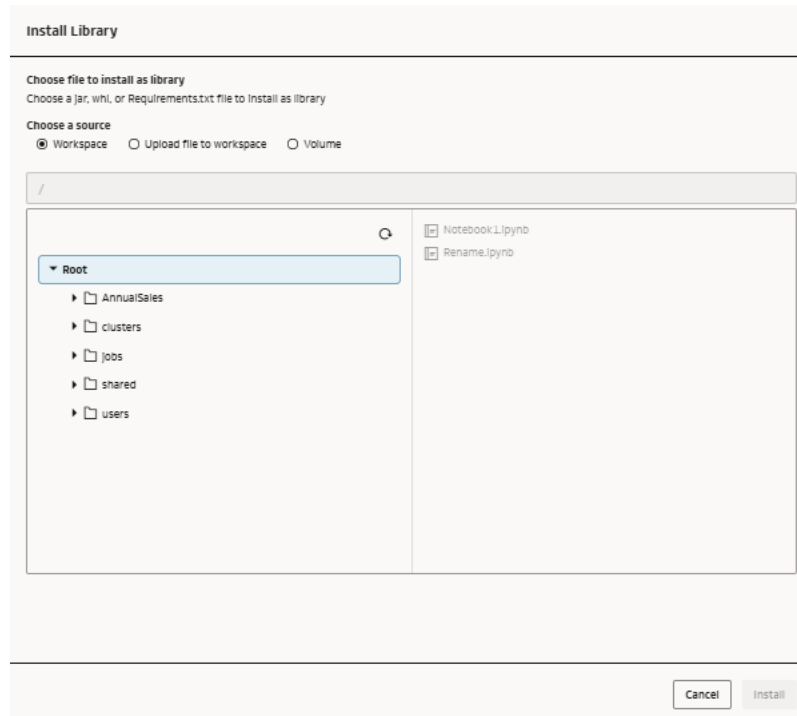
Install a Library from a Workspace or Volume

You can install a library that is in your workspace or volume to expand cluster options for attached notebooks and jobs.

Libraries can only be added from a workspace or a volume where you have appropriate permissions. You can view libraries that are installed on a cluster at any time from the cluster's Library tab.

If the library file you want to install is not already available in your workspace or volume, you can upload the library from your local machine to your workspace first and then install at the cluster.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Library** tab.
3. Click  **Install Library**.
4. Select whether your library is part of a **Workspace** or **Volume**.



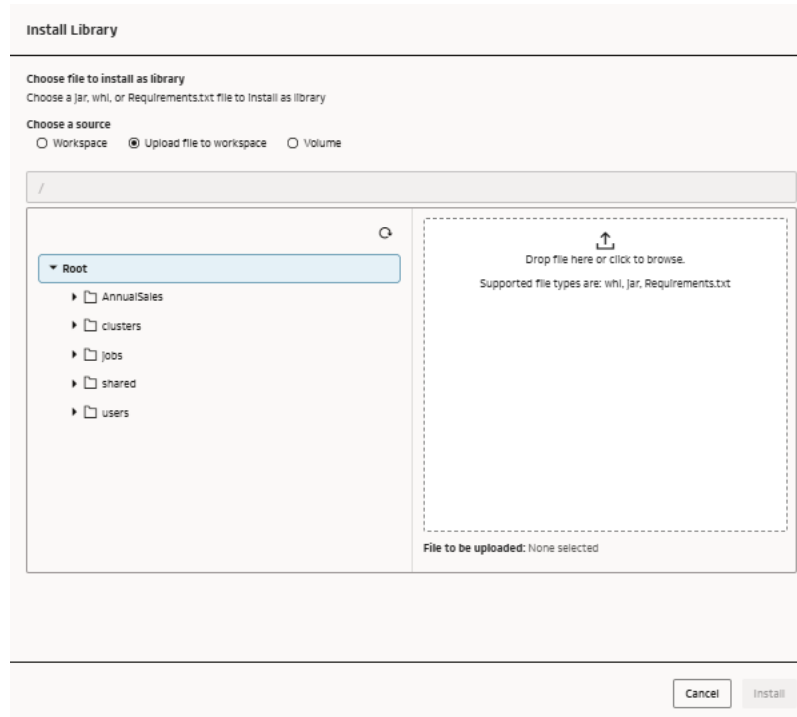
5. Navigate to the library and select it. Click **Install**.
6. Once the library is installed, restart the cluster by clicking **Actions**, then **Restart**.

When the cluster status is Active again, you can use the library in your code inside a notebook or workflow job.

Install a Library from an Uploaded File

You can install a library to your workspace from an uploaded file to expand cluster options for attached notebooks and jobs.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Library** tab.
3. Click  **Install Library**.
4. Select **Upload file to workspace**.



5. Browse to the file that contains your library or drag and drop it into the window.
Your library file must be a .whl or .jar format or a text file with the name requirements.txt. For more information on the requirements.txt file, see [Requirements File Format](#).
Here is an example of a requirements.txt file:

```
plotly==6.0.1  
pandas==2.2.3  
matplotlib==3.10.1
```

6. Click **Install**.
7. Once the library is installed, restart the cluster by clicking **Actions**, then **Restart**.
When the cluster status is Active again, you can use the library in your code inside a notebook or workflow job.

Uninstall a Library

You can uninstall an unwanted or no longer needed library from clusters you own.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Library** tab.
3. Next to the library, click  **Actions** then click **Uninstall**.

4. Click **Uninstall**.

Notebook Scoped Libraries

You can use notebook scoped libraries in Oracle AI Data Platform for creating custom Python environments specific to a Notebook.

Notebook scoped libraries are not cluster scoped, so they do not apply to other notebooks sharing the same cluster. Notebook scoped libraries only apply to current notebook and all workflow jobs that use the current notebook as 'Notebook tasks'.

Note

pip-based commands are currently only supported in notebooks with *.ipynb extensions. Python files with *.py extensions are not supported and installing JAR files as notebook scoped libraries is not supported.

Supported Commands

You can use the following commands for managing your notebook scoped libraries:

Install and uninstall a package by using the package name:

```
pip install <package>
pip uninstall <package>
```

Install and uninstall all the libraries from a requirements.txt file:

```
pip install -r /Workspace/path/to/requirements.txt
pip uninstall -r /Workspace/requirements.txt
```

Install and uninstall wheel files with *.whl extension from a workspace path:

```
pip install /Workspace/path/to/my-package.whl
pip uninstall /Workspace/path/to/my-package.whl
```

List the already installed notebook scoped libraries in the current notebook:

```
pip list
```

Limitations

You can run only one %pip command in a notebook cell. To install multiple packages, run each %pip command in a separate cell.

%pip uninstall <package> does not support the -y option or other additional options.

Installing packages from Git repositories is not supported. This includes Git-based package URLs such as git+https://....

Monitor Compute

This section explains the different methods and metrics you can use to monitor compute in your Oracle AI Data Platform.

Topics:

- [View Spark UI](#)
- [View Driver and Worker Logs](#)
- [View Metrics](#)
- [View Event Logs](#)
- [View Notebooks](#)

View Spark UI

You can view the Spark Web UI to see to monitor the status and resource consumption of your all-purpose compute clusters.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Spark UI** tab.
3. Optional: Click to pop-out button on the top right to view the Spark UI in a separate window.

View Driver and Worker Logs

You can view the Driver and Worker Logs of your All Purpose Compute Clusters for troubleshooting or debugging.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Logs** tab.
3. Filter your logs to see more specific information.

Cluster node	Worker #	Log level	Time frame
Driver ▼	1 ▼	Debug ▼	Last hour ▼

4. Click  **Download** to save a local copy of your filtered data.

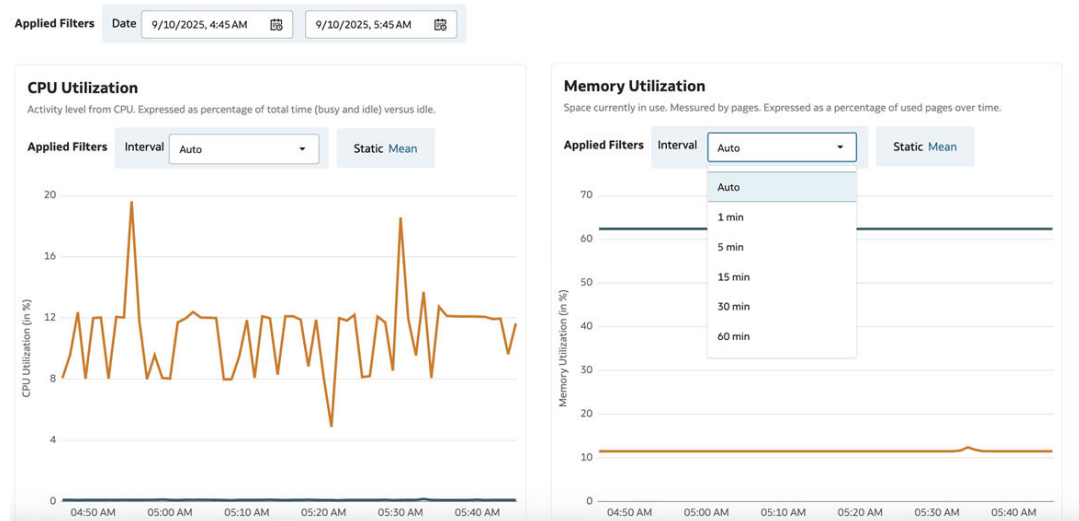
View Metrics

You can monitor the infrastructure metrics of your compute clusters for troubleshooting or for making any sizing adjustments.

You can view status and history for the following metrics:

- CPU Utilization
- Memory Utilization
- Disk read

- Disk write
 - File system utilization
 - Garbage Collector CPU utilization
 - Network received
 - Network transmitted
 - Active tasks
 - Total failed tasks
 - Total task tasks
 - Total completed tasks
 - Total number of tasks
 - Total shuffle read bytes
 - Total shuffle write bytes
 - Total task duration in seconds
 - SQL: Peak concurrent queries
 - SQL: Peak concurrent connections
1. Navigate to your workspace and click **Compute**.
 2. Click your cluster, then click the **Metrics** tab.



3. Select time frames using the **Date** filter to view metrics over a specific period.
4. Select an option from the **Interval** dropdown to filter information for a specific metric.

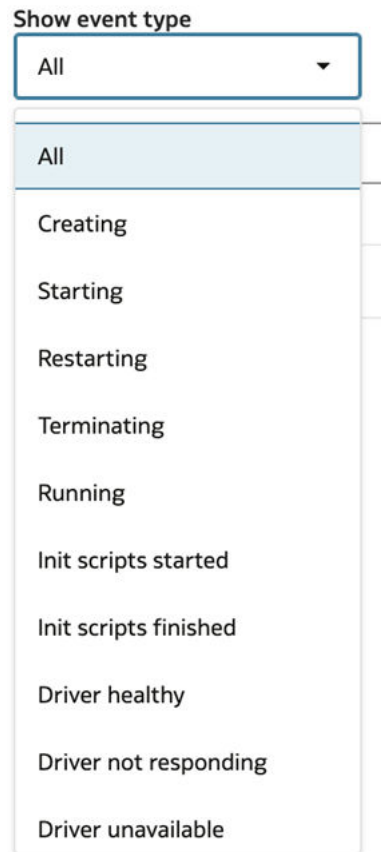
View Event Logs

You can view the Event Logs to monitor different cluster related operations, like creation of clusters, restarts of clusters, init script execution, or monthly maintenance updates.

Oracle AI Data Platform retains the last 14 days of event logs.

1. Navigate to your workspace and click **Compute**.

2. Click your cluster, then click the **Event Logs** tab.
3. Filter your logs to see more specific information.



View Notebooks

You can view all the notebooks the current cluster is attached with. This view includes notebook count, notebook status, and provides you a quick way to navigate to the appropriate notebooks.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Notebooks** tab.



The notebook state is **Active** if code is running from that notebook. The notebook state is **Idle** if no code is running from that notebook.

3. Click the name of a notebook to go to it.

AI Compute

This chapter covers the use of specialized computing resources for powering AI agents in Oracle AI Data Platform.

An AI compute cluster lets you host agents in AI Data Platform workspaces. AI computes can be attached to agents to run the playground experience and host deployed agents for production workloads.

Where to Find AI Compute

You manage AI compute from the Compute page, accessible from the left navigation pane in your AI Data Platform. Click the AI compute tab to see the AI compute clusters available in your workspace.



The AI compute tab is the main landing page for AI compute resources in the currently selected workspace. It provides table actions, filtering, sorting, and status information.

Column or Control	Purpose
Filter	Search the table for a resource by name or visible text.
Add (+)	Start creation of a new AI compute.
Delete	Remove selected AI computes.
AI compute name	Name of the compute resource. Click on the AI compute name to open the resource details page.
State	Current lifecycle state, such as Creating, Active, or Updating.
# of replicas	Number of compute replicas associated with the AI compute.
# of agents	Number of agents hosted in the AI compute. The count is shown as a link when agents are present.
Configuration	Summary of the compute shape of the replicas in the AI Compute, such as 1 OCPU, 16 GB or 2 OCPU, 32 GB.
Updated by / Updated on / Created on	Audit information for the most recent update and creation time.
Action menu (...)	Open resource-specific actions for the selected row.

Note

A green check indicates an Active resource. A spinner indicates an in-progress operation such as Creating or Updating. Wait for a resource to be Active before relying on it for production work.

Clicking the name of an AI compute in the AI compute tab allows you to view detailed information on that AI compute. The AI compute view has four tabs: Agents, Details, Compute utilization, and Permissions.

Agents Tab

The Agents tab lists agents that are hosted in the AI Compute resource. If an agent is not deployed, the AI compute is hosting the playground experience necessary for iterative development and testing. If the agent is deployed, the AI computed is hosting the agent deployment endpoint for production workloads.



Agent name	Authoring mode	Deployment	URI	URI State	AI Compute	Created on	Updated
memory_0000	Visual	Not Deployed				Thu, Jun 05, 2026 at 07:08:04 PST	Thu, Jun 05, 2026 at 07:08:04 PST
test_memory_0000	Visual	Not Deployed				Thu, Jun 04, 2026 at 14:08:04 PST	Thu, Jun 04, 2026 at 14:08:04 PST
test_memory_0001	Visual	Not Deployed				Thu, Jun 04, 2026 at 13:47:04 PST	Thu, Jun 04, 2026 at 13:47:04 PST
test_memory_0002	Visual	Not Deployed				Thu, Jun 04, 2026 at 13:25:04 PST	Thu, Jun 04, 2026 at 13:25:04 PST
test_memory_0003	Visual	Not Deployed				Thu, Jun 04, 2026 at 13:06:04 PST	Thu, Jun 04, 2026 at 13:06:04 PST
test_memory_0004	Visual	Not Deployed				Thu, Jun 04, 2026 at 12:52:04 PST	Thu, Jun 04, 2026 at 12:52:04 PST

Agents Tab Item	Description
Filter by agent name	Search the associated agent list.
Agent name	Name of the agent associated with the compute resource.
Authoring mode	How the agent was authored. Values of Visual or Code.
Deployment	Deployment state. Agents hosted in the AI compute can either be not deployed (playground is hosted in AI compute) or deployed (when the compute is hosting the deployment endpoint of the agent)
URI / URI State	Endpoint URI information when an agent is deployed and has a URI.
AI Compute	The compute associated with the production endpoint of the agent.
Created on / Updated on / Updated by	Creation and update metadata for the agent.

Details Tab

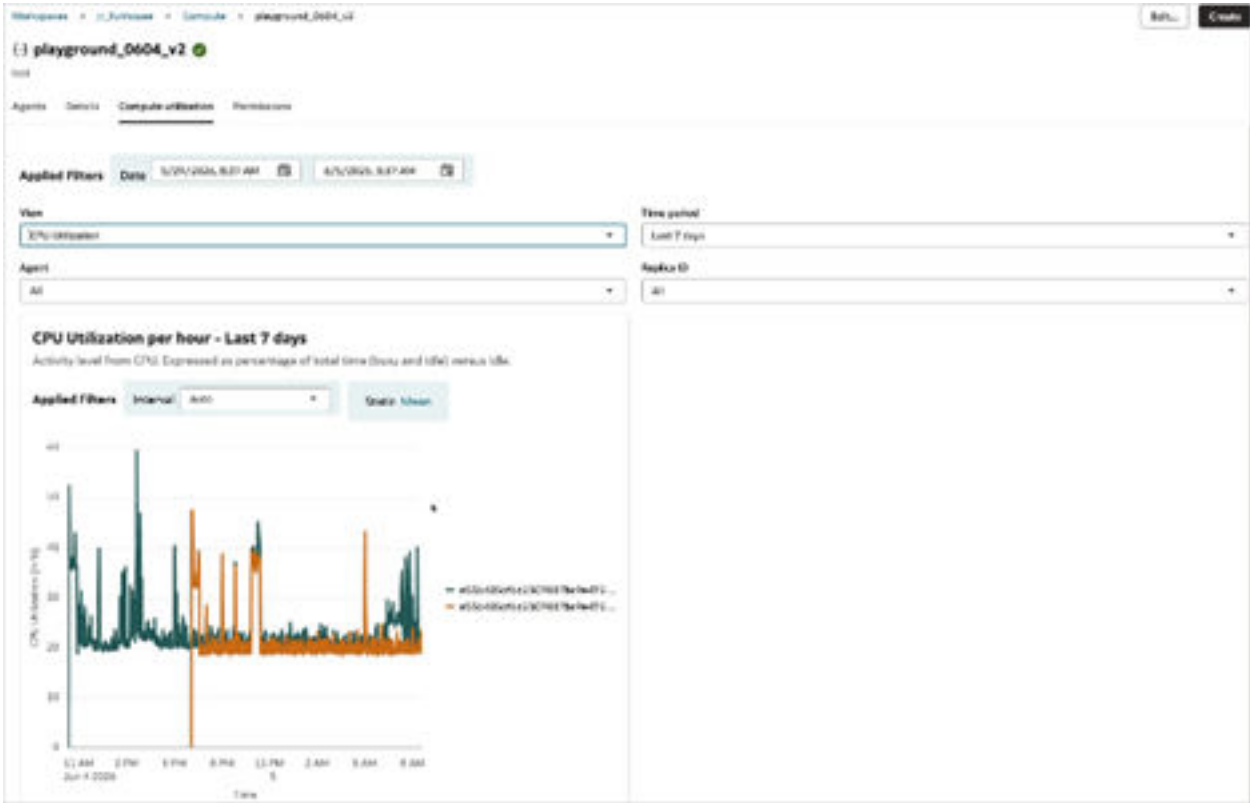
The Details tab summarizes the configured shape and storage for each replica of the AI compute.



Details Tab Item	Description
Driver shape	The GPU shape selected for the AI Compute. AMD is the driver shape for all AI Compute.
OCPUs	The selected number of OCPUs for the AI Compute.
Memory (GB)	The AI compute memory in GB.
Block volume	The block volume size of the AI compute in GB.

Compute Utilization Tab

Use the Compute utilization tab to review CPU, memory, and network utilization over time. This is useful when deciding whether to scale replicas, OCPUs, or memory up or down. There is no downtime during a scale out/in AI compute operation.



Note

Each time series corresponds to a different replica. For example, if you create an AI compute with two replicas, two time series will be displayed just like the screenshot above shows.

Compute Utilization Tab Item	Description
Applied Filters	Choose the start and end date/time for the utilization view.
View	Choose the metric. Options are: - CPU Utilization - Memory Utilization - Network Received Bytes - Network Transmitted Bytes
Time period	Choose between a custom time period or the last 7 days.
Replica ID	Filter metrics to a specific replica ID, or show All compute replicas.
Interval	Choose the aggregation interval for the chart. Options are Auto, 1min, 5 min, 15min, 30min, and 60min.

Permissions Tab

The permission tab provides a table of all the principals that have permission to this AI compute. Use this tab to add principals, change the permission level, or otherwise modify permissions.

Principal name	Principal type	Permission	Will be inherited
jgauth	USER	ADMIN	No
AI_DATA_PLATFORM_ADMIN	ROLE	ADMIN	No

Permissions Tab Item	Description
Principal name	Name of the principal or role.
Principal type	Principal type. USER or ROLE.
Permission	Permission level.
Will be inherited	If permissions granted for a parent object grant permissions to contained objects.

Create an AI Cluster

You can create AI compute clusters to run AI agents in your Oracle AI Data Platform.

1. On the Home page, navigate to your workspace.
2. Click **Create** then click **AI Compute**.
3. Provide a name and description for your AI compute cluster.
4. Set the number of compute replicas. Replicas scale the compute pool. Each replica hosts a copy of all the agents hosted in the AI compute.

Note

The maximum number of replicas is 10. Contact your Oracle representative if you need to increase your replica limit.

5. Set the number of OCPUs in each replica. You can choose 1, 2, 4, 8, 16, 32, or 64 OCPUs.
6. Set the memory (in GB) per replica. The allowed range of memory varies based on the number of OCPUs:

OCPUs	Memory
1 OCPU	16 or 32 GB
2 OCPUs	16, 32, or 64 GB
4 OCPUs	32, 64, or 128 GB
8 OCPUs	32, 64, 128, or 256 GB
16 OCPUs	64, 128, 256, or 512 GB
32 OCPUs	128, 256, or 512 GB
64 OCPUs	256, 512, or 1024GB

7. Review the corresponding AIDP Units that such an AI compute configuration costs per hour.
8. Click **Create**. The new resource appears in the list in the Creating state. The state changes to Active after provisioning completes.

Create an AI Cluster for an Agent

You can create a new AI cluster directly from the agent interface and attach it immediately.

For more information, see [AI Compute](#).

1. On the Home page, navigate to your agent.
2. Click **Compute** then click **Create a new AI compute**.
3. Provide a name and description for your AI compute cluster.
4. Select the number of OCPUs and memory size for your AI compute cluster.
5. Click **Create**.

Attach an Existing AI Cluster to an Agent

You can attach an AI cluster you've previously created to an agent on which you have at least USE permissions.

1. On the Home page, navigate to your agent.
2. Click **Compute** then click **Attach to AI Compute**.
3. Click on the cluster you want to use from the list.

Your agent shows **Alcompute: (ClusterName) running** when it has been successfully attached. This can take up to several minutes.

Edit an AI Cluster

You can modify the configuration settings of an AI compute through the Edit operation.

Changing the OCPU or memory for you AI compute cluster can take several minutes to complete. You may also interrupt long running queries in your agent.

1. On the Home page, navigate to your workspace.
2. Click on Compute then click on the AI Compute tab.
3. Next to the AI compute cluster you want to modify, click ... **Actions** then click **Edit**. You can also click the name of the AI compute then click **Edit** in the top-right.

The screenshot displays the Oracle AI Data Platform interface. At the top, there are tabs for 'Clusters' and 'AI compute'. Below the 'AI compute' tab, there is a search bar and a table listing AI compute clusters. The table has columns for 'AI compute name', 'State', '# of replicas', '# of agents', 'Configuration', 'Updated by', 'Updated on', and 'Created by'. The clusters listed are 'louis_compute', 'test2', 'test_no_memory_agent', 'playground_0604_v1', 'playground_0604_v2', and 'playground_0604'. The 'test2' cluster is highlighted. To the right of the 'test2' row, there is a three-dot menu icon. A dropdown menu is open, showing options: 'Edit...', 'Start', 'Restart', 'Stop', 'Copy ID', and 'Delete'. The 'Edit...' option is highlighted with a red box. Below the table, there is a breadcrumb trail: 'Workspaces > j_runhouse > Compute > test2'. Below the breadcrumb, there is a section for 'test2' with a green status icon. Below this, there is a section for 'Agents' with tabs for 'Agents', 'Details', 'Compute utilization', and 'Permissions'. Below the 'Agents' tab, there is a search bar for 'Filter by agent name'.

4. Modify the name, description, number of compute replicas, OCPU count per replica, and memory per replica. Changes to AI compute are completed with no downtime.
5. Click **Update**. Your AI compute state displays as Updating while the update is in progress.

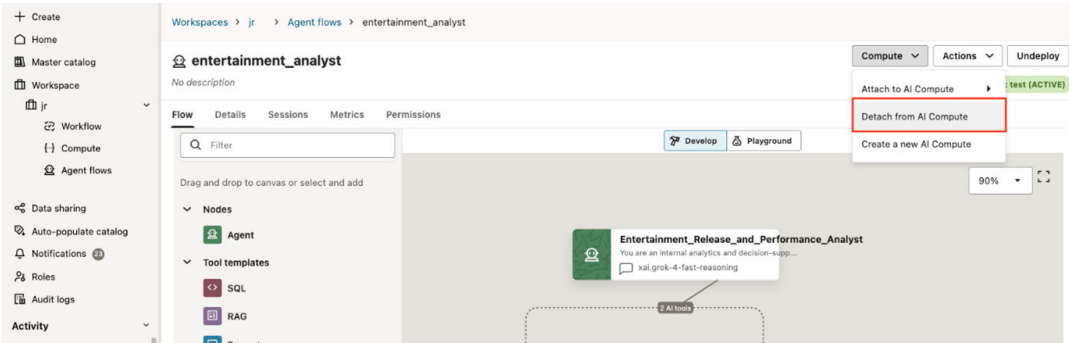
Detach an Agent from AI Compute

You can detach an agent from an AI compute. Detaching the AI compute removes the agent code running on the attached compute and prevents testing.

Note

Detaching an agent from AI compute does not delete the agent. You can resume building and testing the agent by attaching it to the same or different AI compute.

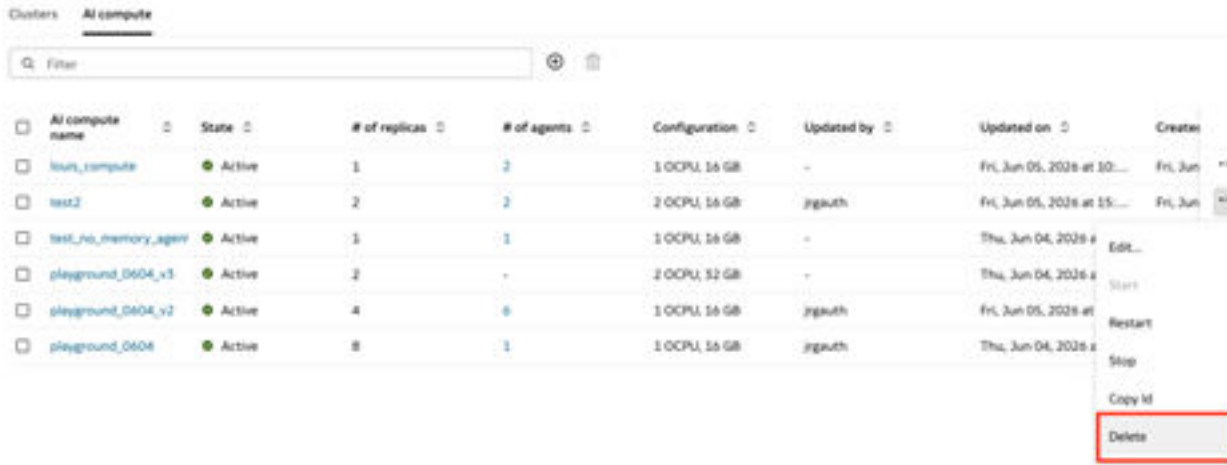
- 1. On the Home page, navigate to your agent.
- 2. Click **Compute** then click **Detach from AI Compute**.



Delete an AI Cluster

You can delete AI compute clusters that are unused or no longer needed.

- 1. Navigate to your workspace and click **Compute** then click the AI Compute tab.
- 2. Next to the AI cluster you want to delete, click ... **Actions** and click **Delete**.
- 3. Click **Delete**.

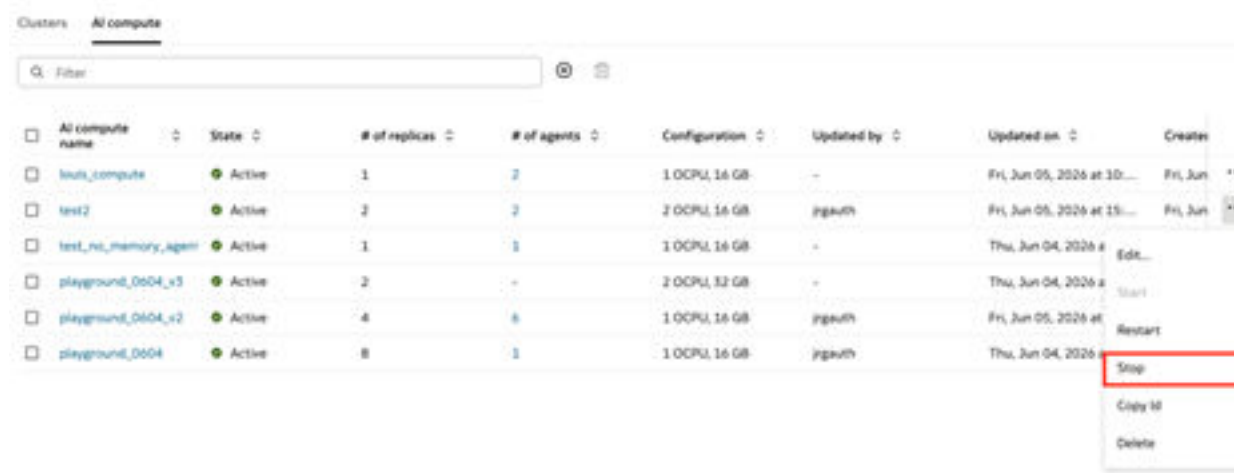


Stop an AI Cluster

You can stop an AI compute to stop all agents running on the AI compute. The compute is freed and the metering stops.

We recommended stopping an AI compute whenever the agents hosted on the compute are not in use.

- 1. On the Home page, navigate to your workspace.
- 2. Click on Compute then click on the AI Compute tab.
- 3. Next to the AI cluster you want to stop, click ... **Actions** and click **Stop**.



Start an AI Cluster

You can start an AI compute that has been previously stopped.

- 1. On the Home page, navigate to your workspace.
- 2. Click on Compute then click on the AI Compute tab.
- 3. Next to the AI cluster you want to start, click ... **Actions** and click **Start**.

Clusters AI compute

☐	AI compute name	State	# of replicas	# of agents	Configuration	Updated by	Updated on	Creator
☐	louis_compute	Active	1	2	1 OCPU, 16 GB	-	Fri, Jun 05, 2026 at 10:...	Fri, Jun ...
☐	test2	Active	2	2	2 OCPU, 16 GB	jgauth	Fri, Jun 05, 2026 at 15:...	Fri, Jun ...
☐	test_no_memory_agent	Active	1	1	1 OCPU, 16 GB	-	Thu, Jun 04, 2026 at 1:...	Thu, Ju ...
☐	playground_0604_v3	Stopped	2	-	2 OCPU, 32 GB	-	Fri, Jun 05, 2026 at 15:...	Thu, Ju ...
☐	playground_0604_v2	Active	4	6	1 OCPU, 16 GB	jgauth	Fri, Jun 05, 2026 at	Edit...
☐	playground_0604	Restarting	8	1	1 OCPU, 16 GB	jgauth	Fri, Jun 05, 2026 at	Start

Restart

Stop

Copy ID

Delete

Restart an AI Cluster

You can restart an active AI compute to pick up recent changes or updates.

📘

Note

Agents hosted on active compute will be interrupted when restarting.

- 1. On the Home page, navigate to your workspace.
- 2. Click on Compute then click on the AI Compute tab.
- 3. Next to the AI cluster you want to start, click ... **Actions** and click **Restart**.

Clusters AI compute

☐	AI compute name	State	# of replicas	# of agents	Configuration	Updated by	Updated on	Creator
☐	louis_compute	Active	1	2	1 OCPU, 16 GB	-	Fri, Jun 05, 2026 at 10:...	Fri, Jun ...
☐	test2	Active	2	2	2 OCPU, 16 GB	jgauth	Fri, Jun 05, 2026 at 15:...	Fri, Jun ...
☐	test_no_memory_agent	Active	1	1	1 OCPU, 16 GB	-	Thu, Jun 04, 2026 at	Edit...
☐	playground_0604_v3	Active	2	-	2 OCPU, 32 GB	-	Thu, Jun 04, 2026 at	Start
☐	playground_0604_v2	Active	4	6	1 OCPU, 16 GB	jgauth	Fri, Jun 05, 2026 at	Restart
☐	playground_0604	Active	8	1	1 OCPU, 16 GB	jgauth	Thu, Jun 04, 2026 at	Stop

Restart

Stop

Copy ID

Delete

10

Connect to Compute

This section covers connecting compute in your Oracle AI Data Platform to other business intelligence tools.

Topics:

- [Connections](#)
- [Connect Oracle Analytics to Oracle AI Data Platform](#)
- [Download JDBC Driver](#)
- [Connect Tableau to AI Data Platform using JDBC](#)
- [Download ODBC Driver](#)

Connections

You can connect your Oracle AI Data Platform with Oracle Analytics Cloud or other business intelligence tools.

You can connect to AI Data Platform from different business intelligence tools using a custom JDBC or ODBC provided by AI Data Platform. AI Data Platform also supports connections from third party commercial BI tools (Tableau, Power BI) or open source BI tools like DBeaver. In order to connect from these tools, you need the connection details of the compute cluster in your AI Data Platform, which you can find in the Connection details tab. The required information varies depending on the product you are connecting from.

[Workspaces](#) > [Data-Playground](#) > [Compute](#) > Dev_cluster_1

Actions ▾

Edit...

Dev_cluster_1 ✓
No description

Details

Connection details

Library

Spark UI

Permissions

Connect with BI Tool


Oracle Analytics Cloud



Tableau



Power BI


JDBC driver

Hostname	gateway.preprod.datalake.us-phoenix-1.oci.oraclecloud.com	
JDBC URL	jdbc:spark://gateway.preprod.datalake.us-phoenix-1.oci.oraclecloud.com/default;SparkServerType=DFI;httpPath=cliservice/44d5e6ee-700e-4ddb-934e-4915a9473b51	
Download JDBC Driver		

ODBC driver

Hostname	gateway.preprod.datalake.us-phoenix-1.oci.oraclecloud.com	
Port	443	
Default database	default	
Spark server type	DFI	
Download ODBC Driver ▾		

Connect Oracle Analytics to Oracle AI Data Platform

You can connect to the catalog or tables managed by your Oracle AI Data Platform instance from an Oracle Analytics Cloud instance.



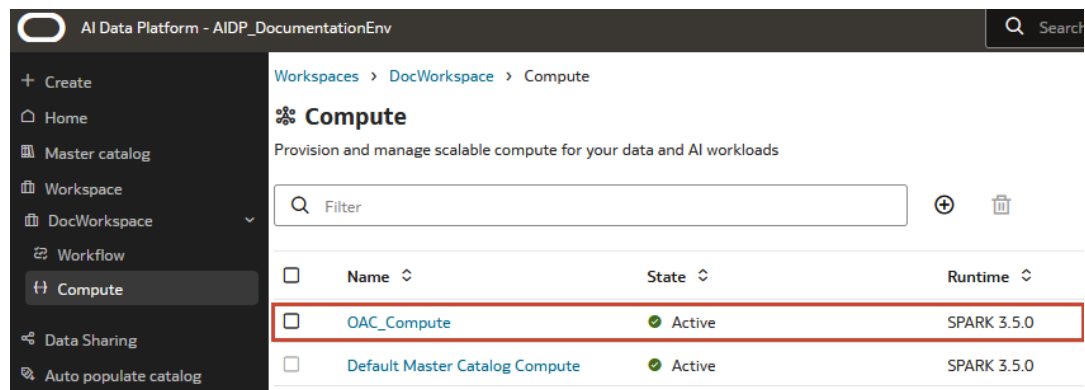
In order to create a connection from Oracle Analytics to AI Data Platform, you need a connection file from AI Data Platform and an API key from Analytics Cloud.

Get an Oracle Analytics Connection Configuration File

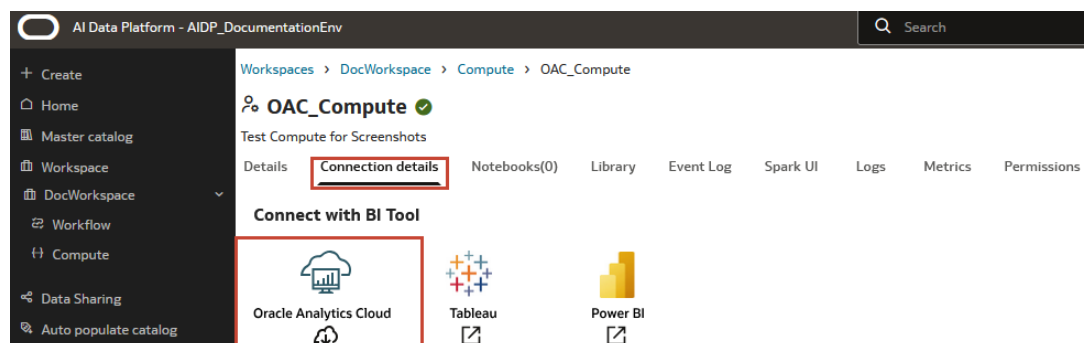
You need to download a config.json file from your Oracle AI Data Platform instance to connect it to Oracle Analytics.

You need a compute cluster with at least 2 OCPUs, 32 GB of memory for both driver and worker nodes, and at least 2 workers.

1. Navigate to your workspace and click **Compute**.



2. Click the cluster you want to connect to Oracle Analytics.
3. In the **Connection details** tab, click the **Download** icon underneath the Oracle Analytics Cloud logo. A config.json file is downloaded to your machine.



Connect to Oracle Analytics

You can connect to the catalog or tables managed by your Oracle AI Data Platform instance from an Oracle Analytics instance.

Ensure you have a connection configuration file from AI Data Platform before creating your connection in Oracle Analytics. See [Get an Oracle Analytics Connection Configuration File](#).

Oracle Analytics connections to AI Data Platform instances only support one catalog per connection. You need to create additional connections to use more than one catalog.

1. On the Oracle Analytics home page, click **Create** then click **Connection**.
2. Click **Oracle AI Data Platform**.
3. Enter a name and description for your connection.
4. In **Location**, specify where to save the connection. To share the connection with other users, save it in /Shared Folders.
5. In Oracle Analytics Cloud, for Connection Details, click **Select**, navigate to your config.json file, and click **Open**.
6. Select **API Key** as the authentication type.
7. Under **Public API Key**, click **Generate**, then **Copy**.
8. In OCI, add an API Key.
 - a. In OCI, click the user icon, then click **User Settings**.
 - b. In My profile, click **Tokens and keys**, then click **Add API key**.
 - c. Select **Paste a public key**, and copy in the key that you generated in Step 7, then click **Add**.

Wait two to three minutes for the API key to synchronize.

9. In **Catalog**, specify the Catalog name to be used in your datasets.

If the **Catalog** list doesn't load and you see the alert message "User does not have sufficient privileges to perform the operation", wait a little longer for the API key to synchronize.

10. Click **Save**.
11. Test your new connection by creating a dataset.

If you have followed the instructions and are still **having issues saving your connection or listing your schemas**, a cluster restart may be required.

Download JDBC Driver

You can download the JDBC driver from the Connection details tab and configure it for different BI tools.

1. Navigate to your workspace and click **Compute**.
2. Click the cluster you want to connect to a JDBC compatible BI tool and click **Connections**.
3. Click **Download JDBC Driver**.

Configure DBeaver

You need to install DBeaver and prepare files downloaded from Oracle AI Data Platform before you can connect it to a compute cluster in your AI Data Platform.

1. [Download](#) and install the DBeaver client. You can use either Community or Enterprise versions, but they must be at least 22.x. DBeaver is only available for Windows, Mac OS X, Eclipse Plugin, and Linux.
2. Unzip the file from [Download JDBC Driver](#).
3. From the unzipped files, unzip the `simbaSpark.zip` driver file.

Configure DBeaver with the Spark Simba JDBC Driver

You can configure DBeaver to connect to a compute cluster in your Oracle AI Data Platform.

You must have installed DBeaver and extracted the `simbaSpark.zip` files downloaded from AI Data Platform.

1. Open DBeaver.
2. Click **Database Navigator**.
3. Click **Driver Manager**.
4. Click **New**.
5. Enter AI Data Platform as the **Driver Name**.
6. Click **Libraries**.
7. Click **Add Folder**.
8. Browse to the location of the `sparkSimba.zip` extract.
9. Click **Find Class**.
10. Select **OK**.
11. Set **Driver class** to `com.simba.spark.jdbc.Driver`.
12. Click **Settings**.
13. Set **Class Name** to `com.simba.spark.jdbc.Driver`.
14. Click **OK**.

Create a Database Connection in DBeaver

To connect DBeaver to a compute cluster in your Oracle AI Data Platform, you need to first create a database connection in DBeaver.

You must have configured DBeaver with the Spark Simba JDBC driver downloaded from AI Data Platform. For more information, see [Configure DBeaver with the Spark Simba JDBC Driver](#).

1. Open DBeaver.
2. Click **Database**.
3. Click **New Database Connection**.
4. Click **All**.

5. Select **AI Data Platform**.
6. Click **Next**.
7. Enter the URL of the JDBC driver. You can find the JDBC URL on the **Connection details** tab of your compute cluster in AI Data Platform.
8. Click **Finish**.

Connect DBeaver to Oracle AI Data Platform using JDBC

Once DBeaver is configured and has a database connection to your Oracle AI Data Platform, you can complete the connection between DBeaver and AI Data Platform.

1. Open DBeaver.
2. Click **Connect**.
3. Choose to connect with an authorization token or an API key.
 - **Connect using authorization token**
 - Use a token by not specifying any profile in the URL if you don't have a DEFAULT profile. For example: `jdbc:spark://gateway.aidp.me-riyadh-1.oci.oraclecloud.com/default;SparkServerType=AIDP;httpPath=cliservice/cf18b4ef-b83e-41dd-82b6-8d391584f6c5`
The URL opens a browser window.
Sign in to the tenancy where the AI Data Platform instance is created.
For more information, see [Token-based Authentication for the CLI](#).
 - **Connect using an API key**, by specifying the OCI Profile with `ociProfile=<profile_name>` in the connection URL.
 - Use API key authentication to connect to an AI Data Platform instance.
Use API key by specifying the OCI Profile with `ociProfile=<profile_name>` in the connection URL. For example, to use OCI Profile name Demo: `jdbc:spark://gateway.aidp.me-riyadh-1.oci.oraclecloud.com/default;SparkServerType=AIDP;httpPath=cliservice/cf18b4ef-b83e-41dd-82b6-8d391584f6c5 ;ociProfile=Demo`
For more information, see [Required Keys and OCIDs](#).
4. DBeaver creates a connection for reading metadata, and a connection for all the other operations. If you're limited for connections, you can disable the second one so that DBeaver uses one connection for all operations.
 - a. Click **Preferences**.
 - b. Click **Common**.
 - c. Click **Metadata**.
 - d. Deselect **Open separate connection for metadata reads**.

Download ODBC Driver

You can download the ODBC driver from the Connection details tab and configure it for different BI tools.

1. Navigate to your workspace and click **Compute**.

2. Click the cluster you want to connect to an ODBC compatible BI tool and click **Connections**.
3. Click **Download ODBC Driver**.
4. Select the appropriate OS from the list.

ODBC driver

Hostname

Port

Mac

Windows 32 bit

Windows 64 bit

Linux i686

Linux x86_64

Download ODBC Driver ▼

11

Streaming

The section covers the use of streaming data or continuously produced data in Oracle AI Data Platform.

Topics:

- [About Streaming](#)
- [Configuring Spark Structured Streaming using Workflows](#)

About Streaming

You can process streaming data or continuously produced data in near real-time in Oracle AI Data Platform using the Apache Spark Structured Streaming capability.

Both notebooks and workflows support Apache Spark structured streaming. You can use the following sources and sinks for reading stream data from, writing stream data to, and for checkpoint locations.

Table 11-1 Supported Sources and Sinks

Source or Sink	Supported?
Volume path (/Volume/bronze/bucket1)	Supported for all formats
Workspace path (/Workspace/folder1/)	Supported for all formats
Tables in catalogs with three part names (catalog.schema.table)	Supported for Delta format only Not supported for Parquet, CSV, JSON, ORC formats Example 1: Supported code • <code>streaming_df = spark.readStream.format("delta").table('stdcatalog.stdschema.deltatable')</code> • <code>streaming_df.writeStream.format("delta").outputMode("append").option("checkpointLocation", "/Volumes/checkpoints1/").toTable("stdcatalog.stdschema.deltatable")</code> Example 2: Unsupported code • <code>spark.readStream.option("withEventTimeOrder", "true").format("format").table("stdcatalog.stdschema.samplecsv")</code>
Kafka	Supported for any Kafka compatible streams without three-part-naming convention Not supported for Kafka based catalog following three-part-naming convention)
OCI Streaming service	Supported
OCI Object storage path (using oci://)	Unsupported

Table 11-1 (Cont.) Supported Sources and Sinks

Source or Sink	Supported?
Oracle Autonomous AI Lakehouse, Oracle AI Database, Oracle Autonomous AI Transaction Processing	Unsupported for streaming (readStream or writeStream)

Structured Streaming Using Notebooks

You can write Python code to process stream data in a notebook. Either volume paths or workspace paths are valid as a checkpoint location, but object Storage paths (oci:// format) are not supported as a checkpoint location. We recommend using volume paths as a checkpoint location.

```

# Trigger Once example
# Source => Volume File Source, Sink => Volume File Sink, checkpoint location => volume location
# Place the csv file in your volume location ex: </Volumes/default/default/manVol/data/>

# Define input and checkpoint paths
input_path = "/Volumes/default/default/manVol/data/*.csv"
output_path = "/Volumes/default/default/manVol/TriggerOnceOutput"
checkpoint_path = "/Volumes/default/default/manVol/TriggerOnceOutputCheckpoint"

userSchema = spark.read.option("header", "true").csv(input_path).schema

# Read streaming data from a file source
df = spark.readStream \
    .format("csv") \
    .schema(userSchema) \
    .option("header", "true") \
    .load(input_path)

word_counts = df.groupBy("vendor_id").count()

display(word_counts)

```

```

# Trigger Once example
# Source => Volume File Source, Sink => Volume File Sink, checkpoint location => volume location
# Place the csv file in your volume location ex: </Volumes/default/default/manVol/data/>

# Define input and checkpoint paths
input_path = "/Volumes/default/default/manVol/data/*.csv"
output_path = "/Volumes/default/default/manVol/TriggerOnceOutput"
checkpoint_path = "/Volumes/default/default/manVol/TriggerOnceOutputCheckpoint"

userSchema = spark.read.option("header", "true").csv(input_path).schema

# Read streaming data from a file source
df = spark.readStream \
    .format("csv") \
    .schema(userSchema) \
    .option("header", "true") \
    .load(input_path)

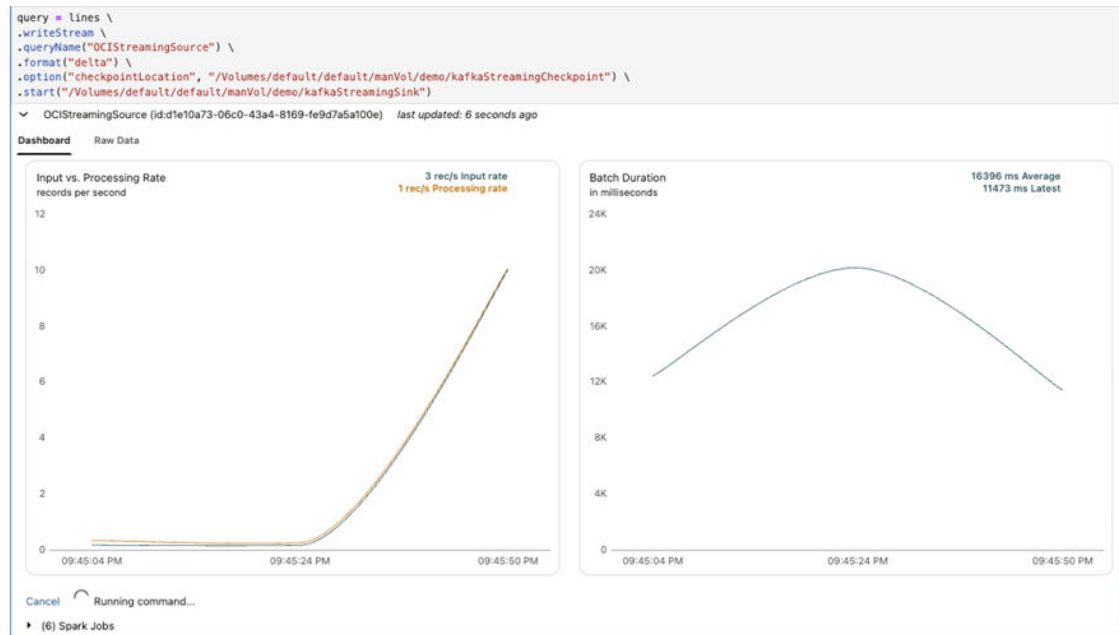
word_counts = df.groupBy("vendor_id").count()

#display(word_counts)

# Write the streaming data to a console sink with Trigger.Once()
query = word_counts.writeStream \
    .outputMode("complete") \
    .format("delta") \
    .trigger(once=True) \
    .option("checkpointLocation", checkpoint_path) \
    .start(output_path)

```

You can see Apache Spark streaming-related events, like input rate, processing rate, and batch duration from the Dashboard tab in your notebook while running streaming code.



You can also view the raw streaming-related events from the Raw Data tab while you incrementally develop your code.

```

streaming_df = spark.readStream.format("rate").load()
display(streaming_df)

```

Cancel Running command...

display_1757433453123631802 (id:eff49446-990b-419f-9399-cf748148c9a4) last updated: just now

Dashboard Raw Data

```

{
  "batch_id": 28,
  "duration_ms": {
    "add_batch": 38,
    "commit_offsets": 24,
    "get_batch": 0,
    "latest_offset": 0,
    "query_planning": 4,
    "trigger_execution": 84,
    "wal_commit": 18
  },
  "id": "eff49446-990b-419f-9399-cf748148c9a4",
  "input_rows_per_second": 100,
  "name": "display_1757433453123631802",
  "num_input_rows": 1,
  "processed_rows_per_second": 11.904761904761903,
  "run_id": "d66df18c-0283-4edc-980e-87b06422b4ac",
  "sink": {
    "description": "MemorySink",
    "num_output_rows": 1
  },
  "sources": [
    {
      "description": "RateStreamV2[rowsPerSecond=1, rampUpTimeSeconds=0, numPartitions=default",
      "end_offset": 29,
      "input_rows_per_second": 100,
      "latest_offset": 29,
      "num_input_rows": 1,
      "processed_rows_per_second": 11.904761904761903,
      "start_offset": 28
    }
  ],
  "state_operators": [],
  "timestamp": "2025-09-09T15:58:02.487000+00:00"
}

```

Configuring Spark Structured Streaming using Workflows

You can configure a streaming task inside a workflow for continuous processing of stream data.

You first need to create a job and then add one Notebook or Python task to that job to begin using workflows with streaming in Oracle AI Data Platform.

1. Navigate to your workspace and click **Workflow**.
2. Click **Create Job**.
3. Provide a name and description for your job.
4. Click **Browse** and select the location to save the job in your AI Data Platform. Click **Select**.
5. Enter 1 for **Max Concurrent Runs**.
6. Click **Create**.
7. Click the job you just created.
8. Click **Add task**.
9. Provide a name for your task.
10. Select **Notebook** or **Python** for **Task type**.
11. Click **Browse** and navigate to the Notebook or Python script you want to add as a Streaming task. Click **Select**.
12. Select a compute cluster for the Notebook or Python task, if one is not already attached.
13. Select the **Streaming** checkbox. Selecting Streaming disables execution timeout and task dependencies as options.

Task details saved at Fri, Aug 29, 2025 at 12:03:46 UTC

Task name
Task_350339

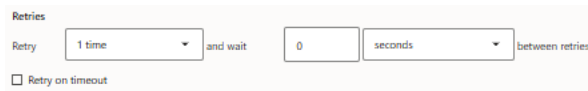
Task type
Notebook task

File location
/Workspace/code/wf_notebook.ipynb **Browse**

Cluster
INTL2

☒ **Streaming**

14. Select the number of retries a task should attempt on failure. If you select more than 0, you must also specify how much time the job run should wait between retries and if retries should be attempted on timeout.



Retries

Retry 1 time and wait 0 seconds between retries

☐ Retry on timeout

15. Click **Run Now**.

After a Streaming task is started, it continues to run until you manually stop it. During regular monthly maintenance, the Streaming task is stopped and restarted by the service without requiring any action from your end.

Part III

Integration

You can configure Oracle AI Data Platform to integrate seamlessly with other Oracle services. This section covers enabling services integrated with AI Data Platform and leveraging the shared data to build advanced AI/ML pipelines.

Topics:

- [Fusion Data in Oracle AI Data Platform with BICC](#)
- [Oracle Cloud Infrastructure GoldenGate Integration](#)
- [Git Integration \(Preview\)](#)

Oracle Cloud Infrastructure GoldenGate Integration

Oracle AI Data Platform integrates with Oracle Cloud Infrastructure GoldenGate to enable real-time replication of operational data into your AI Data Platform for analytics and AI-driven workloads.

With the integration of OCI GoldenGate, you can continuously stream and merge transactional changes from source systems into AI Data Platform target tables.

For full details on configuration parameters, prerequisites, and troubleshooting, please refer to the OCI GoldenGate documentation [Connect to Oracle AI Data Platform](#).

How it Works

- OCI GoldenGate captures change data from source databases and uses stage and merge data flow to enable high throughput, low impact data replication into AI Data Platform.
- OCI GoldenGate can run initial instantiation and sync the initial load with cdc replication without any data loss. During initial instantiation or cdc replication, OCI GoldenGate can automatically create target tables.

This allows AI Data Platform to always reflect the most up-to-date view of your operational data for downstream analytics, machine learning, and data engineering.

Requirements to Integrate with OCI GoldenGate

- An OCI tenancy with an active AI Data Platform instance
- Database privileges on target AI Data Platform tables
- The SIMBA JDBC driver for Apache Spark, downloadable from the AI Data Platform console.

Configuration Summary

When you set up OCI GoldenGate integration with AI Data Platform as part of the replicat process, you need to do the following:

- Select **AIDP** from the **Target** dropdown.
- Configure the replicat properties marked as *TODO* in the properties file.
- Set `gg.target=aidp` to enable AI Data Platform at the target.
- Specify a staging location. If no staging location is specified, Oracle Cloud Infrastructure Object Storage is used.

OCI Generative AI (Pretrained Foundation Models)

Oracle Cloud Infrastructure Generative AI is a fully managed Oracle Cloud Infrastructure service that provides a set of state-of-the-art, customisable large language models (LLMs) that cover a wide range of use cases, including chat, text generation, summarization, and creating text embeddings.



Oracle AI Data Platform users can access OCI Generative AI models if they have the requisite permissions and the pre-trained model is hosted in the same region as the data lake. For more information on permissions, see [Getting Access to Generative AI](#).

You can use OCI Generative AI models in AI Data Platform for the following use cases:

- Use the pre-trained chat models to create text for any purpose.
- Extract specific pieces of data from text.
- Generate executive summaries for documents that are too long to read, or summarize any type of text.
- Classify text into predefined categories.

You can also run batch inferences on Spark Data Frames using the pre-trained models in a language of your choice, like SQL or Python. For more information on pretrained models, see [Pretrained Foundational Models in Generative AI](#).

Note

Check to see if the OCI Generative AI service you want to use is available in your region. Models are region-specific and scripts that refer to a model in an unsupported region may not work. For more information, see [Generative AI Regions](#).

Prerequisites for OCI Generative AI

You must meet the following prerequisites to use OCI Generative AI in AI Data Platform:

- User must have USE permissions on the base models
- AI Data Platform is in the same region where the OCI Generative AI models are hosted

If the prerequisites are met, the models are listed in the `default.oci_ai_models` schema. You can then list the models in the catalog explorer while working in a notebook and drag drop the models to generate sample code or use the model for batch inference. Alternatively, you can choose to write your code in an AI Data Platform notebook to invoke the model.

You can use the following methods to invoke a OCI Generative AI model:

SQL

```
select *, query_model(model_name, concat("What is the sentiment for this  
review: ", review)) as sentiment from  
<<catalog_name>>.<<schema_name>>.<<table_name>>
```

Where:

- *model_name* is the generative AI model you want to invoke:
default.oci_ai_models.<model_name>
- *review* is the column name that is used to create the prompt
- *sentiment* is the output column name
- <<catalog_name>>.<<schema_name>>.<<table_name>> is the table in 3-part name pattern

PySpark

```
df.withColumn("sentiment", query_model(model_name, "What is the sentiment for  
this review: "+review))
```

Where:

- *model_name* is the generative AI model you want to invoke:
default.oci_ai_models.<model_name>
- *review* is the column name that is used to create the prompt
- *sentiment* is the output column name
- *df* is the input data frame

Request Limit

Description	Limit Name	Service Limit
Maximum number of chat requests per minute allowed per compartment for on-demand inferencing	max-on-demand-chat-request-per-minute-count	500

Part IV

Machine Learning

Oracle AI Data Platform provides machine learning (ML) lifecycle management using MLflow concepts and APIs, specifically experiments, runs, and a model registry.

These capabilities are deeply integrated into AI Data Platform across multiple surfaces, including Workspaces, Experiments, and Catalog, so teams can track work where it happens and promote outcomes into shared, governed assets.

ML Lifecycle

End-to-end ML lifecycles typically follows these steps:

1. **Data prep:** Clean and format raw inputs
2. **Exploratory data analysis (EDA):** Explore data to find patterns
3. **Feature engineering:** Create variables for models
4. **Experiment:** Iteratively train using multiple approaches (each iteration is a run)
5. **Validate and store:** Identify the best run and register the model for reuse
6. **Run inferences:** Use a registered model version for batch inference from notebooks
7. **Monitor:** Track basic production performance and availability for deployed models

Core capabilities

Experiment tracking per team workspace

- Experiments are scoped to a workspace to separate teams and organize work.
- MLflow-compatible autologging captures parameters, metrics, and artifacts for each run, creating a reproducible record that supports reruns with controlled changes.

Run comparison and registration

- Runs can be filtered and compared to identify a candidate model.
- A run can be registered into the Master Catalog-backed Model Registry, carrying versioning, tags, and custom fields. Version management is handled by the platform when updated models are registered.

From registry to notebook inference

- Models can be loaded in notebooks by latest or by explicit version, enabling consistent reuse.
- Batch inference workflows can reference registry versions directly, reducing manual handling between experimentation and inference.

Lineage for auditability

- Registered models link back to the originating experiment run, including run conditions such as hyperparameters, environment variables, metrics, and artifacts.
- This supports review and audit by making the provenance of each model explicit.

Why Use MLflow?

AI Data Platform uses MLflow as the foundation for its MLOps framework because it provides an open, extensible, and framework-agnostic approach to managing the end-to-end machine learning lifecycle.

MLflow supports the core capabilities required for operationalizing machine learning at scale, including experiment tracking, model packaging, artifact management, model versioning, registry-based and governance. Its ability to capture parameters, metrics, artifacts, and run metadata in a consistent way makes it well suited for improving reproducibility, auditability, and collaboration across data science and engineering teams.

A key reason for selecting MLflow is its broad compatibility with popular machine learning frameworks such as TensorFlow, PyTorch, and scikit-learn. This allows AI Data Platform to support diverse model development patterns without forcing teams into a single framework or toolchain. MLflow's plugin architecture and deployment flexibility also make it easier to extend the platform and integrate with existing enterprise infrastructure.

By standardizing on MLflow, AI Data Platform can provide a consistent MLOps experience across experimentation, model registration, lifecycle management, while retaining the flexibility needed to evolve with different AI/ML use cases.

Experiments

This chapter provides information on creating, managing, and developing experiments in your workspace.

Experiments in Oracle AI Data Platform provide means for data scientists, machine-learning engineers, and managers to collaborate in the development of models. AI Data Platform enables you to:

- Track experiments for building best models via performance analysis, collaboration, analysis of experimental conditions (hyper-parameters, input dataset, feature engineering etc.)
- Ensure old experiments aren't repeated with same experimental conditions
- Rerun old experiments with different experimental conditions for likely performance gains
- Reproducibility of earlier experiments
- Facilitate retraining with newer datasets when performance degrades or retraining is warranted
- Load a model from Model Registry into notebook and compare model performance versus a new model being developed

① Note

If you haven't previously used Experiments or Models in your AI Data Platform, you need to either restart your associated compute cluster or create a new one to use with Experiments and Models.

Limitations

Experiments are currently not supported by ARM-based compute clusters. Ensure the attached compute cluster is either Intel- or AMD-based.

Create an Experiment

You can create experiments directly in the Experiments page of Oracle AI Data Platform.

1. On the home page, click **Experiments**.
2. Click **Create**.
3. Provide a name and description for your experiment.
4. Optional: Provide additional metadata in the form of free-form or defined tags. Click **Add** to create additional tags.
5. Click **Create**.

Edit an Experiment

You modify details for experiments in your Oracle AI Data Platform workspace.

1. On the home page, click **Experiments**.
2. Next to the experiment you want to edit, click **Edit**.
3. Modify details for your experiment.
4. Click **Save changes**.

View Experiment Run Details

You can see a history of past experiment runs, compare them, and see the details for specific runs for experiments in your Oracle AI Data Platform workspace.

1. On the home page, click **Experiments**.
2. Click the name of the experiment you want to view run details for.
3. Use the drop-down lists and search bar to filter the displayed experiment runs.
4. Click **List** or **Compare** to change the experiment runs are displayed.
 - **List** displays experiment runs based on your filters in ordered rows with metrics about the experiment run shown in each column. You can sort by metric by clicking that column header.
 - **Compare** displays bar graph comparisons of key metrics for every experiment run currently displayed by your filters.
5. Click the name of an experiment run to view details for that experiment run.

Create an Experiment Run in a Notebook with Sample Code

You can create runs for an experiment within a notebook by modifying sample code with existing experiment details.

1. Navigate to your notebook where you want to create a run for an experiment.
2. Click the Experiments tab.
3. Click **Sample code**.
4. In the sample code block, replace experiment name="Customer Churn Prediction" with experiment name="<your_experiment_name>". You can also copy this code and modify it with your experiment name:

```
import mlflow
experiment_name = "experiment name" #Replace this with your own experiment
name
mlflow.set_experiment(experiment_name)

mlflow.autolog()

with mlflow.start_run():
    # training code goes here
```

```
# OPTIONAL - Log additional items after training
mlflow.log_metric("custom_metric", 99.9)
```

5. Autologs automatically record a default set of metrics, depending on the model selected. To manually specify your own metrics, you can modify this code to invoke `mlflow.log_metric("<metric_name>", <metric_variable>)`:

```
import mlflow
import numpy as np
from sklearn.tree import DecisionTreeRegressor
from sklearn.metrics import mean_squared_error, mean_absolute_error,
r2_score

experiment_name = "Customer Churn Prediction"
mlflow.set_experiment(experiment_name)

mlflow.autolog()

with mlflow.start_run():
    # training code goes here
    model = DecisionTreeRegressor(random_state=42, max_depth=5)
    model.fit(X_train, y_train)

    preds = model.predict(X_test)

    mse = mean_squared_error(y_test, preds)
    rmse = float(np.sqrt(mse))

    mlflow.log_metric("test_rmse", rmse)
    mlflow.log_metric("test_mae", float(mean_absolute_error(y_test,
preds)))
    mlflow.log_metric("test_r2", float(r2_score(y_test, preds)))
    # OPTIONAL - Log additional items after training
    mlflow.log_metric("custom_metric", 99.9)
```

6. Run the code block from your notebook. The run is now registered to specified experiment.

Note

Multiple runs for an experiment are automatically logged with different names. For parameter sweep scenarios, Oracle AI Data Platform automatically captures all the runs and specified metrics with different names to the specified experiment.

Register a Model from Experiment Run Details

You can register a model in your Oracle AI Data Platform workspace from a specific experiment run by clicking through to the details of that experiment run.

1. On the home page, click **Experiments**.
2. Click the name of the experiment you want to view run details for.
3. Click the name of the experiment run you want to register a model from.
4. Click **Register**.

5. Provide a name and description for your model.
6. Choose the location in your master catalog to create the model.
7. From the **Models** drop-down list, select the appropriate model.
8. Optional: Provide additional metadata in the form of free-form or defined tags. Click **Add** to create additional tags.
9. Optional: Provide additional information in the form of custom fields, like model type or use cases. Click **Add** to create additional custom fields.
10. Click **Register**.

Delete Experiment

You modify details for experiments in your Oracle AI Data Platform workspace.

Note

You can't delete an experiment if there is a registered model based off a run of that experiment.

1. On the home page, click **Experiments**.
2. Next to the experiment you want to delete, click **Delete**.
3. Click **Delete**.

Models

Oracle AI Data Platform enables you to develop and refine machine-learning (ML) models that you can use for running inferences and predictions for your data.

Machine-learning models in AI Data Platform enable you to build optimized models through performance analysis, collaboration, and analysis of experimental conditions like hyper-parameters, input datasets, and feature engineering.

The Model registry in AI Data Platform exists in the Master Catalog and experiment runs developed in your workspace can be selected and registered as models. Models in the registry can be loaded into notebooks to compare performance against models in development or scheduled through workflows to run batch inferences using the loaded model.

You can register models in your AI Data Platform workspace by creating one based on an existing experiment or by uploading a file that contains the necessary information. You can also register a model directly from an experiment in your notebook.

Note

If you haven't used previously used Experiments or Models in your AI Data Platform, you need to either restart your associated compute cluster or create a new one to use with Experiments and Models.

Once a model is created, you can deploy it for use as an endpoint from the Deployments tab. Once deployed, you can activate the deployment to make it available through the endpoint, deactivate it temporarily, change the model version used by the deployment, run a test query on an active deployment, modify the deployment name, description, and tags, or delete the deployment.

A model requires an attached AI compute to deploy. For more information, see [AI Compute](#).

Note

Oracle recommends using different AI compute clusters for AI agent workloads and model deployment workloads for best performance. We also recommend using different AI compute clusters for each model deployment.

Limitations

Model deployment is limited to models sized 100 MB or less.

Register a Model from an Existing Experiment Run

You can register a model from an existing experiment run in your Oracle AI Data Platform workspace.

1. On the home page, click **Models**.

2. Click **Register**.
3. Enter a name, description, and version for your model.
4. Select **Existing experiment run**.
5. From the **Workspace** drop-down list, select the workspace where your experiment is located.
6. From the **Experiment** drop-down list, select the experiment to use for your model.
7. From the **Experiment run** drop-down list, select the experiment run to use.
8. From the **Models** drop-down list, select the model to use.
9. Optional: Provide additional metadata in the form of free-form or defined tags. Click **Add** to create additional tags.
10. Optional: Provide additional information in the form of custom fields, like model type or use cases. Click **Add** to create additional custom fields.
11. Click **Register**.

View Model Details

You can see details for your registered models, like versions, metrics, artifacts, and lineages.

1. On the home page, click **Models**.
2. Click the name of the model you want to view details for.
3. Click the Overview, Details, Versions, or Artifacts tabs to see more details for your model.
4. Click Versions, then click the Run name to see the run and experiment that produced this model.

Edit Model Details

You can modify the name, description, and metadata tags for your models after creation.

1. On the home page, click **Models**.
2. Next to the model you want to edit, click **Edit**. You can also click the name of the model and click the **Actions** menu in the top right, then click **Edit**.
3. Modify the model name or description. You can also add and remove tags. Click **Save changes**.

Delete a Model

You can delete models you no longer need from your workspace.

1. On the home page, click **Models**.
2. Next to the model you want to delete, click **Delete**. You can also click the name of the model and click the **Actions** menu in the top right, then click **Delete**.
3. Click **Delete**.

Create a Model Deployment

Creating a model deployment allows you to use the model as an endpoint for agents and other functions.

Note

Ensure you have an active AI compute before you create a new model deployment. To create a new AI compute, see [Create an AI Cluster](#).

1. On the home page, click **Models**.
2. Click the name of the model you want to create a deployment for.
3. Click **Create deployment**.
4. Provide a name and description for the deployment.
5. From the **Version** drop-down menu, select which version of the model to use for this deployment.
6. Select the workspace to deploy to.
7. Select an active AI compute to use for this model deployment from the drop-down.
8. Select the authentication method.
 - For **OAuth**, you need to provide the audience claim, issuer claim, and URI for the `jwt` .JSON.
9. Optional: Provide tags to help other users and agents identify the use of this deployment. Click **Add** to create multiple tags.
10. Click **Create**.

Deploy a Model

Once a model deployment is created you need to deploy it to access it through an endpoint. You can also redeploy a model that was previously deactivated.

1. On the home page, click **Models**.
2. Click the name of the model you want to deploy.
3. Click the Deployments tab.
4. Click the Actions menu, then click **Activate**.
5. Optional: Enter a reason for activating this model deployment.
6. Click **Activate**.

Undeploy a Model

You can deactivate model deployments to make their endpoints inaccessible to others. Deactivated model deployments can be reactivated when needed.

1. On the home page, click **Models**.
2. Click the name of the model you want to deactivate the deployment for.

3. Click the Deployments tab.
4. Click the Actions menu, then click **Deactivate**.
5. Optional: Enter a reason for deactivating this model deployment.
6. Click **Deactivate**.

Send a Query to a Model Deployment Endpoint

You can test a model deployment by sending a query to the model deployment endpoint and verifying the response.

1. On the home page, click **Models**.
2. Click the name of the model you want to send a query to the deployment endpoint for.
3. Click the Deployments tab.
4. Click the Actions menu, then click **Send query**.
5. Select the query method.
6. In the **Request** field, enter the test query.
7. Click **Send request**. Verify in the **Response** field that the request was received and accepted.

Roll Forward a Model Deployment Version

Model deployments versions can be rolled forward to more recent versions as needed. Changing the model deployment version does not alter its endpoint address.

1. On the home page, click **Models**.
2. Click the name of the model you want to roll forward the version for.
3. Click the Deployments tab.
4. Click the Actions menu, then click **Roll forward**.
5. From the **Target version** drop-down menu, select the desired version.
6. Click **Roll forward**.

Roll Back a Model Deployment Version

Model deployments versions can be rolled back to older versions as needed. Changing the model deployment version does not alter its endpoint address.

1. On the home page, click **Models**.
2. Click the name of the model you want to roll back the version for.
3. Click the Deployments tab.
4. Click the Actions menu, then click **Rollback**.
5. From the **Target version** drop-down menu, select the desired version.
6. Click **Rollback**.

Edit a Model Deployment

You can modify the name, description, and metadata tags on a model deployment after it has been created.

1. On the home page, click **Models**.
2. Click the name of the model you want to edit.
3. Click the Deployments tab.
4. Click the Actions menu, then click **Edit**.
5. Modify the name, description, and tags.
6. Click **Save changes**.

Delete a Model Deployment

You can delete model deployments you own that are no longer needed.

Note

You can only delete a deactivated model deployment. See [Undeploy a Model](#).

1. On the home page, click **Models**.
2. Click the name of the model you want to delete the deployment for.
3. Click the Deployments tab.
4. Click the Actions menu, then click **Delete**.
5. Click **Delete**.

View Model Deployment Activity

You can view a history of activity on a model deployment to see what actions were taken by users and the status of those actions.

1. On the home page, click **Models**.
2. Click the name of the model you want to see the deployment activity for.
3. Click the Deployments tab.
4. At the bottom of the page, review the Activity section.
5. In the **Actions** column, click  **View details** to see the specific details for that action.

16

Customer-managed MLflow Servers

You can choose to use your own MLflow tracking servers in your Oracle AI Data Platform notebooks.

Note

Built-in integrations like model catalogue, experiments do not work when you use your own MLflow tracking server.

You can choose to integrate your own MLflow tracking server with your notebooks by including this code at the beginning of your notebook:

```
os.environ.pop("MLFLOW_TRACKING_AUTH", None)
os.environ["MLFLOW_TRACKING_URI"] = "http://<your-mlflow-host>:5000"
```

This code disables the `MLFLOW_TRACKING_AUTH` variable and sets the tracking URI to your tracking server *<your-mlflow-host>* prior to running any MLflow commands in the notebook.

Part V

AI Agents

This chapter provides information on creating, testing, deploying and monitoring agents in your workspace.

Agents are end-to-end agentic applications. Agents are defined through a graph of steps represented by nodes of different types (triggers, agents, guardrails, or tools). Agents can be defined through a no-code visual flow builder and through code via third-party libraries, such as LangGraph.

Oracle AI Data Platform offers multiple tool templates that can be configured to access your data and fit your use cases. The supported tools are:

- **Custom Tool:** The Custom Code tool lets agent developers extend AI Data Platform with their own Python code. You package your tool implementation as a ZIP file, upload it to your workspace, and configure it. The agent calls your code as a tool, with parameters supplied by the LLM at runtime.
- **HTTP:** The HTTP Request tool lets your agent call any HTTPS REST API. You configure the request, including method, URL, headers, query parameters, request body, authentication, and optionally, a response optimization step. The agent then invokes the endpoint at runtime. The HTTP request tool is available in both the visual builder and the code builder. In the code builder, the tool is configured through the `aidpUtils` Python library.
- **Prompt:** The prompt tool allows the AI developer to define a parametrized prompt that can be issued to an LLM for their choice. Common use cases for a prompt tool include email drafting tasks, translation tasks, style conversion, git commit message, and code explanations.
- **Remote MCP Server:** Agent developers can connect their agents to remote model context protocol (MCP) servers using the Remote MCP Server tool.
- **RAG:** The RAG tool lets agents pull relevant external knowledge before generating a response. In AI Data Platform, the RAG tool queries a knowledge base (23ai Vector Search) and retrieves semantically relevant document chunks. Those chunks are then passed to the agent for response generation.
- **SQL:** The SQL tool enables agents to execute SQL queries against structured data sources registered via external catalogs, such as Oracle Autonomous AI Lakehouse, Oracle Autonomous AI Transaction Processing, or Oracle Autonomous AI Database. The tool is intended for scenarios where the SQL queries are predefined and can be parametrized. The objective is to let an agent assign values to the parameters. This tool is not an NL2SQL tool that generates a SQL query based on a natural language prompt.

Note

The SQL tool only performs queries against data in an external catalog. It does not support data stored in a standard catalog.

 Note

You must attach an AI Compute to your agent before you can test a system tool. If no compute is attached, the Test tab is disabled.

Creating agents in AI Data Platform generates an agent artifact file (.aflow) in the workspace folder that you select. This file can't be modified.

Agent Creation

This section covers the creation of AI agents through the visual flow builder or through code.

Topics:

- [Multi-agent Systems and Supervisor Patterns](#)
- [About the Visual Flow Canvas](#)
- [Agents Through Code](#)
- [Agent Testing](#)

Multi-agent Systems and Supervisor Patterns

A multi-agent system is an AI application design in which a user request is handled by multiple cooperating agents instead of one large, all-purpose agent.



Each agent has its own role, instructions, model configuration, memory policy, and allowed tools. The flow defines how the request moves between those agents and how the final answer is produced.

This design is useful when a workflow naturally separates into specialist responsibilities. For example, one agent can retrieve data, another can call an API, another can summarize findings, and a supervisor can decide which specialist to use and combine the results into a single response.

Note

As a design principle, it's best to start with the smallest agent design that meets the requirements. Add multiple agents when separation of concerns improves reliability, security, maintainability, or observability more than it increases cost and complexity.

Benefits of Multi-agent Systems

Multi-agent systems are best suited for:

- **Specialization:** give each agent a focused job, prompt, and tool set instead of one crowded instruction block.
- **Routing and decomposition:** let a supervisor interpret the request, split it into subtasks, and choose the right specialist for each subtask.
- **Tool and data isolation:** expose sensitive or high-impact tools only to the agents that are responsible for using them.
- **Governance and troubleshooting:** make handoffs, tool ownership, memory settings, and failure points easier to inspect.

When to Choose Multi-Agent or Single Agent Designs

A single agent with more tools is often the right first design. It is simpler to test, cheaper to run, and easier to reason about when the task has one clear objective and one permission model. Use a multi-agent design when the workflow benefits from explicit roles, bounded tool access, or a supervisor that can coordinate multiple specialist outputs.

Design Question	Use single agents when...	Use multi-agents when...
Task shape	The request has one main objective and one response stile.	The request must be decomposed, routed, verified, or synthesized across specialties.
Tools and data	The same instruction set and permission model can safely govern all tools	Different agents need different tools, data sources, or access boundaries.
Instructions	The prompt remains clear even with all business rules and tool guidance in one place.	Instructions are easier to maintain as smaller, role-specific prompts.
Cost and latency	You want the shortest path from user message to answer.	The reliability, governance, or maintainability benefits justify extra orchestration.
Troubleshooting	Failures are simple to debug in one trace.	You need explicit handoffs, state isolation, and clearer ownership for each step.

Supported Pattern: Orchestrator/Supervisor

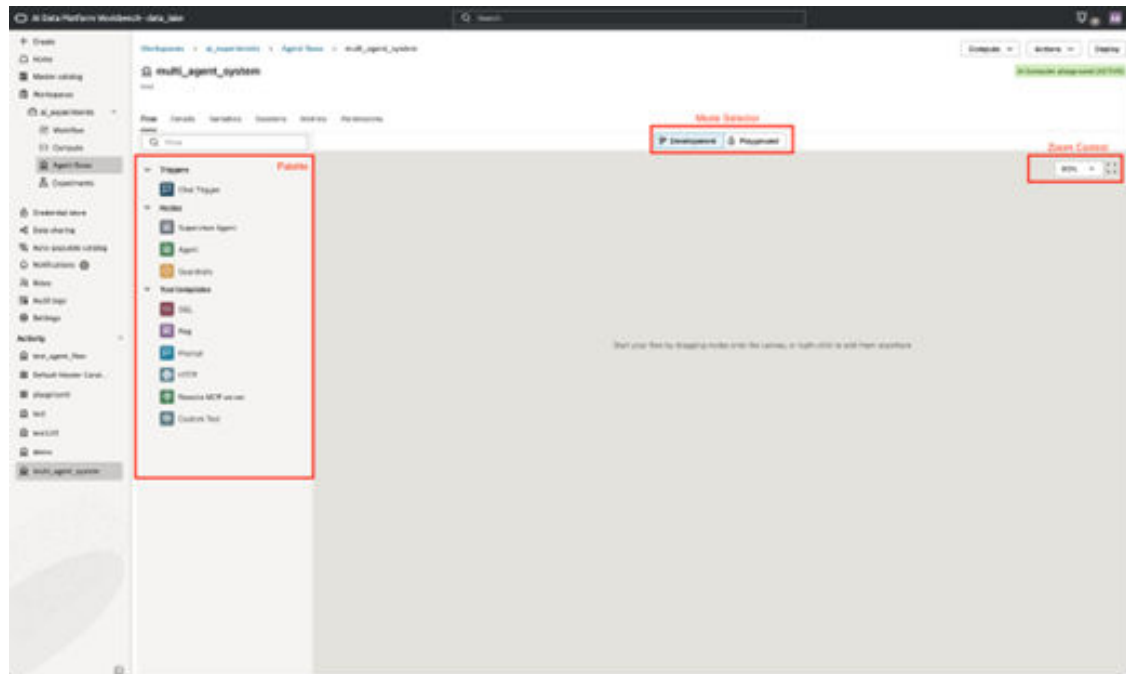
The current canvas experience supports the orchestrator / supervisor pattern. In this pattern, the Chat Trigger receives the user message, optional Guardrails evaluate the input, and a Supervisor Agent acts as the orchestrator for the rest of the flow.

The supervisor should focus on planning, routing, delegation, and final response synthesis. It decides which agent should handle a task, sends that agent a scoped instruction, reviews the result, and then either delegates another step or returns the final response. Agents should be narrower specialists: they do the assigned work, use their attached tools, and return useful results to the supervisor.

About the Visual Flow Canvas

An agent is assembled by dragging nodes and tool templates from the left palette onto the canvas, then connecting the nodes in the order the request should travel.

Selecting a node opens a configuration panel at the bottom of the screen.



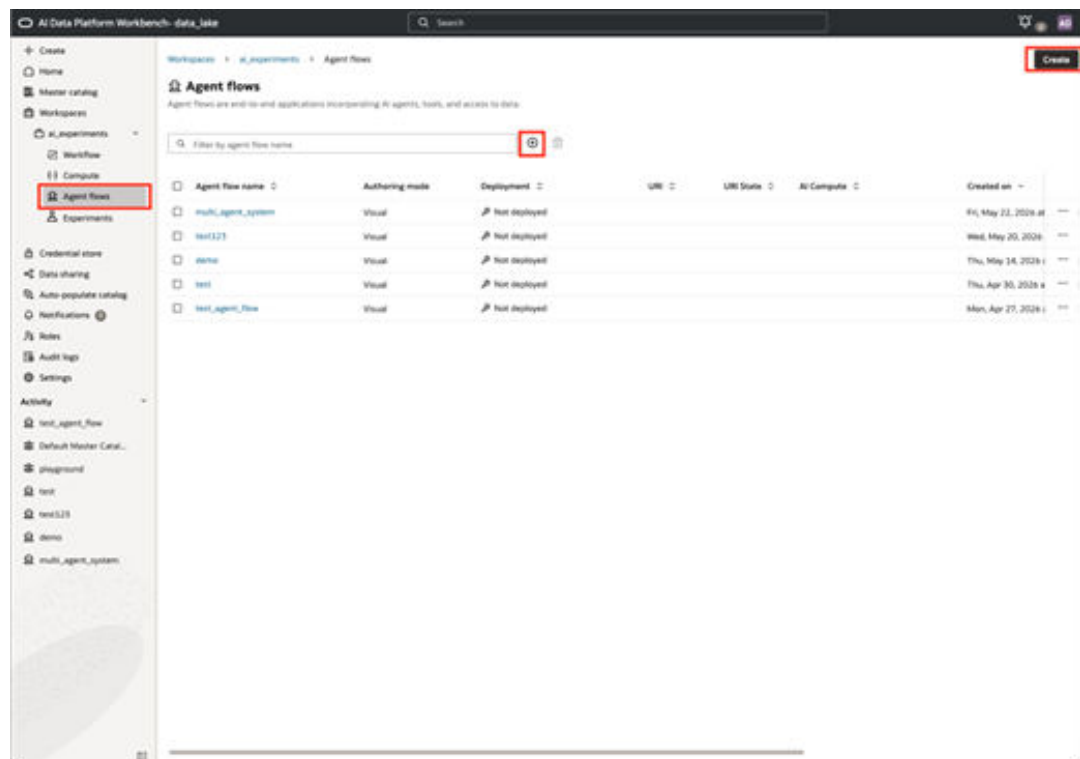
Canvas Element	Purpose
Chat Trigger	Entry point for a user message. In the screenshot this node is labeled Message and typically sits at the top of the flow. A chat trigger node can be connected to an agent, a supervisor agent, or a guardrails node. Only one chat trigger is allowed per canvas. Chat trigger nodes can be configured to enable agents to read provided files and write output in specified file formats. The file types supported depend on the model used. For supported file types, look for specific models in Offered Pretrained Foundational Models in Generative AI.
Guardrails	Optional policy and safety layer placed before or after model work. The guardrails policies include PII, content moderation, and prompt injection detection. A guardrails node can filter traffic between a chat trigger and an agent node, between a supervisor and executor agents, or between agent and tool nodes. We recommend a single guardrails node between the chat trigger and the agent node.
Supervisor Agent	The orchestrator. It receives the user request, decides which executor agent or tool should handle each task, and coordinates the final answer. Only one supervisor agent is allowed in a canvas.
Agent	An executor agent. Each executor should have a clear specialty, such as data retrieval, API lookup, summarization, or document question answering. Use an agent / executor agent for a single agent system.

Canvas Element	Purpose
Tool templates	Reusable capabilities that can be attached to individual executor or supervisor agent. Tool templates include SQL, RAG, Prompt, HTTP, Remote MCP server, and Custom Tool.
Development / Playground	Mode selector above the canvas. Development is used while editing the agentic system; Playground is used to initiate test sessions and inspect agent behavior. Playground requires that an AI compute is attached to your agent.
Zoom control	Canvas zoom selector. The screenshots show 60 percent and 90 percent zoom levels.

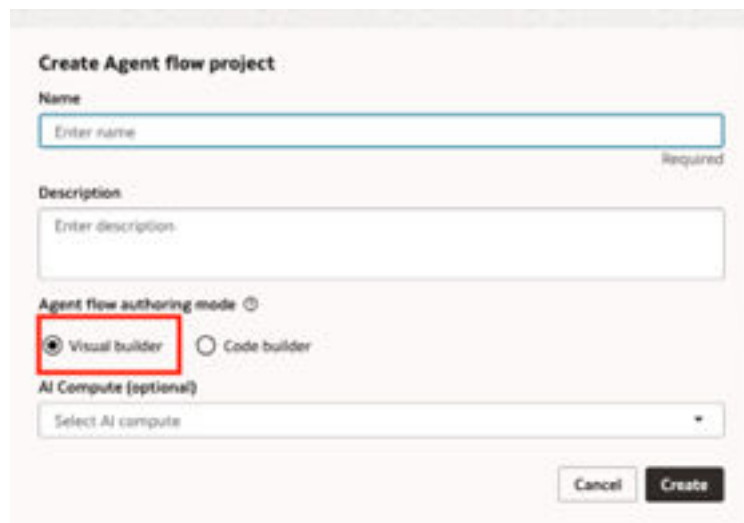
Create an Agent

You can create an agent in a workspace where you have Manage permission.

1. On the Home page, navigate to your workspace.
2. Click on **Agents** in the left navigation pane.
3. Click **Create Agent** or click **Create** in the top right.



4. Provide a name and description for the agent.
5. For **Agent flow authoring mode**, select **Visual builder**.



6. Optional: From the **AI Compute** drop-down menu, select a compute to use for the agent.
7. Click **Create**. Start building your agent by dragging a node from the palette to the canvas.

Note

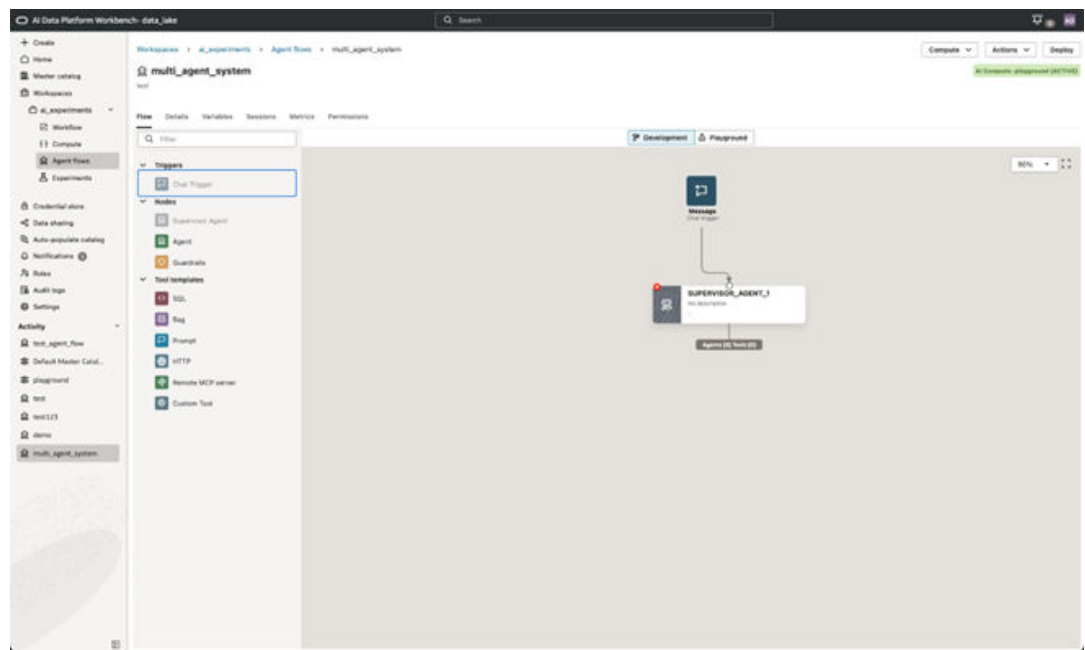
Start your first agent build simple: one Chat Trigger, one Agent. Add complexity after a successful run of your first build, such as guard rails, additional tools, or even multi-agent system design.

Add Chat Trigger and Agent to Visual Builder Canvas

Your first step after creating an agent with the Visual Builder should be to add a chat trigger and a supervisor agent.

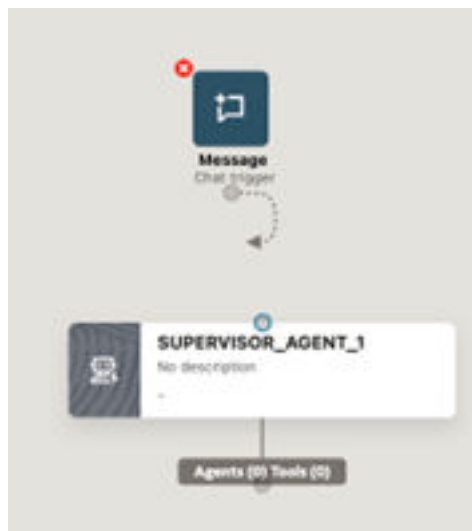
The trigger receives the user message. The supervisor interprets the request, plans the work, and delegates to executor agents or tools. You can drag nodes in the canvas, configure them, and connect them later.

1. Navigate to your agent in your workspace.
2. Click and drag a **Chat Trigger** from the palette to the canvas. The node appears on the canvas as a Message.
3. Click and drag a **Supervisor Agent** to the canvas.



4. Click and drag the connector handle on the Chat Trigger node to connect it to the Agent node.

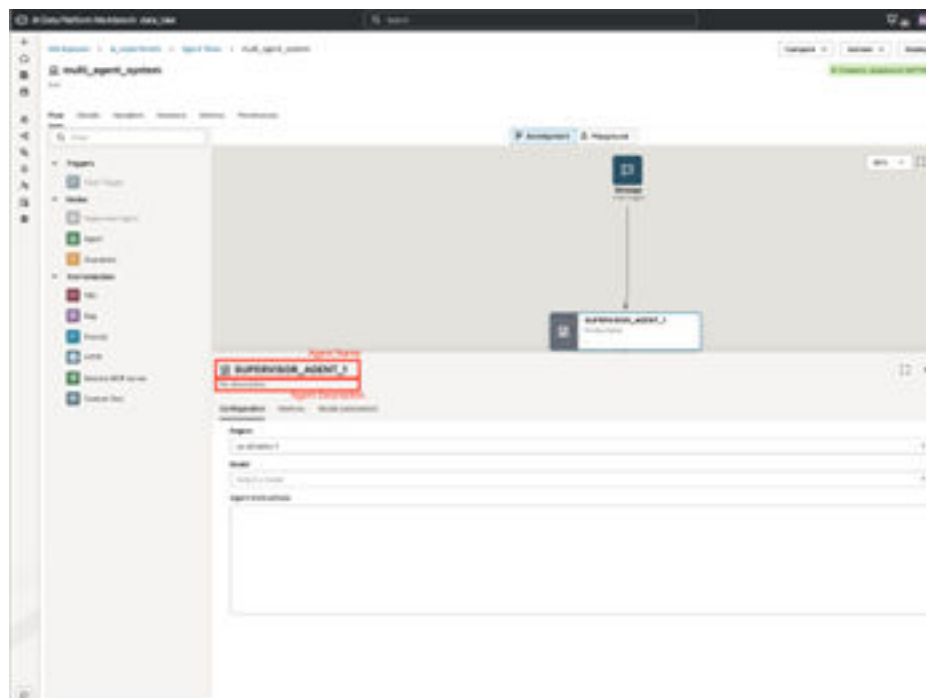
The Supervisor Agent badge displays how many agent and tools are connected. In a new build, the Supervisor Agent shows: Agents (0) Tools (0).



Configure a Supervisor Agent

You need to configure a supervisor agent added to your Visual Builder canvas with instructions outlining the supervisor role.

You configure a supervisor agent with the following fields.



Field	Configuration
Agent Name	Provide a descriptive name for your supervisor agent. A good, descriptive name will be beneficial when debugging the system behavior through traces and logs.
Agent Description	Provide a description of the agent purpose, role, and general behavior. Useful for documentation purposes.
Region	Choose the region where the OCI Generative AI model used by Supervisor Agent is hosted. See Generative AI Models by Region .
Model	Choose the OCI Generative AI service model used by the supervisor. The dropdown lists the models available in the region you selected.
Agent instructions	Describe the supervisor role, routing rules, delegation policy, tool-use expectations, and final response format.

1. Navigate to the agent in your workspace.
2. Click the **Supervisor Agent** node on your canvas.
3. Provide a meaningful name and description for your supervisor agent.
4. Enter the region and model for the OCI Generative AI service model used by the supervisor.
5. Provide the agent instructions for your Supervisor Agent.

Suggested Supervisor Instructions

You should use the Instructions field for a Supervisor Agent to make the supervisor responsible for orchestration, not for doing every task itself.

Keep the instructions concrete so routing decisions are predictable. See the following for an example of a set of Supervisor instructions:

You are the supervisor for a multi-agent system.

Responsibilities:

- Understand the user's request and break it into subtasks.
- Select the most appropriate agent or tool for each subtask.
- Do not perform specialist work yourself when an agent is available.
- Ask for clarification only when required information is missing.
- Combine agent outputs into a concise final answer.
- Mention important assumptions, limits, or failed tool calls in the final answer.

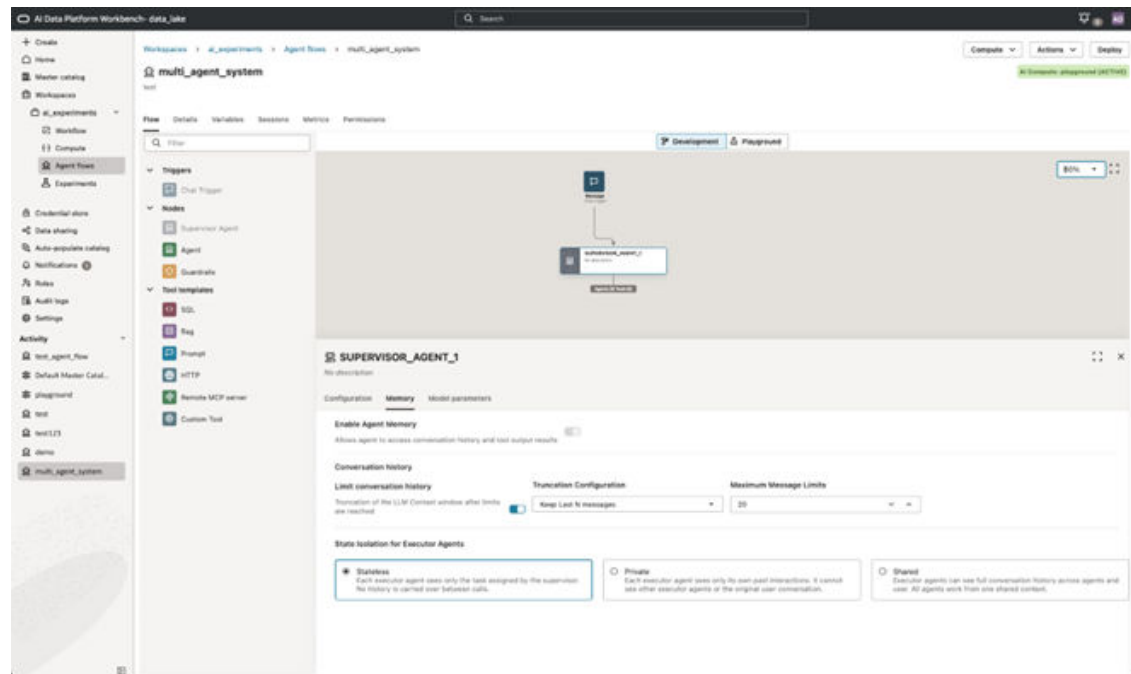
Routing rules:

- Use the SQL agent for structured data questions.
- Use the HTTP agent/tool for external API lookups.
- Use the RAG agent/tool for document or knowledge-base questions.
- Use the prompt tool for reusable prompt-only transformations.

Configure Supervisor Agent Memory and State Isolation

The Memory tab for a Supervisor Agent controls how much conversation and tool-output history is available to the supervisor and how much context is shared with executor agents.

You configure the memory and isolation state for your Supervisor Agent with the following fields.



Field	Configuration
Enable Agent Memory	Enable when users need multi-turn continuity. Disable for isolated, one use tasks. This field cannot be disabled for Supervisor Agents.
Limit conversation history	Enable to truncate the LLM context window after the specified limit is hit. Disable to show the full history.
Truncation configuration	If Limit conversation history is enabled, use this field to set the conditions for truncating the the context window. Options are: - Keep last N messages - Token budget - Both
Maximum Message Limits and Token Budget	One or both of these options is displayed, depending on your choice for Truncation Configuration . Default values are 20 messages and 5000 tokens. We recommend starting with moderate values and adjusting as needed.
State Isolation for Executor Agents	Select Stateless , Private , or Shared . Stateless: Each executor agent sees only the task assigned by the supervisor. No history is carried over between calls. Select this if you want the strongest isolation and least cross-agent context. Private: Each executor agent sees only its own past interactions. It cannot see other executor agents of the original user conversation. Select this if your executor needs continuity across its own tasks but should not share context with other agents. Shared: Executor agents can see full conversation history across agents and users. All agents work from one shared context. Select this if you need broad context sharing and have reviewed privacy and prompt-injection risks.

1. Navigate to the agent in your workspace.
2. Click the **Supervisor Agent** node on your canvas.
3. Click the Memory tab.
4. Choose whether to enable **Limit conversation history**. Select a **Truncation Configuration** and set limits, if enabled.
5. Choose an option for **State Isolation for Executor Agents**.

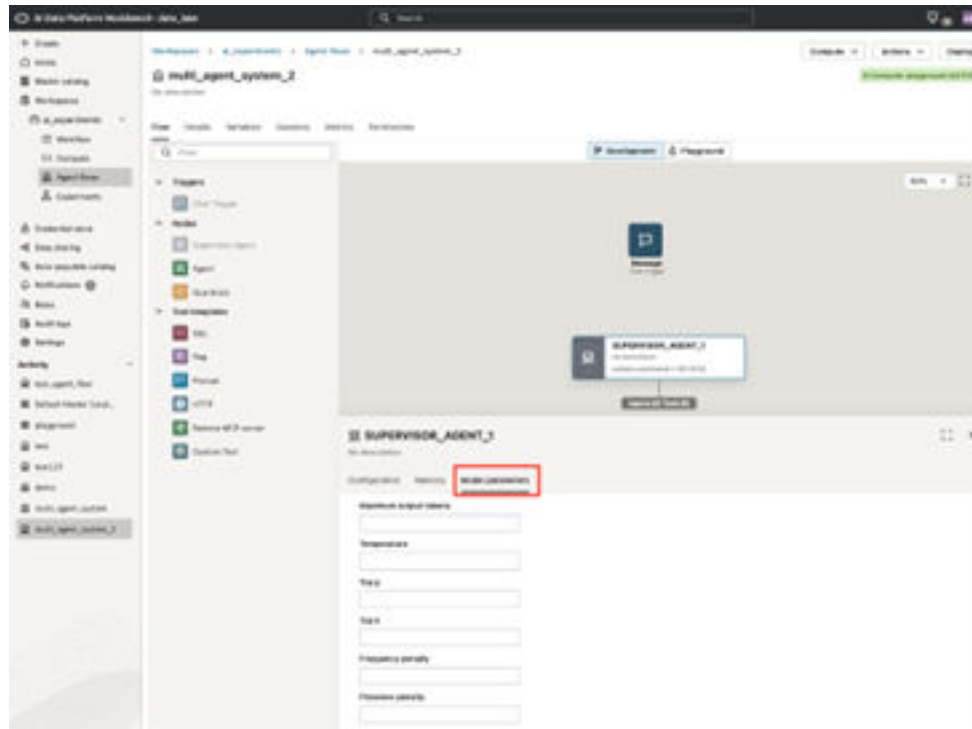
Models Parameter Tab

The model parameters tab lets you configure model-specific parameters that are available for the selected model.

Model parameters can be configured separately for supervisor and executor agents. Parameters you can use include temperature, top K, top P, and frequency penalty.

Note

Only a subset of models exposes configurable parameters. Furthermore, the parameters vary across model families.



Add Guardrails to an Agent

You can add additional layers of protection to your agents by adding one or more guardrail nodes to your canvas.

By default, no guardrails are applied to your agentic systems beyond what the selected model provider is offering out-of-the-box for their models. Guardrails can be placed between the Chat Trigger and the Supervisor Agent so policies are applied before a request reaches the supervisor agent and before the supervisor agent returns a response to the caller.

Guardrail	Options	When to Use
Personal identifiable information (PII)	- Input and Output tabs - Checkboxes for Person, Address, Telephone number, Email	Use when the flow must block or mask sensitive personal data before or after model processing.
Content moderation prevention	Input and Output rows with Block, Inform, and Allow options.	Use to define how the flow handles hate, sexual, violent, toxic, derogatory, or harassing content.
Prompt Injection detection	Input row with Block and Allow options.	Use to reduce the chance that malicious instructions override the system or agent instructions.

For more information on guardrail settings, see [Guardrails](#).

1. Navigate to the agent in your workspace.
2. Drag a **Guardrails** node from your palette to your canvas. Place it between your Chat Trigger node and Supervisor Agent node.
3. Delete the connection between Chat Trigger and Supervisor Agent by hovering over the connection and clicking the red X.

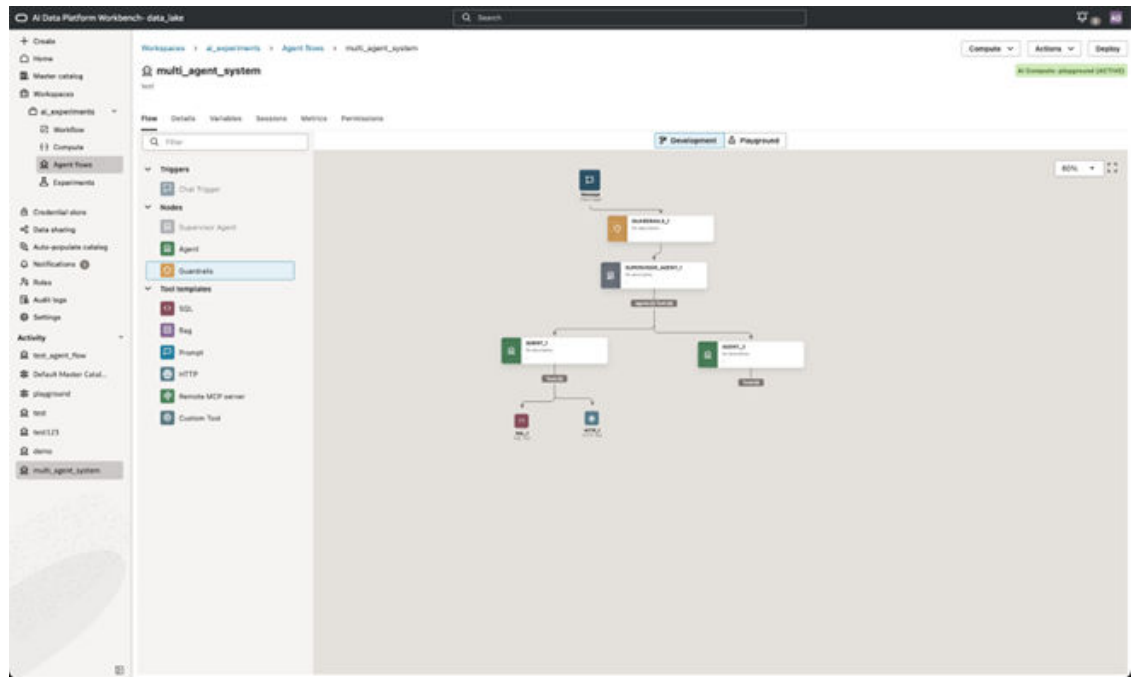


4. Click and drag the connector handle on the Chat Trigger to the Guardrail node. Then click and drag the connector handle from the Guardrail node to the Supervisor Agent.
5. Click the Guardrail node to open the Configuration page.
6. Configure the guardrails to select the desired action for input and output checks.

Add Executor Agents and Tools to an Agent

You can add executor agents to tools to perform specialized work for the supervisor agent.

In the example below, the Supervisor Agent delegates to AGENT_1 and AGENT_2. AGENT_1 is connected to SQL_1 and HTTP_1 tools.

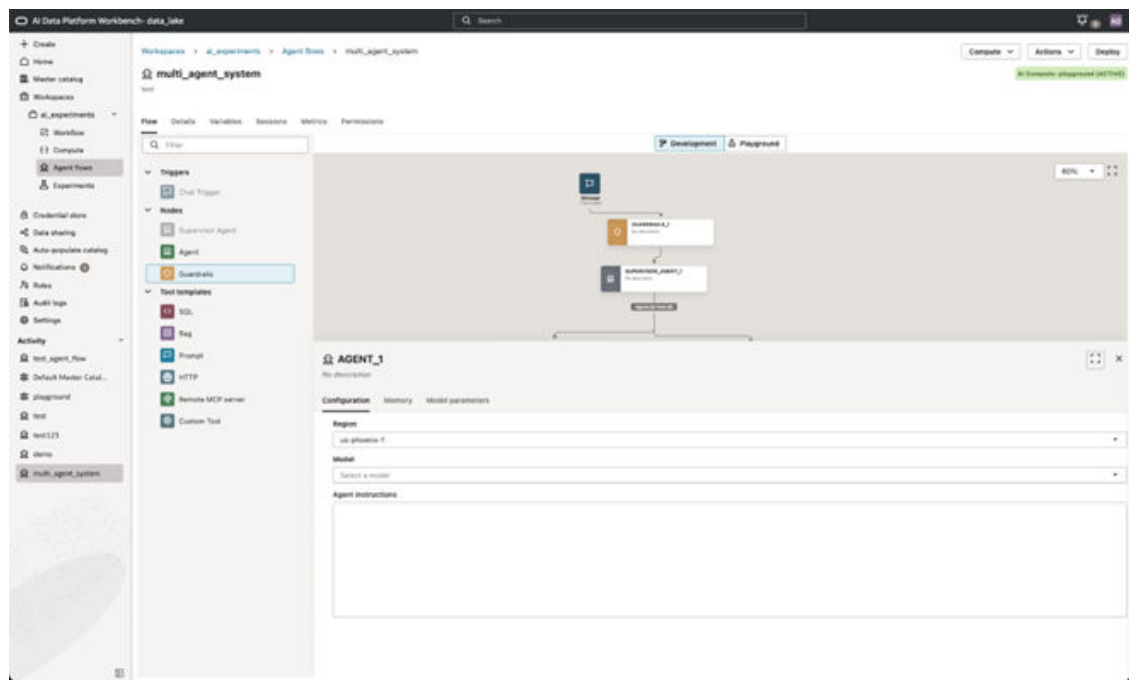


1. Navigate to the agent in your workspace.
2. Drag an Agent node from the palette to your canvas. Agent nodes should be placed below a Supervisor Agent.
3. Drag Tools from the palette to your canvas.
4. Click and drag the connector handle on your Supervisor Agent to connect to the Agent nodes.
5. Click and drag the connector handle on your Agents to connect to the Tool nodes.

Executor Agent Configuration

Agent nodes can be configured by modifying settings on their Configuration, Memory, and Model tabs to help you define the purpose of each agent.

Agents should be configured narrowly, given a specific function and goal, so the supervisor agent can route work reliably.

**Table 17-1 Agent Configuration Tab**

Field	Configuration
Agent Name	Best practice is to name each executor agent according to its specialty, such as SQL_AGENT, DOCUMENT_AGENT, API_AGENT, or SUMMARY_AGENT. The name of each executor agent is visible to the supervisor agent, so use descriptive names.
Agent Description	Provide a detailed description of each executor agent. The description of each executor agent is visible to the supervisor agent.
Region	Choose the region where the OCI Generative AI model used by the Agent is hosted. See Generative AI Models by Region .
Model	Choose the OCI Generative AI service model used by the agent. The drop-down menu lists the models available in the region you selected. Select a model that fits the executor task. Executor agents don't need to use the same model as the supervisor agent.
Agent instructions	Describe exactly what the executor should do, what tools it may use, and what output structure it should return.

Executor Agent Memory Tab

In the case of executor agents connected to a supervisor agent, the memory for executors is configured in the supervisor node and applied to all executor agents.

Field	Configuration
Enable Agent Memory	Enable when users need multi-turn continuity. Disable for isolated, one use tasks.
Limit conversation history	Enable to truncate the LLM context window after the specified limit is hit. Disable to show the full history.
Truncation configuration	If Limit conversation history is enabled, use this field to set the conditions for truncating the the context window. Options are: • Keep last N messages • Token budget • Both
Maximum Message Limits and Token Budget	One or both of these options is displayed, depending on your choice for Truncation Configuration. Default values are 20 messages and 5000 tokens. We recommend starting with moderate values and adjusting as needed.
State Isolation for Executor Agents	Select Stateless, Private, or Shared. • Stateless: Each executor agent sees only the task assigned by the supervisor. No history is carried over between calls. Select this if you want the strongest isolation and least cross-agent context. • Private: Each executor agent sees only its own past interactions. It cannot see other executor agents of the original user conversation. Select this if your executor needs continuity across its own tasks but should not share context with other agents. • Shared: Executor agents can see full conversation history across agents and users. All agents work from one shared context. Select this if you need broad context sharing and have reviewed privacy and prompt-injection risks.

Executor Agent Model Parameters Tab

The model parameters tab lets you configure model-specific parameters that are available for the selected model.

Note

Only a subset of models exposes configurable parameters. Parameters also vary across model families.

Examples of parameters include temperature, top K, top P, and frequency penalty. Model parameters can be configured separately for supervisor and executor agents.

Suggested Executor Instructions

You are the SQL executor agent.

Responsibilities:

- Translate the supervisor's task into safe SQL tool usage.
- Use only the SQL tools attached to this agent.
- Return a concise answer plus any important query assumptions.
- Do not invent data. If the tool cannot answer, say what is missing.
- Return structured output with: answer, evidence, assumptions, and follow_up_needed.

Checklist for Agents through Visual Builder

Use this list as a guide to ensure you've included and configured every necessary component for an agent built using the Visual Builder.

Build Checklist

- The agent has exactly one expected entry point: Chat Trigger / Message.
- Guardrails are connected in the intended position and enabled where required. We recommend inserting guardrails between the trigger message and the agent.
- The Supervisor Agent has a selected region, selected model, and orchestration instructions. Same for executor agents.
- Configure memory of the multi-agent system in the Memory tab of the Supervisor agent. Select executor state isolation that matches the privacy and continuity requirements.
- Each executor Agent has a clear specialty and narrow instructions.
- Each tool is attached only to the agent that should use it.
- No node is disconnected.
- An AI compute is attached to the agentic system to test individual tools and for running the Playground experience.

Table 17-2 Common Issues

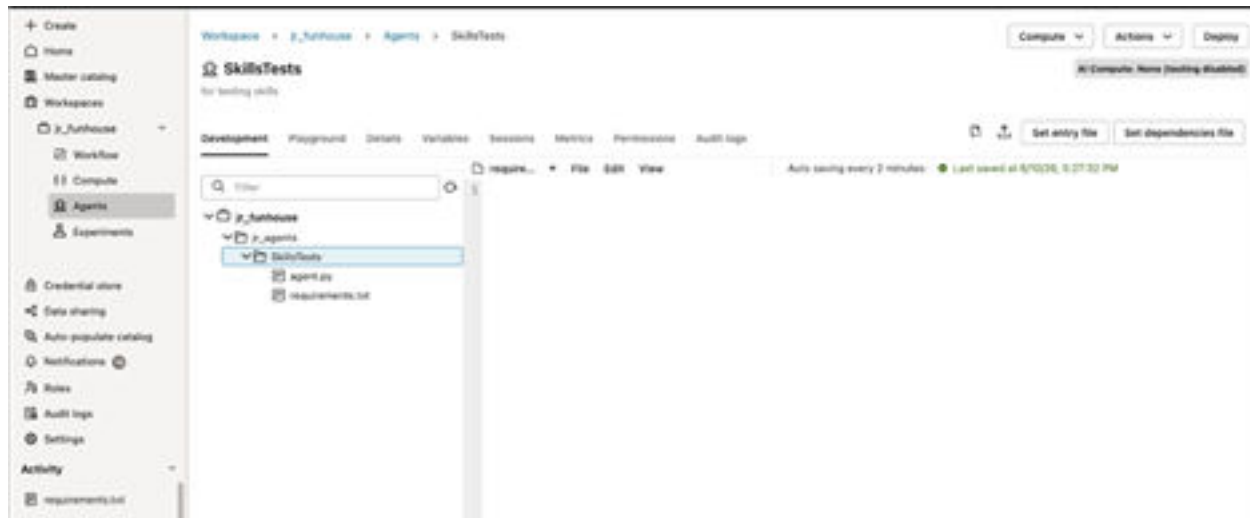
Issue	Likely Cause	Suggested Action
Supervisor does not call an executor	Supervisor instructions are too vague or no executor is connected.	Add explicit routing rules and confirm the executor node is connected to the supervisor.
Executor returns broad or off-topic answers	Executor instructions are too general.	Make the executor role narrower and define the required output structure.
Tool is not used	Tool is disconnected or attached to the wrong agent.	Check the tool connection and the agent tool count badge.
Guardrail does not fire	Guardrail section is configured but not enabled.	Open the guardrails node and confirm the section toggle is on.
Context leaks across agents	State isolation is set to Shared or memory is broader than intended.	Use Stateless or Private isolation for stricter separation.
Follow-up questions lose context	Memory is disabled or truncation is too aggressive.	Enable memory and adjust the maximum message limit.

Agents Through Code

You can bring your own LangGraph code base to AI agents in Oracle AI Data Platform or create a brand new LangGraph agent directly on the platform through the agent coding experience.



You can use the AI Data Platform utility Python library `aidutils` to configure your foundational model and import system tools to your agent. For `aidutils` API reference, see [Aidp-utils API for Oracle AI Data Platform](#).

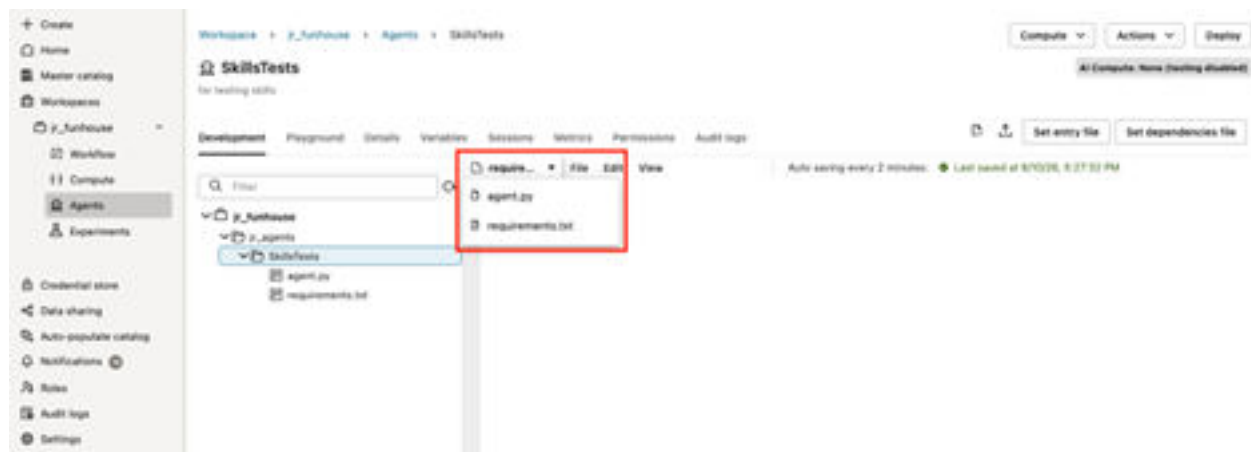


You create an agent through code by either uploading an existing code file or creating code files directly in your agent through the inline editor.

The inline code editor in agents supports the following code file types:

- Python (.py)
- JSON
- TXT
- CSV
- PSV
- SH
- Folder

You can see and navigate through the available code files by clicking the file selector drop-down list.



Entry and Dependency Files

Entry files are code files that have the class with setup and invoke methods expected for an agent defined as code. AI Data Platform requires you to set an entry file for agents through code.

Dependency files are files that include third-party libraries required by your agent defined as code. Dependency files are typically `requirements.txt` files that contain a list of the required third-party libraries.

Note

Third-party libraries are installed when you test your code in the editor by clicking the Play button or when you test the agent through the Test tab. We recommend install third-party libraries by testing the code first. Errors during installation of the libraries are displayed in the output cell.

Agent Class

AgentBasic is a template class for setting up and invoking a simple conversational agent using a stateful LangGraph workflow. It demonstrates the structure required for minimal agent development with two main methods:

- `setup()`: Initializes the agent workflow and defines the graph.
- `invoke(user_query, **kwargs)`: Runs the agent on a user message and returns the response.

It can be directly run and tested using a `main()` function before integration into a larger system.

Definition

```
class AgentBasic:
    def __init__(self) -> None:
        self.graph = None
    def setup(self) -> None:
        self.graph = StateGraph(MessagesState)
        self.graph.add_node(mock_llm)
        self.graph.add_edge(START, "mock_llm")
        self.graph.add_edge("mock_llm", END)
```

```

self.graph = self.graph.compile()
system_prompt = "Be a helpful assistant."
async def invoke(self, user_query: str, **kwargs):
    user_message = HumanMessage(content=user_query)
    messages = {"messages": [dict(user_message)]}
    try:
        return self.graph.invoke(messages)
    except Exception as e:
        import traceback
        logger.error(f"Exception while calling invoke {e}", exc_info=True)
        print("Stack trace:\n", traceback.format_exc())

```

Test Invocation

This test invocation is ideal for initial functional testing.

Note

Include a main entry point for stand-alone testing.

```

import asyncio

async def main():
    test_agent = AgentBasic()
    test_agent.setup()
    result = await test_agent.invoke("Hi there")
    print("Agent response:", result)
if __name__ == "__main__":
    asyncio.run(main())

```

How it works:

- The script creates an agent, sets it up, and sends a sample user message.
- The agent responds (`{"messages": [{"role": "ai", "content": "hello world"}]}` in this example).

Usage Guide

Create an Agent class with the setup and invoke methods.

setup()	Initializes the agent workflow	agent.setup()
invoke()	Runs the agent with a user message	await agent.invoke("Your question")

- **Asynchronous:** invoke() is an async method; use it with await or run in an async loop.
- **Testing:** The included main() guard (if __name__ == "__main__":) makes it easy to test the agent before deployment.

Build an Agent Through Code by Upload

You can build your end-to-end agent application with existing code by uploading your LangGraph code base.

Oracle AI Data Platform supports LangGraph version 1.0.1.

Note

You can upload individual files and folders up to a maximum of 500 files, each file can have a maximum size of size of 500MB. The upload is limited to a total size of 5GB.

1. Navigate to your agent in your workspace. Click the agent name.
2. Click **Upload**.



3. Drag and drop a file into the pane or click to browse to select a file.
4. Click **Upload**.

Build an Agent Through Code by Creating New Code

You can build your end-to-end agent application with existing code by creating code directly in your agent through the code editor.

Code editor supports the following file types:

- Python (.py)
- JSON
- TXT
- CSV
- PSV
- SH
- Folders

1. Navigate to your agent in your workspace. Click the agent name.
2. Click **Add new file**.

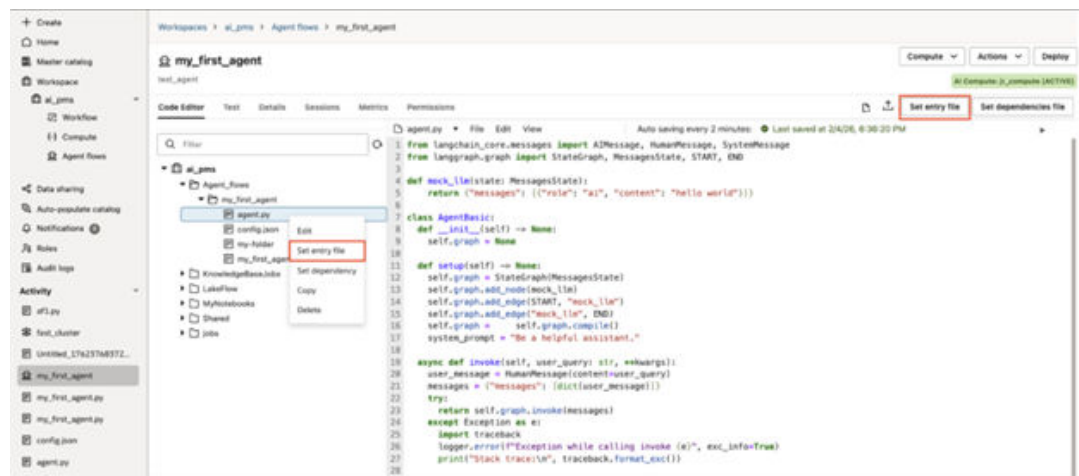


3. Enter a name for your code file.
4. Select a file type from the drop-down list.
5. Click **Create**.

Set an Entry File for Agents through Code

Your AI agent through code requires an entry file that has the required class, setup, and invoke methods expected for your agent.

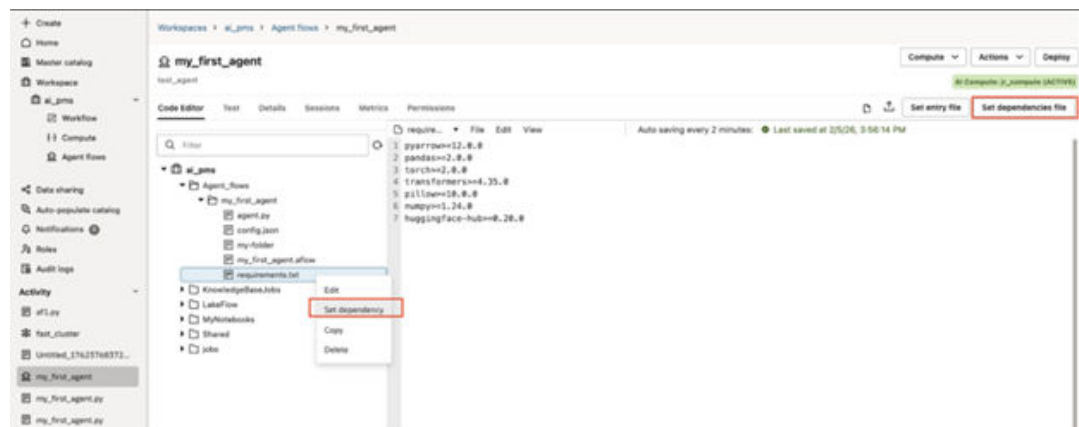
1. Navigate to your agent in your workspace. Click the agent name.
2. In the Code Editor tab, locate the entry file in the left navigation pane. If the file does not exist, you can upload it by clicking **Upload** or create it by clicking **Add new file**.
3. Right-click the entry file and click **Set entry file**. You can also select the file and click the **Set entry file** button in the top right of the code editor.



Set a Dependency File for Agents through Code

You need to set a dependency file for agents flows through code that contains any third-party libraries your code is dependent on.

1. Navigate to your agent in your workspace. Click the agent name.
2. In the Code Editor tab, locate the dependency file in the left navigation pane, typically `requirements.txt`. If the file does not exist, you can upload it by clicking **Upload** or create it by clicking **Add new file**.
3. Right-click the dependency file and click **Set dependency**. You can also select the file and click the **Set dependencies file** button in the top right of the code editor.



Test Agent Code

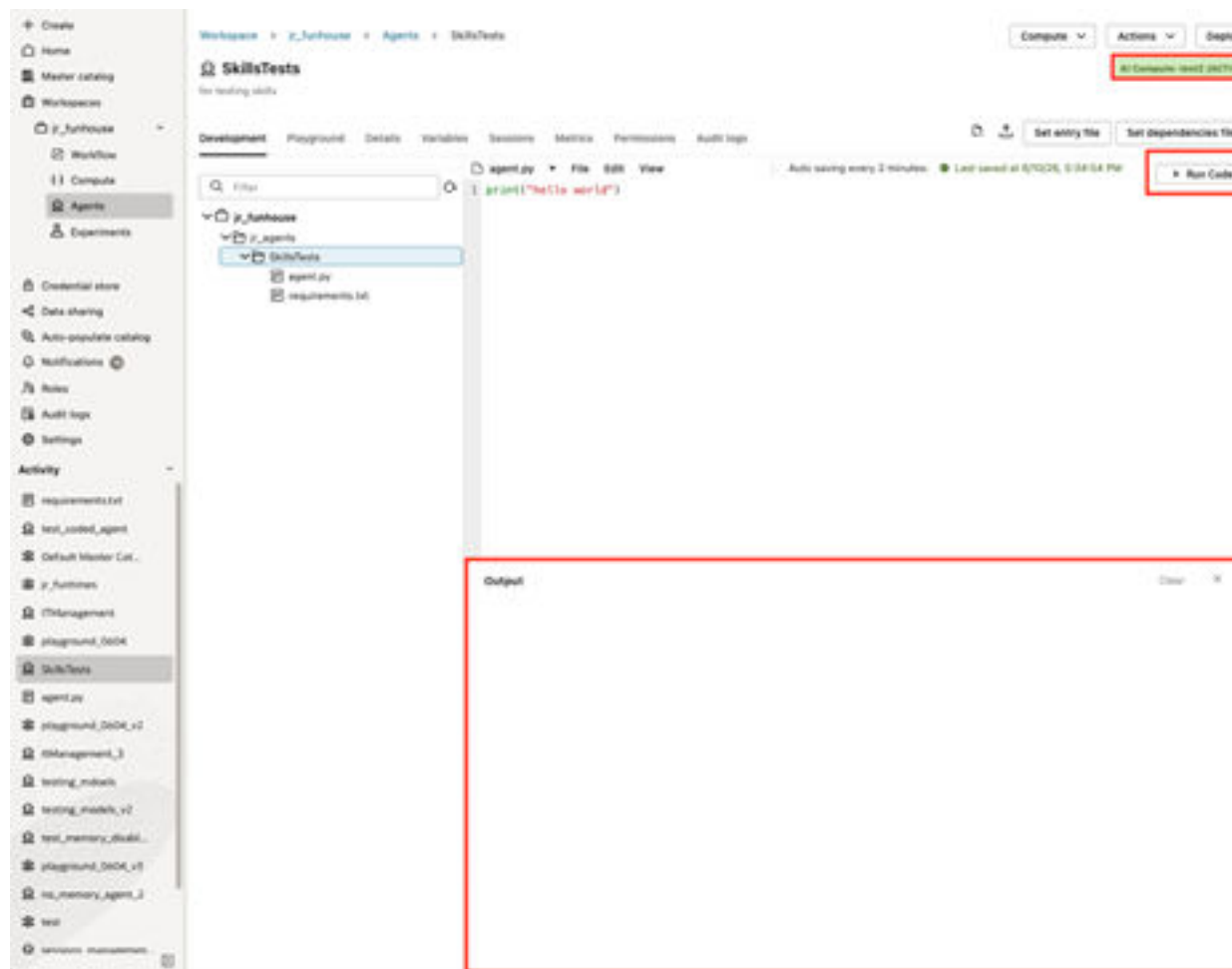
You can test the code used for your agent from the Test tab to validate and debug code.

You must have an AI compute attached to your agent to test.

1. Navigate to your agent in your workspace. Click the agent name.
2. Click the **Playground** tab.



3. Click **Play** to test the selected code file.



An output cell in the lower half of the code editor window displays the outputs of print or logging statements in your code. Errors are also displayed in the output cell.

Agent Skills in Coding Experience

Agent Skills let an agent discover and use task-specific instructions, reference files, templates, assets, and optional executable scripts without hardcoding that domain knowledge into the agent's instructions.

A skill is stored as a folder in your agent code base. Each skill has a required `SKILL.md` file that describes what the skill does and how the agent should use it. A skill can also include supporting files such as schemas, examples, prompts, templates, assets, or scripts.

For more information, see [Agent Skills Overview](#).

Agent skills support a progressive disclosure model:

1. The agent discovers that a skill exists.
2. The agent activates the skill only when it is relevant.
3. The agent loads additional files from the skill folder only when needed.
4. The agent may run an explicitly declared skill entrypoint, if the skill allows it.

When to use Agent Skills

You should use Skills when you want to package reusable agent capabilities such as:

- Domain-specific instructions
- Coding or data analysis workflows
- SQL generation guidance
- Business process playbooks
- File templates
- Schema references
- Reusable scripts for safe calculations, transformations, or lookups

Skills are useful when the agent should have access to specialized and re-usable knowledge, but you do not want to place all that knowledge directly in the agent prompt.

How Skills Work at Runtime

At runtime, the host application determines which skill directories are available, such as project-level and user-level skill folders. The platform loads each skill's metadata from SKILL.md and builds a catalog keyed by skill name.

The agent can then use skill-related tools:

Tool	Purpose
<code>activate_skill(name)</code>	Loads the skill instructions from SKILL.md.
<code>list_skill_files(name, path)</code>	Lists files available inside a skill folder.
<code>load_skill_file(name, path)</code>	Loads a supporting file from the skill folder.
<code>run_skill_entrypoint(name, entrypoint, args_json, timeout_seconds)</code>	Runs an explicitly declared Python entrypoint, if allowed by the skill.

Some environments may also embed a summary of available skills directly into the system prompt. In that setup, the agent can discover available skills from the prompt, then use `activate_skill` when it needs the full instructions.

Skill Folder Structure

A skill uses an Agent Skills-style folder layout:

```
<skills_dir>/
  some-skill/
    SKILL.md
    references/
    ...
    scripts/
    ...
    assets/
    ...
```

Only SKILL.md is required. The other folders are optional.

Folder or File	Required	Purpose
SKILL.md	Yes	Main skill metadata and instructions.
references/	No	Supporting documentation, schemas, examples, or templates.
scripts/	No	Python scripts that may be run only when explicitly declared as entrypoints.
assets/	No	Static assets used by the skill.

Writing SKILL.md

Each skill must include YAML frontmatter at the top of SKILL.md, followed by Markdown instructions.

Basic Example

```
---
name: sql-helper
description: Helps the agent write safe SQL queries using project schemas.
license: internal
compatibility: "agent-platform"
metadata:
  owner: data-platform
  domain: analytics
allowed-tools: "analyzeQuery inspectSchema"
---
```

SQL Helper

Use this skill when the user asks for SQL generation, query review, or schema-aware analysis.

Before writing SQL:

1. Inspect the relevant schema files in `references/`.
2. Prefer explicit column names.
3. Avoid destructive statements unless the user explicitly asks for them and the environment allows them.

Table 17-3 Supported Frontmatter Fields

Field	Required	Description
name	Yes	Unique skill name used by the catalog and tools.
description	Yes	Short description used for discovery and routing.
license	No	License or usage policy for the skill.
compatibility	No	Compatibility note for supported runtimes or platforms.
metadata	No	String-to-string metadata map.

Table 17-3 (Cont.) Supported Frontmatter Fields

Field	Required	Description
allowed-tools	No	Space-separated list of tools this skill permits.
entrypoints	No	List of executable entrypoints declared by the skill.

Adding Supporting Files

Supporting files let a skill keep detailed content outside the main instructions. This keeps SKILL.md focused while still giving the agent access to richer context. For example:

```
skills/
  sql-helper/
    SKILL.md
  references/
    warehouse_schema.md
    query_style_guide.md
    examples.md
```

The agent can inspect these files with:

```
list_skill_files("sql-helper", "references")
load_skill_file("sql-helper", "references/warehouse_schema.md")
```

Use supporting files for content such as:

- Database schemas
- API examples
- Prompt templates
- Style guides
- Domain glossaries
- Step-by-step playbooks
- Test cases or examples

Creating an Executable Skill

A skill can optionally expose reusable executable behavior through `run_skill_entrypoint`. This is intended for controlled operations such as calculations, transformations, validation, or fetching structured data.

Executable skills must meet two requirements:

1. The skill must include `run_skill_entrypoint` in `allowed-tools`.
2. The script must be explicitly declared in the `entrypoints` section of `SKILL.md`.

Example Executable Skill

```
skills/
  statistics-helper/
```

```
SKILL.md
scripts/
    summarize_numbers.py
```

SKILL.md

```
---
name: statistics-helper
description: Computes basic summary statistics for numeric data.
allowed-tools: "load_skill_file list_skill_files run_skill_entrypoint"
entrypoints:
  - name: summarize_numbers
    script: scripts/summarize_numbers.py
    func: run
    description: Returns count, min, max, mean, and median for a list of
numbers.
---
```

Statistics Helper

Use this skill when the user asks for basic descriptive statistics.

```
scripts/summarize_numbers.py:
from statistics import mean, median
```

```
def run(*, values: list[float]) -> dict:
    if not values:
        raise ValueError("values must not be empty")

    return {
        "count": len(values),
        "min": min(values),
        "max": max(values),
        "mean": mean(values),
        "median": median(values),
    }
```

Example invocation:

```
run_skill_entrypoint(
    name="statistics-helper",
    entrypoint="summarize_numbers",
    args_json="{\"values\": [10, 20, 30, 40]}",
    timeout_seconds=10
)
```

The runner returns structured output that includes `exit_code`, `stdout`, `stderr`, and a best-effort parsed result when the script prints or returns JSON.

Rules for Executable Entrypoints

Executable entrypoints are intentionally constrained. The platform only runs Python files that are:

- Located under the skill's `scripts/` directory
- Declared in the skill's entrypoints frontmatter
- Allowed by the skill's `allowed-tools` setting

The platform does not provide general-purpose arbitrary script execution. Scripts that are not declared in SKILL.md are not runnable.

The script runner uses a timeout, defaults to 10 seconds, runs Python with isolated-mode behavior, and applies path restrictions. However, subprocess-based execution is not a full operating system sandbox. For production use, higher isolation such as containers, restricted filesystems, or network controls should be considered.

Tool Permissions with allowed-tools

allowed-tools acts as a skill-level permission gate. For a documentation-only skill, you may allow only file-reading tools:

```
allowed-tools: "load_skill_file list_skill_files"
```

For a skill that can execute declared scripts, include run_skill_entrypoint:

```
allowed-tools: "load_skill_file list_skill_files run_skill_entrypoint"
```

Do not add run_skill_entrypoint unless the skill genuinely needs executable behavior.

How to Let Your Agents Discover and Use Skills

To supplement your agent with skills, you must instantiate a skill catalog, a skill middleware and convert skills into tools using the following objects from the aidpUtils library:

Tool	Purpose
discover_skill_catalog	Determine default skill search locations (project + user) Build a SkillCatalog from discovered directories
SkillMiddleware	Append available skills summary and routing rules to system prompt. Provide factory helpers for workspace-driven middleware construction.
make_skill_tools	This method returns the skill discovery tools – activate_skill, list_skill_files, load_skill_file, and run_skill_entrypoint. These tools can be used by the agent to activate and run different skills.

Here's an example of what your entry file would include:

```
from aidputils.agents.skills.discovery import discover_skill_catalog
from aidputils.agents.skills.middleware import SkillMiddleware
from aidputils.agents.skills.tools.factories import make_skill_tools
...
class SchoolGradeAgentWithEmbeddedSkills:
    ...
    def init(self) -> None:
        ...
        self.catalog = discover_skill_catalog(skill_folder_whitelist=None)
        self.skill_middleware = SkillMiddleware(self.catalog)
        self.tools = make_skill_tools(self.catalog)
```

You can debug your skills catalog by adding this logger statement to your code. This will print every skill discovered in the skills catalog:

```
for info in self.catalog.list():
    logger.info("skill_id=%s name=%s desc=%s root=%s skill_file=%s",
info.skill_id, info.name, info.description, info.root_dir, info.skill_file)
```

Skill Precedence

The platform can load skills from multiple locations, such as project-level and user-level directories. The catalog aggregates those locations into a single name-keyed list of skills.

When multiple stores contain a skill with the same name, precedence determines which one is used. Later stores override earlier ones, which allows a host application to control whether user-level skills, project-level skills, or workspace-level skills take priority.

Skill Authoring Best Practices

Keep SKILL.md focused

Use SKILL.md for the core instructions the agent needs immediately after activation. Put long schemas, examples, and reference material in references/.

Write clear descriptions

The description field is used for discovery. Make it specific enough for the agent to know when to activate the skill.

Good:

```
description: Helps generate BigQuery SQL using the finance warehouse schema.
```

Less useful:

```
description: Helps with data.
```

Use explicit endpoint names

Endpoint names should describe the operation clearly:

```
entrypoints:
  - name: validate_query
  - name: summarize_numbers
  - name: transform_csv
```

Avoid vague names such as:

```
entrypoints:
  - name: run
  - name: do_it
```

Return structured results

Executable scripts should return JSON-serializable results whenever possible. This makes the output easier for the agent to inspect and use.

Avoid unnecessary execution

Prefer instructions and reference files when possible. Use executable entrypoints only for operations that genuinely require code.

Add a New Skill

You can add new Agent skills by creating a new folder inside the skills directory and adding the necessary files and folders.

1. Create a folder under the skills directory: `.agents/skills/<skill-name>/`.
2. Add a `SKILL.md` file with required frontmatter.

```
---
name: <skill-name>
description: <what this skill helps the agent do>
---
```

3. Write the skill instructions in Markdown below the frontmatter.
4. Add optional supporting files under:

```
references/
assets/
scripts/
```

5. If the skill is executable, add `run_skill_entrypoint` to `allowed-tools`, declare the entrypoints in `SKILL.md`, and place the Python implementation under `scripts/`.

Add a New Executable Capability to an Existing Skill

You can add a new executable operation to an existing skill to expand the capabilities of `SKILL.md`.

1. Add a Python file under the skill's `scripts/` directory.
`.agents/skills/<skill-name>/scripts/my_operation.py`
2. Implement a `run(...)` function.

```
def run(*, input_text: str) -> dict:
    return {
        "length": len(input_text),
        "uppercase": input_text.upper(),
    }
```

3. Add a matching entrypoint to `SKILL.md`.

```
allowed-tools: "load_skill_file list_skill_files run_skill_entrypoint"
entrypoints:
  - name: my_operation
    script: scripts/my_operation.py
    func: run
    description: Processes input text and returns structured output.
```

4. Test the endpoint with a JSON object as arguments.

```
{  
  "input_text": "hello"  
}
```

Agent Skills Troubleshooting

If you encounter issues with implementing Agent Skills, check this list for help resolving your problem.

The agent does not see my skill

Check that:

- The skill folder is located under a configured skills directory.
- The folder contains SKILL.md.
- SKILL.md has valid YAML frontmatter.
- The frontmatter includes both name and description.

The agent activates the wrong skill

Check for duplicate skill names across skill directories. If two skills have the same name, catalog precedence determines which one is used.

A supporting file cannot be loaded

Check that:

- The file is inside the skill folder.
- The path does not include traversals such as ../.
- The file is not hidden.
- The file is not excluded, such as __pycache__ or .pyc.

An endpoint will not run

Check that:

- run_skill_endpoint is included in allowed-tools.
- The endpoint is declared in SKILL.md.
- The script path is under scripts/.
- The script is a .py file.
- The function name in func exists in the script.
- The arguments are a valid JSON object.

An endpoint times out

Increase timeout_seconds only if the operation is expected to take longer. For long-running or resource-intensive operations, consider moving the operation to a dedicated service or more isolated execution environment.

Example: Complete Agent Skill

This example demonstrates what a complete agent skill would look like after implementation.

Folder Structure

```
skills/  
  customer-support-reply/  
    SKILL.md  
    references/  
      tone_guide.md  
      refund_policy.md  
      escalation_rules.md
```

SKILL.md

```
---  
name: customer-support-reply  
description: Helps draft customer support replies using the company tone  
guide and policy references.  
allowed-tools: "load_skill_file list_skill_files"  
metadata:  
  owner: support-operations  
  domain: customer-support  
---  
  
# Customer Support Reply  
  
Use this skill when the user asks for help drafting, reviewing, or improving  
a customer support response.  
  
Workflow:  
  
1. Identify the customer's issue.  
2. Load the relevant policy file from `references/` if needed.  
3. Draft a clear, empathetic response.  
4. Avoid making commitments that are not supported by policy.  
5. Recommend escalation when the request matches the escalation rules.  
This skill does not run code. It gives the agent structured instructions and  
optional policy files that can be loaded only when relevant.
```

Agent Testing

You can test your agents to preview and debug their output. You can also create and manage testing sessions to explore different testing scenarios for your agents.

The first step to test an agent is to attach your agent to an AI compute. The action of attaching an agent pushes a copy of your agent to an AI compute. As long as your agent is attached to an AI compute, any changes you made to your agent are propagated to the attached compute every time you click on the Test button.

The **Run now** command and Playground access your files differently.

- **Run now** runs the agent in your workspace, meaning it automatically has access to paths like `/Workspace/...`
- Deployed agents and agents in the Playground run a packaged copy of the agent on the attached AI compute. The package holds the files and folders of the entry file and subfolders. Parent and sibling workspace folders are not included.

A file opened with an absolute workspace path works with **Run now**, but fails in Playground and deployment scenarios. To avoid this failure and make the agent work in both **Run now** and Playground scenarios:

- **Put the file inside the agent folder:** By putting the file in the file folder or subfolder of the entry, the file is included in the package and is included when the agent is deployed or tested in the Playground.
- **Open the file with a relative path:** A relative path ensures that the file can still be located when the package is extracted to a different location.

```
import json
import os

CONFIG_PATH = os.path.join(
    os.path.dirname(os.path.abspath(__file__)), "config", "settings.json"
)

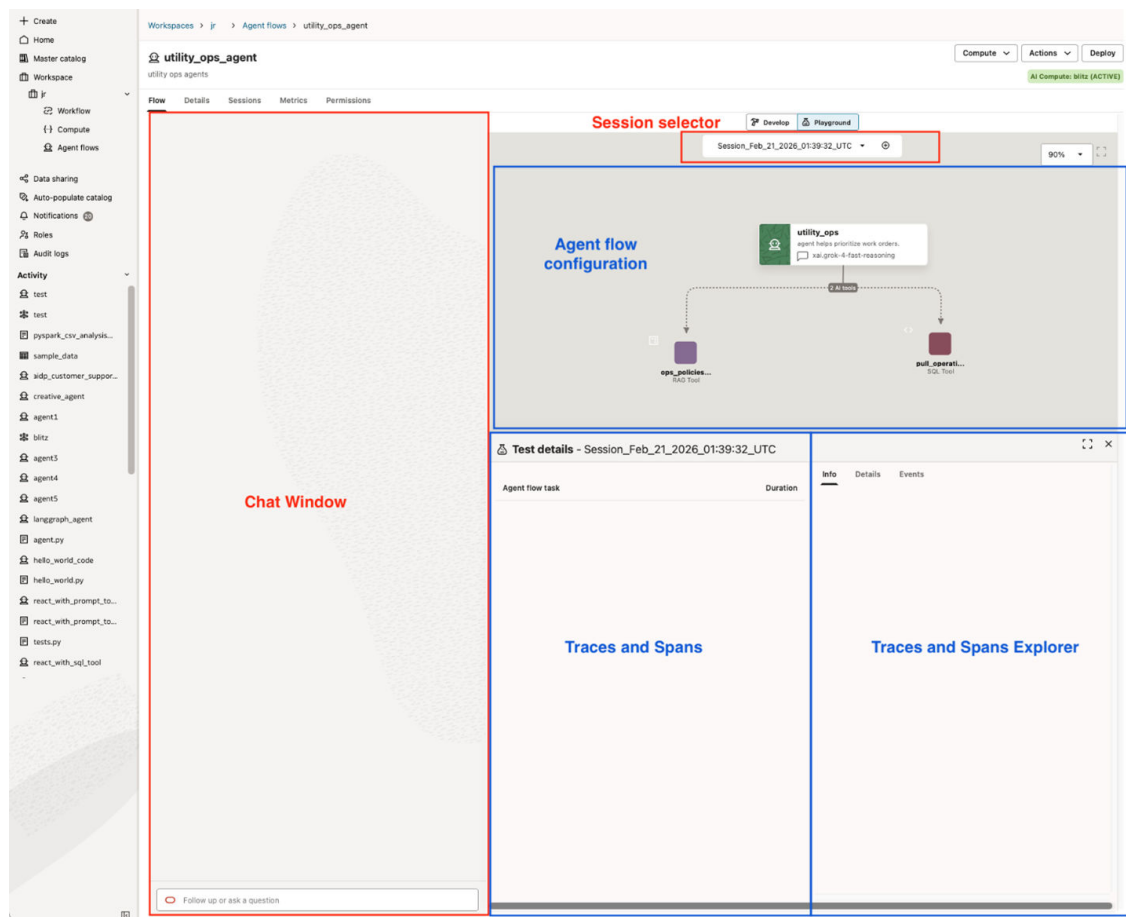
with open(CONFIG_PATH) as f:
    config = json.load(f)
```

- **Use the Skills folder:** For files shared across agents, you can use the skills folder `.agents/skills/` to load skills tools rather than including them by path.
- **Use agent parameters or session variables:** For values that change often, we recommend using agent parameters or session variables for each deployments. You can also use a catalog volume for larger files.

Note

You can use the same methods to ensure files included in a Custom tool are deployed successfully to other locations.

After you click the Test button, you are taken to the test playground.



The test playground has the following components:

- A chat window where you can initiate a session and start chatting with the agent, or resume an existing session
- A graph-based representation of the agent
- A panel showing a tree of traces and spans generated during the session
- A traces and spans explorer panel that displays traces and spans attributes, input/output. The Details tab includes IDs, start and end time, execution time, while the Events tabs highlight any errors during the execution.

The Playground lets you interact and test each agent independently if you wish to do so. By default, the supervisor agent is selected, but you can choose to chat with and test each executor agent independently. This allows you to simulate the behavior of a supervisor agent issuing requests to executor agents. To do this, you select the agent you want to test in the drop-down menu in the chat window.

Traces and spans are displayed in the central panel as soon as you create your first message. Each task corresponds to a different user message. You can click on the left caret to expand the trace and inspect the spans.

Test your Agents in the Playground

You can test visual builder and LangGraph-based agents from the Test playground to validate and debug your agents.

You must have an AI compute attached to your agent to test. You can add a new AI compute cluster by following [Create an AI Cluster for an Agent](#) or attach an existing AI compute cluster by following [Attach an Existing AI Cluster to an Agent](#).

1. Navigate to your agent in your workspace. Click the agent name.
2. At the top of the canvas, click **Playground**. It may take several seconds to push the agent to your attached compute.



Your agent is displayed in the test playground.

Create an Agent Test Session

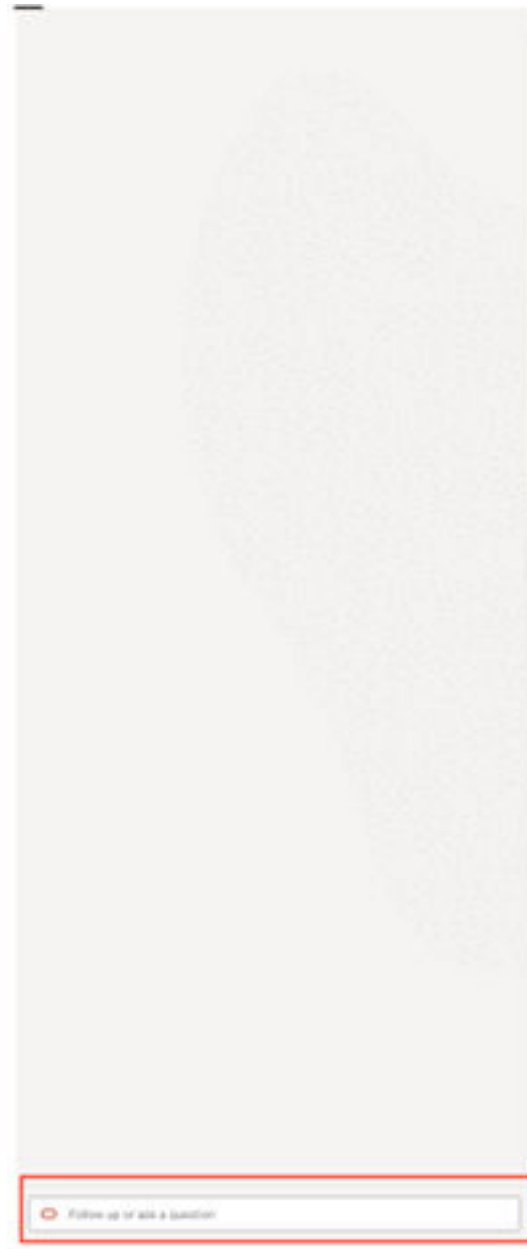
You can create a test session to initiate a new conversation with your agent.

All sessions created in the test playground target the agent and are hosted on the attached compute. Once a session is created, it can be resumed later.

1. Navigate to your agent in your workspace. Click the agent name.
2. At the top of the canvas, click **Playground**
3. In the session selector, click **Create a session**.



4. Start a dialog with your agent by entering a query in the chat box.



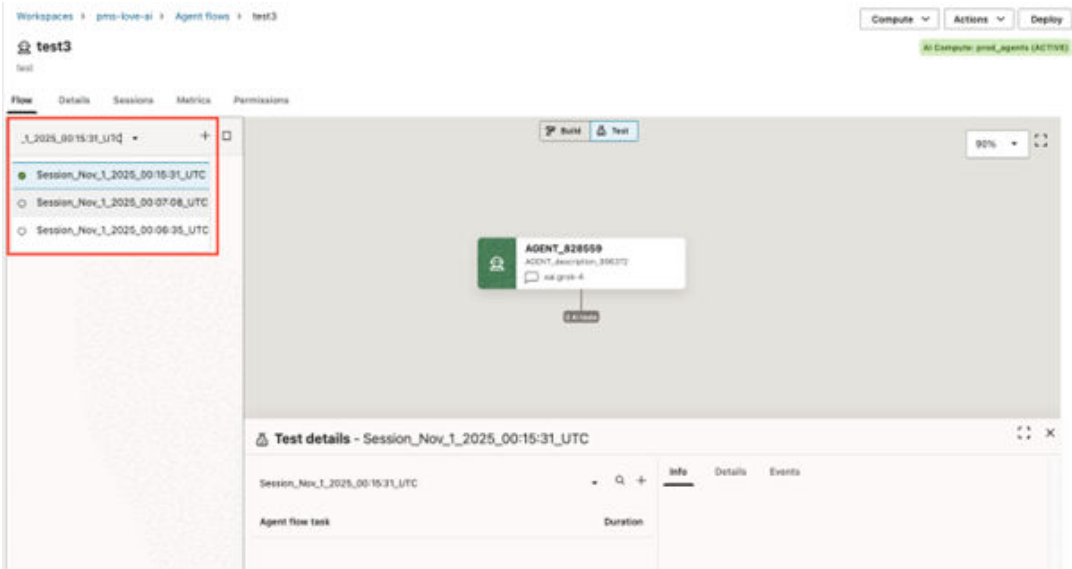
Resume an Agent Test Session

You can resume agent test sessions you have previously created.

Note

You can only resume sessions that you created.

1. Navigate to your agent in your workspace. Click the agent name.
2. At the top of the canvas, click **Playground**
3. From the session drop-down, select a previous session.

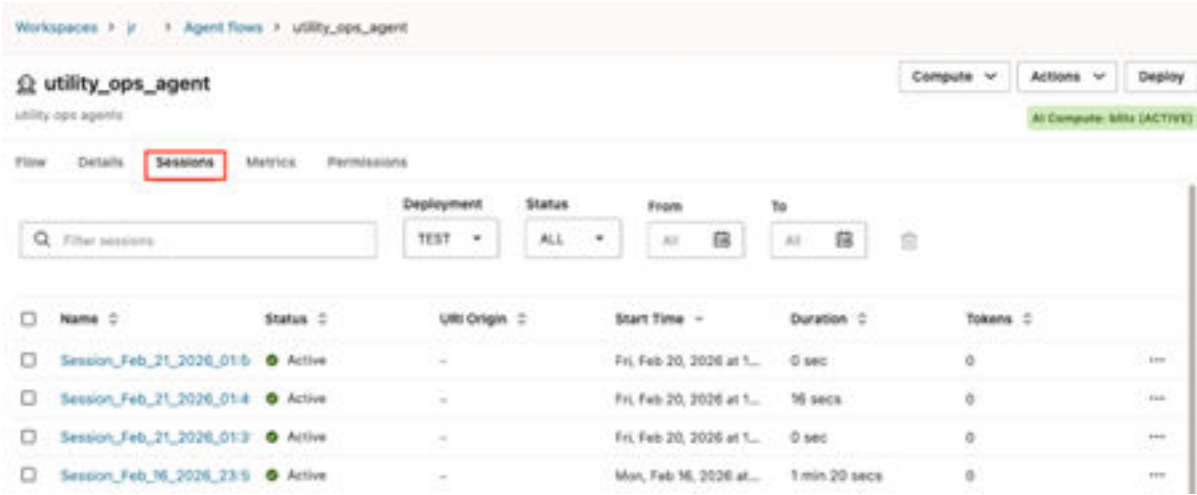


4. Resume your dialog with your agent by entering a query in the chat box.

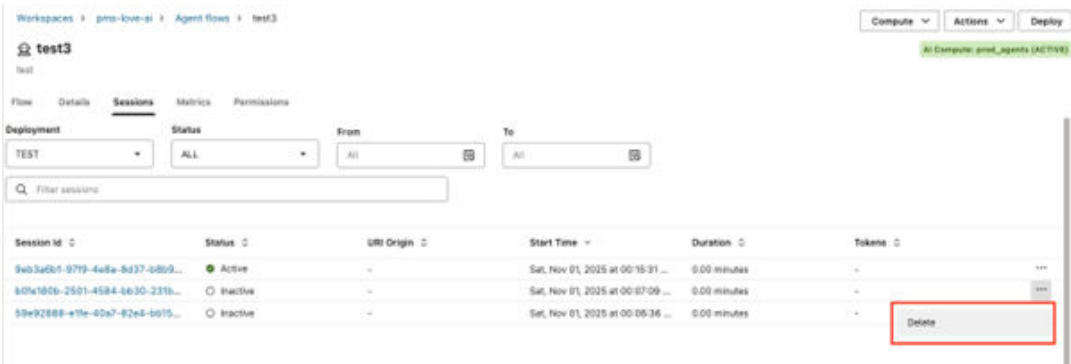
Delete an Agent Test Session

You can delete test sessions for agents hosted on the attached AI compute and sessions that have been created on a deployed agent.

1. Navigate to your agent in your workspace.
2. Click the Sessions tab.



3. Next to the session you want to delete, click ... **Actions** then click **Delete**.



4. Click **Delete**.

About Agent Tools

Oracle AI Data Platform supports tool templates that can be configured to access your data and fit your use cases.

Agents support configurations that consist of a single agent that can interface with one or more tools. AI Data Platform offers three tool templates that can be configured for use through visual flows or code:

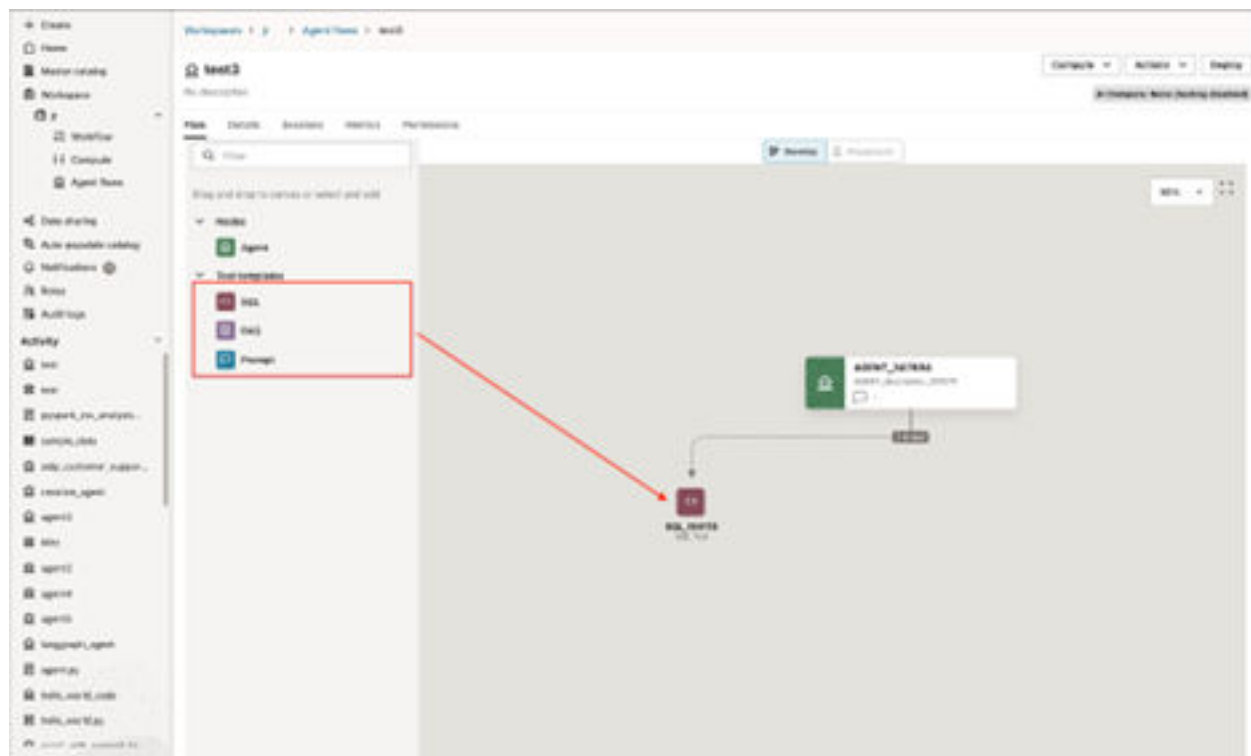
- **Custom Code:** The Custom Code tool allows AI developers to implement their tool using Python. Developers package their tool in a ZIP, upload it to their workspace, and configure it as a node in their agent. Custom Code tools are intended for cases where built in tools don't provide the integration they need.
- **HTTP Request:** The HTTP Request tools lets developers use supported REST API calls in their agents, leveraging the AI Data Platform APIs and the functions they provide. Agents can use REST APIs to create workspace objects, check details, pull lists, or modify existing objects. For a full list of available APIs, see [REST API for Oracle AI Data Platform](#).
- **Prompt:** The prompt tool allows the AI developer to define a parametrized prompt that can be issued to an LLM for their choice. Common use cases for a prompt tool include email drafting tasks, translation tasks, style conversion, git commit message, and code explanations.
- **RAG:** The RAG tool lets agents pull relevant external knowledge before generating a response. In AI Data Platform, the RAG tool queries a knowledge base (26ai Vector Search) and retrieves semantically relevant document chunks. Those chunks are then passed to the agent for response generation.
- **SQL:** The SQL tool enables agents to execute SQL queries against structured data sources registered via external catalogs, such as Oracle Autonomous AI Lakehouse, Oracle Autonomous AI Transaction Processing, or Oracle Autonomous AI Database. The tool is intended for scenarios where the SQL queries are predefined and can be parametrized. The objective is to let an agent assign values to the parameters. This tool is not an NL2SQL tool that generates a SQL query based on a natural language prompt.

Note

The SQL tool only performs queries against data in an external catalog. It does not support data stored in a standard catalog.

Agent Tools through Visual Flow

When you add tools to agents through visual flow, you can find tools under **Tool templates** in your agent. You add a tool to your agent by dragging and dropping it into the visual flow canvas. After dragging the tool node on the canvas, the node automatically connects with the agent.



Each tool can be configured in the Parameters tab and be tested independently of the agent by clicking on the Test tab.

Note

You must attach an AI Compute to your agent before you can test a system tool. If no compute is attached, the Test tab is disabled.

Agent Tools through LangGraph Code

You add tools to your LangGraph coded agents through an instance of the `AIDPToolConf()` class.

```
from aidputils.agents.toolkit.configs import AIDPToolConf
aid_tool = AIDPToolConf(name, description, tool_class, conf, params)
```

The parameters mirror the visual flow experience for tools. For each tool, you must provide:

- **Name:** A descriptive name to help users and the LLM understand the purpose of the tool.
- **Description:** A thorough summary that provides sufficient information for users and LLMs to understand what the tool does.
- **tool_class:** The supported tool type, `PromptTool`, `SQLTool`, `RAGTool`, `HTTPTool`, and `MCPTool`.
- **conf:** The tool configuration. This information is hidden from the LLM.
- **params:** The parameters exposed to the LLM.

Custom Tool

The Custom Code tool lets agent developers extend Oracle AI Data Platform with their own Python code.

You package your tool implementation as a ZIP file, upload it to your workspace, and configure it as a Custom Code tool node in the agent. The agent calls your code as a tool, with parameters supplied by the LLM at runtime.

The Custom Code tool is intended for cases where the built-in tools (HTTP, SQL, RAG, MCP) do not cover the integration you need — for example, when you need to perform local computation, parse a domain-specific format, or compose multiple steps that should appear to the agent as a single tool call.

AI Data Platform has the following limits when uploading a ZIP file with Python code for your custom code tool:

Constraint	Limit
Maximum ZIP size	10 MB
Maximum file size inside the ZIP	10 MB per file
Maximum total uncompressed size	500 MB
Path traversal	Blocked (../ rejected)

 Note

Custom Code tools run on the AI compute attached to your agent. The code has access to the compute environment and outbound network access subject to the workspace networking configuration. Only upload code from sources you trust.

Custom Code Tool Parameters

On the Parameters tab, you configure the static settings for each tool class in the package. The Tool Class dropdown lets you switch between the tools discovered in the package.

The custom code tool Parameters tab includes the following sections:

- **Tool class:** Select the tool class to configure. The drop-down is populated from the classes registered in `tool_implementation.py`.
- **Description:** A clear, concise description of what the tool does. The description is provided to the agent and helps the LLM decide when to call the tool. The default description is read from `tool_config.json` and can be overridden here.
- **Configuration:** The static settings the tool needs at runtime. These are the keys defined in the `conf` object of `tool_config.json`. Examples include `timeout`, `base_dir`, `max_output_lines`, and credential references. Configuration values support `{{variable}}` runtime parameter references. Session variables are not currently substituted into custom-tool configuration; if you need a session value, pass it as a runtime parameter from the agent.
- **AI tool definition:** The schema exposed to the agent, including the tool name, description, and the runtime parameters the agent can pass. The schema is rendered automatically from the `schema` array in `tool_config.json`.

Custom Code Tool Authoring

A Custom Code tool package is a ZIP file with the following structure:

```
my_tool.zip
├── tool_implementation.py    # Required. Contains the tool class(es).
├── tool_config.json         # Required. Tool metadata and schema.
├── requirements.txt         # Optional. Python dependencies.
├── utils/                  # Optional. Helper modules.
│   ├── __init__.py
│   └── helpers.py
├── config/                 # Optional. Static configuration files.
│   └── settings.yaml
├── wheels/                 # Optional. Bundled wheel files for offline
└── install.
    └── humanize-4.15.0-py3-none-any.whl
```

tool_implementation.py

Each tool class extends CustomToolBase and is decorated with @BaseTool.register. The class must implement the _execute_tool class method, which receives the tool configuration, the runtime parameters from the agent, and the system context variables, and returns a value, like dict, str, or list.

The following is a blank example template of a tool_implementation.py:

```

"""Custom Code tool implementation."""
from aidputils.agents.tools.custom_tools.base import CustomToolBase

@BaseTool.register
class MyTool(CustomToolBase):
    """Brief description of what the tool does."""

    @classmethod
    def _validate_config(cls, conf, runtime_params, **context_vars):
        """Optional. Validate configuration before execution.
        Raise ValueError to abort the call.
        """
        # Example: require an api_key in the tool configuration
        if not conf.get("conf", {}).get("api_key"):
            raise ValueError("api_key is required")

    @classmethod
    def _execute_tool(cls, conf, runtime_params, **context_vars):
        """Required. Implement the tool logic.

        Args:
            conf: the AIDPToolConf dict. User configuration values
                  live under conf["conf"] when the tool is invoked from
                  a deployed agent. During a Test run the tool may
                  receive a flat conf dict; the Developer Toolkit example
                  below uses a small _get_cfg helper that tolerates both
                  shapes.
            runtime_params: the runtime parameters passed by the
                           agent at invocation time.
            context_vars: system context (such as datalake_id).

        Returns:
            Any value (dict, str, list, ...). It will be wrapped into
            the MCP response by the framework.

            To signal a failure, raise an exception:
            - ValueError -> INVALID_CONFIG
            - any other exception -> TOOL_EXECUTION_ERROR
            Do NOT return {"error": "..."}; the framework wraps a
            successful return in {"response": ..., "success": True},
            so a returned error dict is treated as a normal payload
            and the agent will not see it as a failure.
        """
        tool_conf = conf.get("conf", conf)
        param_value = runtime_params.get("my_param", "")
        # Tool logic here

```

```

        return {"output": f"Processed: {param_value}"}

    @classmethod
    def _transform_response(cls, response):
        """Optional. Transform the response before MCP formatting."""
        return response

```

tool_config.json

The `tool_config.json` file describes the tools in the package — their display name, description, version, runtime parameter schema, and default configuration values. Each tool registered in `tool_implementation.py` must have a corresponding entry in the tools array.

The following is a blank example template of a `tool_config.json`:

```

{
  "displayName": "My Tool Package",
  "description": "Brief description of the tool package.",
  "tools": [
    {
      "toolClassName": "MyTool",
      "displayName": "My Tool",
      "description": "Clear description of when the agent should call this
tool.",
      "version": "1.0.0",
      "schema": [
        {
          "name": "my_param",
          "type": "string",
          "description": "What this parameter is for."
        }
      ],
      "conf": {
        "timeout": 30
      }
    }
  ]
}

```

Schema Field Types

The Parameters tab in the visual builder accepts string, number, and boolean. The runtime accepts a wider set when authoring `tool_config.json` by hand: int, integer, float, double, number, numeric, bytes, list, array, sequence, dict, map, mapping, set, tuple, none, null, plus generic forms like `list[int]`. These wider types are usable from JSON but are not exposed in the UI dropdown.

requirements.txt

The `requirements.txt` file lists the Python dependencies your tool needs. Standard pip syntax is supported, including version specifiers and comments. The file is optional — if your tool only uses the Python standard library or pre-installed packages, you do not need a `requirements.txt`.

The following is a blank example of `requirements.txt`:

```
# List third-party dependencies one per line.
# Examples:
# humanize>=4.0
# python-dateutil>=2.8,<3.0
# beautifulsoup4==4.12.3
```

AI Data Platform filters the dependencies in `requirements.txt` before installing them on the AI compute, to prevent runtime conflicts with the platform itself. The filtering rules are as follows:

Category	Example	Action
Platform packages	langgraph, langchain-core, langchain-oci, langchain_mcp_adapters, pyyaml	Discarded (would break the agent runtime).
Pre-installed packages	oci, requests, requests-toolbelt, websockets, cryptography, certifi, pyopenssl, urllib3, pydantic, pydantic-core, pydantic-settings, numpy, oracledb, sqlalchemy, aiohttp, httpx, httpx-sse, anyio, jsonschema, orjson	Skipped (already available, no need to declare).
URL or VCS installs	git+https://..., -e ./local_pkg	Blocked (security).
Everything else	humanize, beautifulsoup4, jmespath	Installed.

① Note

Dependencies declared in `requirements.txt` are installed during the full deployment of the agent. Dependencies are not installed during a single test run from the configuration panel. If your tool depends on third-party packages, deploy the agent first and then exercise the tool from the Playground.

For tools that need dependencies which are not pre-installed and where deterministic, offline installation is important, you can bundle `.whl` files inside a `wheels/` directory at the root of the ZIP. The platform installs from the local wheels directory first and falls back to the package index only if needed. This is the recommended approach for production tools.

Bundling wheels for offline install

```
pip download \
  --dest wheels/ \
  --platform manylinux_2_28_x86_64 \
  --python-version 3.11 \
  --only-binary=:all: \
  -r requirements.txt
```

Tool Lifecycle Hooks

Custom Code tools support three lifecycle methods. Only `_execute_tool` is required.

Method	When called	Purpose
<code>_validate_config</code>	Before <code>_execute_tool</code>	Validates the configuration. Raise <code>ValueError</code> to abort the call before it runs.
<code>_execute_tool</code>	On every tool invocation	Required. Implements the tool's behavior. Returns any value (dict, str, list) and raises an exception to signal a failure (<code>ValueError</code> → <code>INVALID_CONFIG</code> , any other exception → <code>TOOL_EXECUTION_ERROR</code>). Do not use a returned <code>{"error": "..."} dict</code> as it is treated as a normal payload.
<code>_transform_response</code>	After <code>_execute_tool</code>	Transform the response before it is wrapped in the MCP format and returned to the agent.
<code>prompt_template</code>	string	Prompt template used by the LLM, with variables in <code>{{variable}}</code> format for dynamic insertion

Configuration values versus runtime parameters

Custom Code tools have two distinct sources of input that are easy to confuse. Configuration values come from the Configuration section of the Parameters tab and are baked into the tool when the agent is deployed. Runtime parameters come from the agent at invocation time and are different on every call.

- Configuration values are accessed via `conf.get("conf", conf)`. Use them for things that do not change between calls — base URLs, credential references, timeouts, output limits.
- Runtime parameters are accessed via `runtime_params.get("name")`. Use them for the values the agent actually decides at call time — the query, the file path, the request body.

Note

Configuration values can pass through template substitution and may arrive as strings even when you defined them as numbers. Always coerce numeric configuration values defensively, for example: `int(tool_conf.get("timeout", 30))`.

Multiple Tools Per Package

A single ZIP can contain multiple tool classes. Each class registered with `@CustomToolBase.register` becomes a separate tool in the agent. The Tools panel on the Package tab lists all discovered tools and lets you enable each one independently. Each tool is configured separately on the Parameters tab via the Tool Class dropdown.

Code Tool through LangGraph Code

From the code builder, a Custom Code tool is registered through the `aidpUtils` Python library by referencing the uploaded package and selecting one of its tool classes.

```
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.configs import AIDPToolConf
```

```
hello_tool_conf = AIDPToolConf(  
    name="hello_tool",  
    description="Returns a hello world greeting.",  
    tool_class="HelloTool", # the class registered with @BaseTool.register  
    conf={},                # values from tool_config.json "conf";  
    supports {{variable}} substitution  
    params=[  
        {"name": "name", "type": "string",  
         "description": "Name to greet."}  
    ],  
)  
  
hello_tool = create_langgraph_tool(hello_tool_conf.model_dump())
```

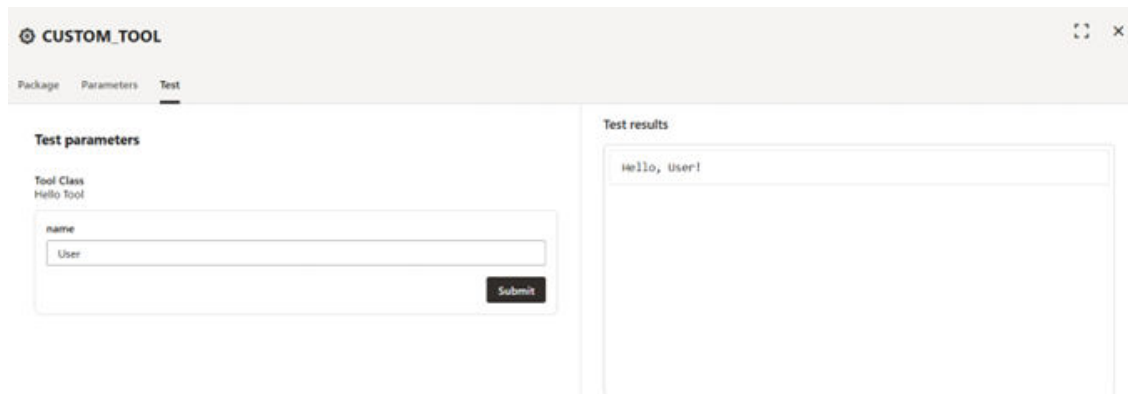
tool_class must be the exact class name registered via `@BaseTool.register` in `tool_implementation.py`. The framework looks the class up in `BaseTool.tool_class_registry[tool_class]`. `conf` mirrors the `conf` object of the matching entry in `tool_config.json`.

Note

Do not place `package_path` or `tool_class_name` inside `conf` as they are not consumed.

Test Agent Custom Code Tools

The Test tab lets you execute the tool without running the full agent. Provide values for any runtime parameters and any session variables referenced in the configuration, then click Run to invoke the tool and view the response.



The screenshot shows a web interface titled "CUSTOM_TOOL" with three tabs: "Package", "Parameters", and "Test". The "Test" tab is active. Under "Test parameters", the "Tool Class" is listed as "Hello Tool". There is a text input field for "name" with the value "User" and a "Submit" button. To the right, the "Test results" section shows the output "Hello, user!" in a text box.

Note

If your tool depends on third-party packages declared in `requirements.txt`, the dependencies are installed during the full deployment of the agent, not during a single test run. To test code that depends on additional packages, deploy the agent first and then invoke the tool from the Playground.

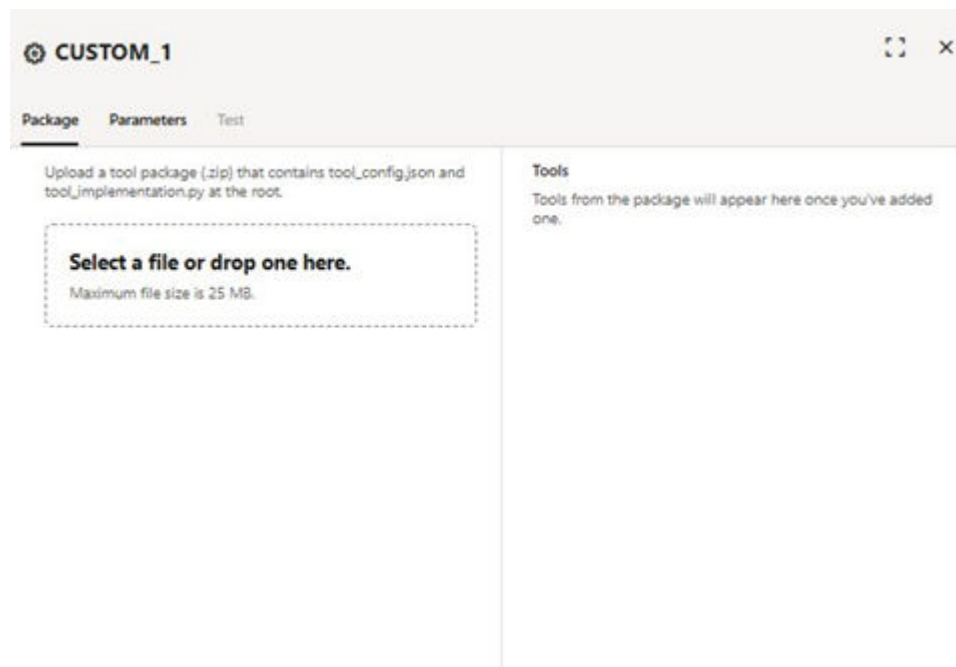
Add a Custom Tool to an Agent

You can add a custom tool to your agents to allow you to use your own Python code to extend Oracle AI Data Platform.

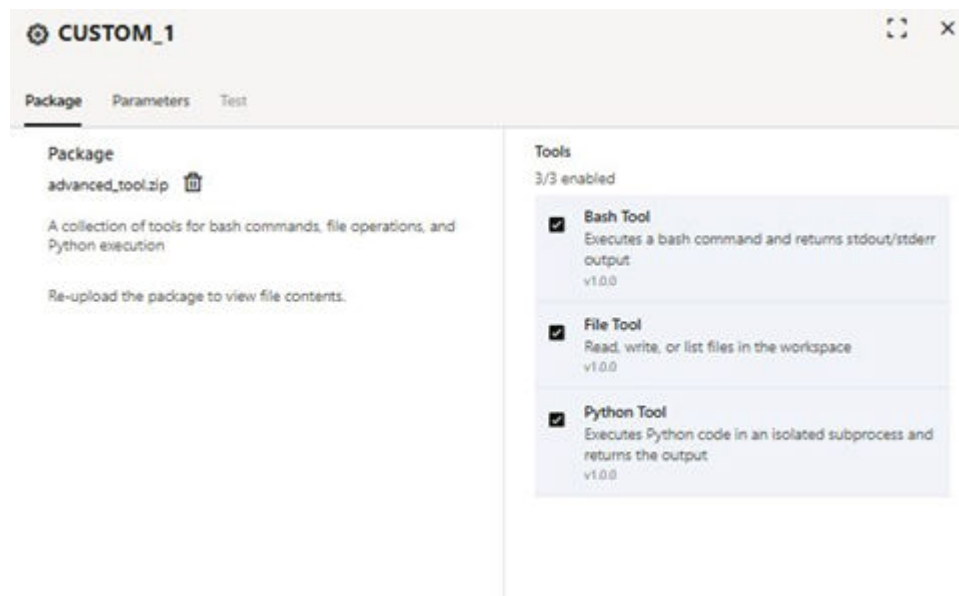
Note

An AI compute must be attached to your agent prior to adding a custom code tool. AI compute is required to install dependencies and run the tool.

1. Navigate to your agent.
2. From Tool templates, drag and drop a Custom tool to your canvas.
3. In the Package tab, click to select the ZIP file with your custom code or drag and drop it on to the screen. Wait for the upload to complete.



4. Review the list of discovered tools in the Tools section of the Package tab. Each tool class found in your `tool_implementation.py` is listed with its class name, description, and version.



5. Select the tools to enable. Disabled tools are not exposed to the agent.
6. Optional: Click the Test tab. Provide test parameters and click **Submit**. See test results in the Test results pane.

Remote MCP Server Tool

Agent developers can connect their agent flows to remote model context protocol (MCP) servers using the Remote MCP Server tool.



[Video](#)

The MCP tool is available in both the visual builder as well as in the code builder experiences. In the code builder experience, the MCP connection can be configured through the `aidpUtils` Python library. In this section, we walk you through both the visual builder and code builder experiences.

Note

This feature supports MCP servers with HTTP-streamable transports (remote servers). Local, stdio-transport MCP servers are not supported.

MCP Credentials in Oracle AI Data Platform Credential Store

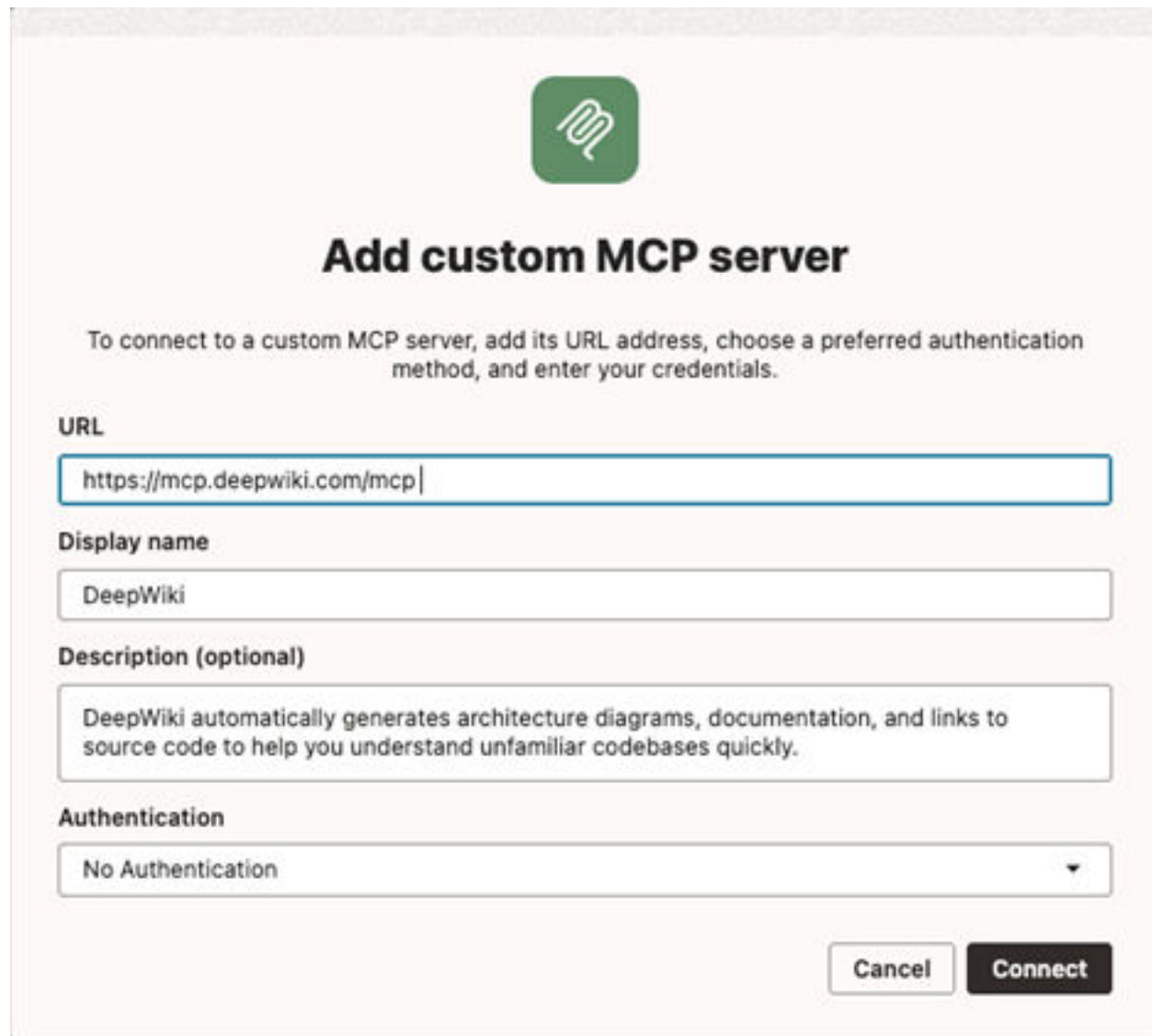
When configuring your MCP server, you need to select whether the remote MCP server requires **No Authentication** or a **Bearer token**. If your MCP server requires an authentication token, that token needs to be added to your Credential Store before it can be referenced by the MCP server.

When creating an MCP server credential, you select the **Secret token** option for **Credential type**, then provide the identifier key, such as an API Key and the token value. For more information, see [Create Credentials \(Preview\)](#).

Note

A single credential can hold multiple keys.

Publicly available MCP servers do not require additional authentication. For example, connecting to `https://mcp.deepwiki.com/mcp` would look like this:



Add custom MCP server

To connect to a custom MCP server, add its URL address, choose a preferred authentication method, and enter your credentials.

URL

`https://mcp.deepwiki.com/mcp`

Display name

DeepWiki

Description (optional)

DeepWiki automatically generates architecture diagrams, documentation, and links to source code to help you understand unfamiliar codebases quickly.

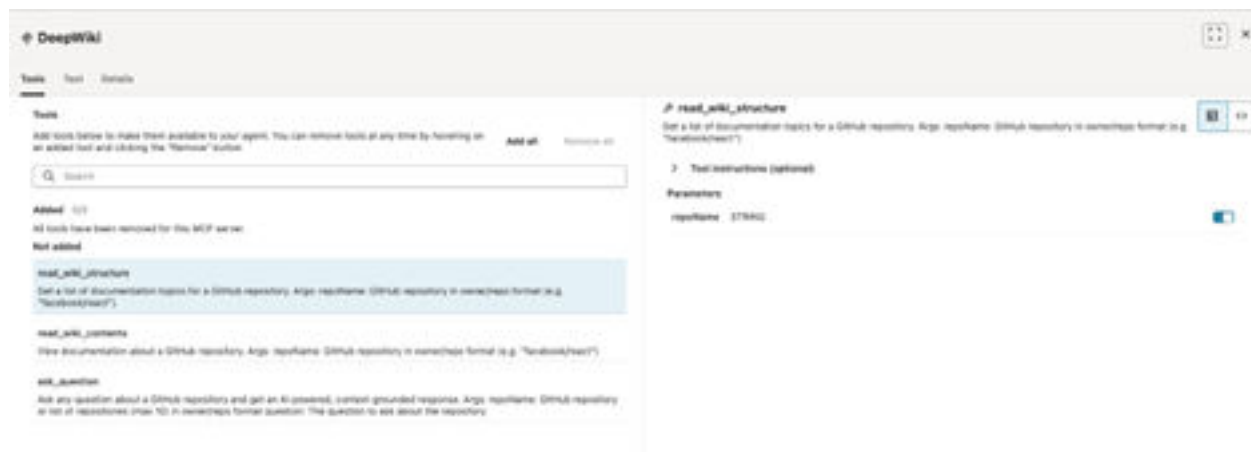
Authentication

No Authentication

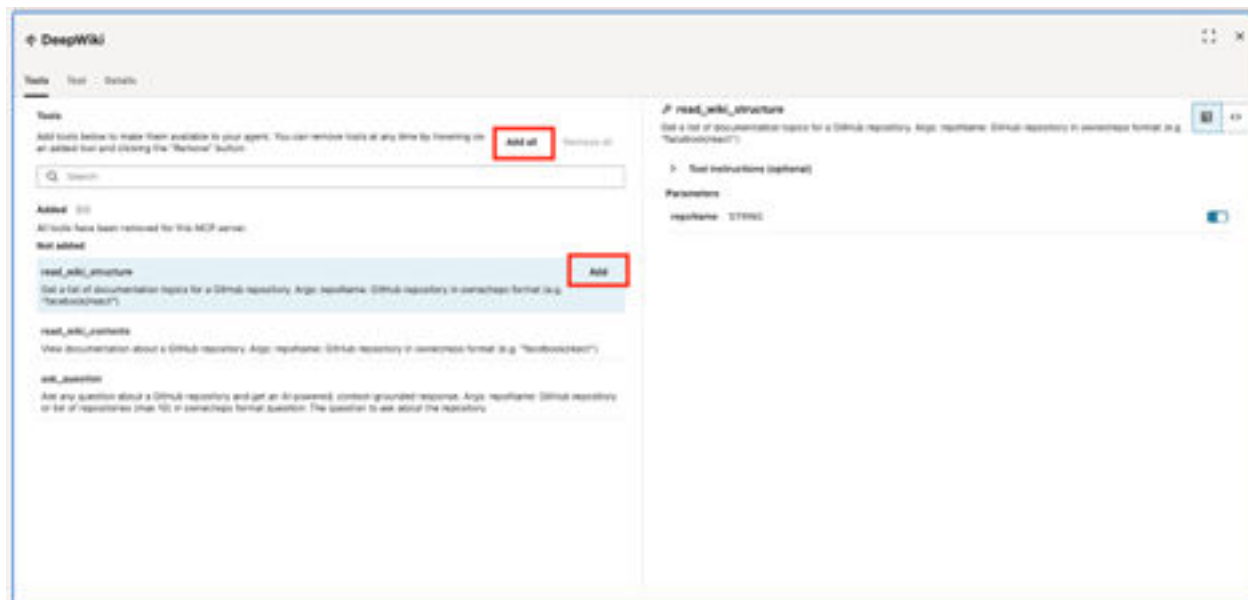
Cancel Connect

How to Expose MCP Tools to the Agent

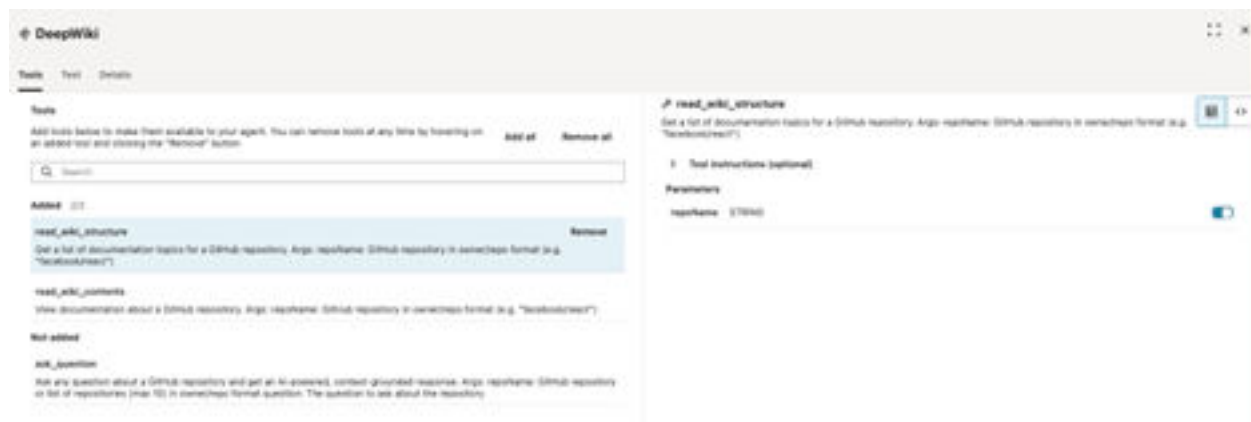
Once a successful connection to the remote MCP server has been established, you can start configuring which tools hosted on the server you want to expose to your agent. The MCP server configuration panel is shown below in the case of the DeepWiki MCP server.



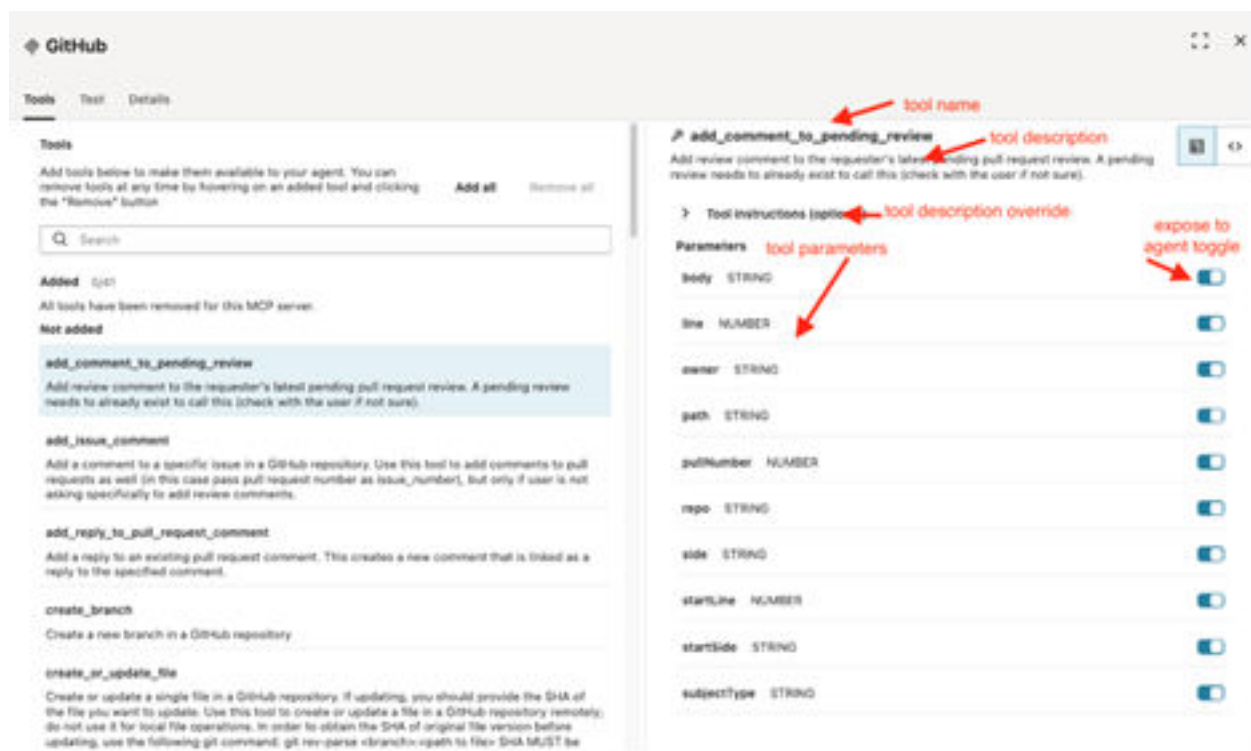
On the left, the Tools tab displays a list of tools available in the MCP server. You must add tools to expose them to your agent. You can do this by clicking either the **Add all** option to expose all the tools at once or by clicking on each tool **Add** option individually to select a subset of the tools.



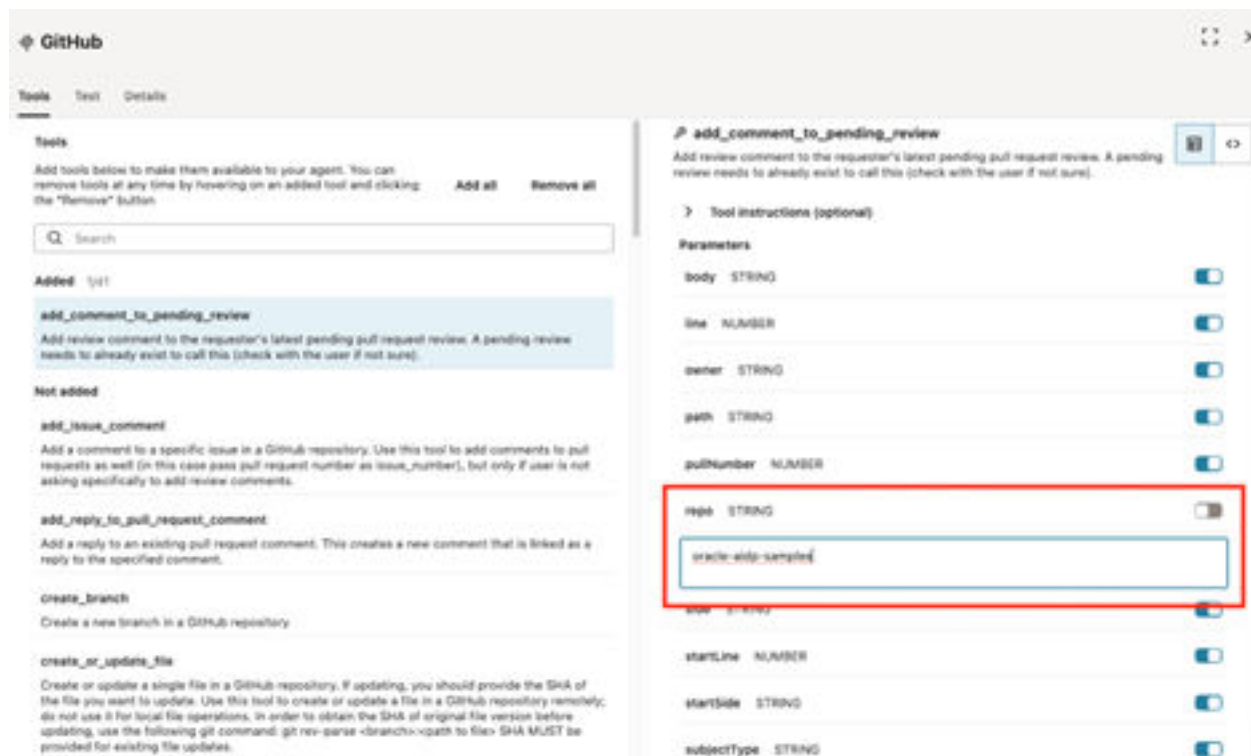
In the example below, we added two tools (`read_wiki_structure`, `read_wiki_structure`). You can remove tools by clicking on **Remove**.



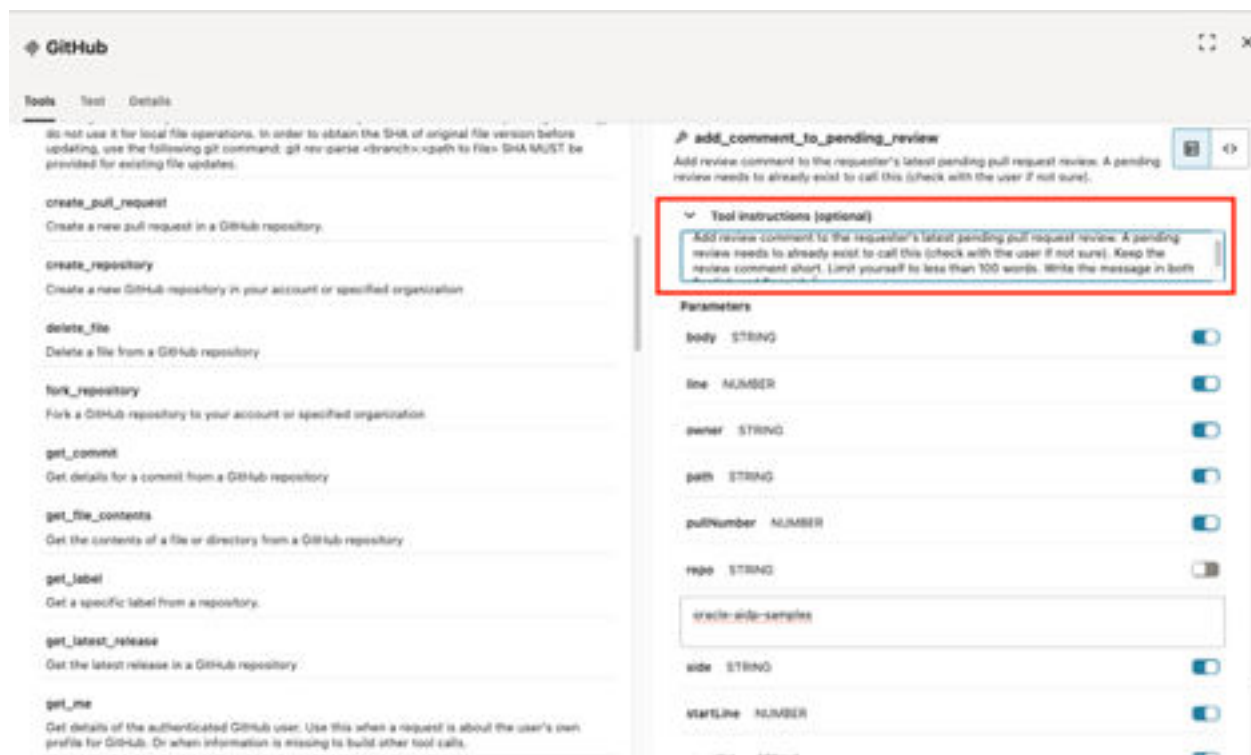
The right panel of the Tools tab provide documentation about each tool including the tool name, tool description as well as the tool parameters. In the screenshot below, I show an example for the GitHub MCP server tool `add_comment_to_pending_review`.



AI Data Platform provides a couple of additional controls over each tool. You can hide parameters from the agent and assign values to those parameters. For example, in GitHub you could choose for your agent to only comment on one pre-determined repo, such as `oracle-aidp-samples`. To achieve this, you disable the **repo** parameter and assign a default value in the text box:



In the **Tool Instructions** field you can also override the tool description and provide an alternative description with additional instructions. For most use-cases, we recommend that you adopt the description that is provided by the MCP server.



Remote MCP Server Tool through LangGraph Code

The `aidputils` Python library provides developers with the ability to select a remote MCP server and expose a subset of its tools to an agent built with LangGraph. For `aidputils` API reference, see [Aidp-utils API for Oracle AI Data Platform](#).

You can build a collection of allowed tools by creating an instance of `build_structured_tools_from_allowed_mcp_tools`:

```
from aidputils.agents.toolkit.tool_helper import
build_structured_tools_from_allowed_mcp_tools

TOOLS = build_structured_tools_from_allowed_mcp_tools(
    allowed_tools=<ALLOWED_MCP_TOOLS>,
    server_name=<MCP_SERVER_NAME>,
    endpoint=<MCP_ENDPOINT>,
    transport="streamable_http",
    auth=<MCP_AUTH>,
    headers={}
)
```

Where:

- `<MCP_SERVER_NAME>` is a display name you want to give to your MCP server. This is used for documentation purpose and is not exposed to the agent.
- `<MCP_ENDPOINT>` is the endpoint of the MCP server (e.g. `https://api.githubcopilot.com/mcp/`)
- `<MCP_AUTH>` is a dictionary with key "authType". This key can take two values: `NO_AUTH` or `BEARER_TOKEN`. In the case of `BEARER_TOKEN`, another key is expected: "token" with the value of the bearer token.
- `<ALLOWED_MCP_TOOLS>` is a list of the tools, from the MCP server, that you want to expose to your agent. Each tool needs a full JSON tool definition following the MCP protocol.

Here's an example:

```
MCP_SERVER_NAME = "test_mcp"
MCP_ENDPOINT = "http://144.25.36.217:9301/mcp"
MCP_AUTH = { "authType": "BEARER_TOKEN", "token": "valid-123" }

{
  "ALLOWED_TOOLS": [
    {
      "tool": {
        "name": "get_current_weather",
        "description": "Get current weather for a given city with advanced
options.",
        "inputSchema": {
          "type": "object",
          "properties": {
            "city": {
              "type": "string"
            }
          },
          "unit": {
```

```

        "type": "string",
        "default": "metric"
    },
    "include_historical": {
        "type": "boolean",
        "default": false
    },
    "detailed": {
        "type": "boolean",
        "default": true
    },
    "timeout": {
        "type": "integer",
        "default": 30
    }
},
"required": [
    "city"
]
},
"instruction": "",
"argOverrides": {}
},
{
    "tool": {
        "name": "get_forecast",
        "description": "Get forecast for a given city with customizable
options.",
        "inputSchema": {
            "type": "object",
            "properties": {
                "city": {
                    "type": "string"
                },
                "days": {
                    "type": "integer",
                    "default": 5
                },
                "unit": {
                    "type": "string",
                    "default": "metric"
                },
                "include_alerts": {
                    "type": "boolean",
                    "default": false
                },
                "detailed": {
                    "type": "boolean",
                    "default": true
                },
                "hourly": {
                    "type": "boolean",
                    "default": false
                }
            }
        }
    }
},

```

```

        "required": [
            "city"
        ]
    },
    "instruction": "",
    "argOverrides": {}
}
]
}

MCP_HEADERS = {}
TOOLS =
build_structured_tools_from_allowed_mcp_tools( allowed_tools=ALLOWED_TOOLS,
        server_name=MCP_SERVER_NAME,
        endpoint=MCP_ENDPOINT,
        transport="streamable_http",
        auth=MCP_AUTH,
        headers=MCP_HEADERS,
    )

```

The TOOLS object can be then used when creating an instance of an agent with `langchain.agent create_agent` in the `setup()` method of your class agent definition:

```

def setup(self):
    logger.info("Initializing TestMcpAgent")

    oci_llm = init_oci_llm(llm_conf)

    system_prompt = textwrap.dedent(
        """
        You're a weather agent. Append 12345 to every response.
        """
    ).strip()

    self.agent = create_agent(
        name="test_mcp_high_code",
        model=oci_llm,
        tools=TOOLS,
        system_prompt=system_prompt,
        debug=True,
    )

    logger.info("Agent ready.")

```

Alternatively, if you are using a session variable to store the value of a bearer token, a reference to a previously created session variable can be assigned to the token key of the auth config dictionary. For example:

```

test_mcp_auth_config = { "authType": "BEARER_TOKEN", "token" :
    "{{sessionvariables.cred.mcp.test_mcp.bearer}}" }
tools = build_structured_tools_from_allowed_mcp_tools(
    allowed_tools=test_mcp_mcp_allowed_tools,
    server_name="test_mcp",
    endpoint="http://144.25.36.217:9301/mcp",

```

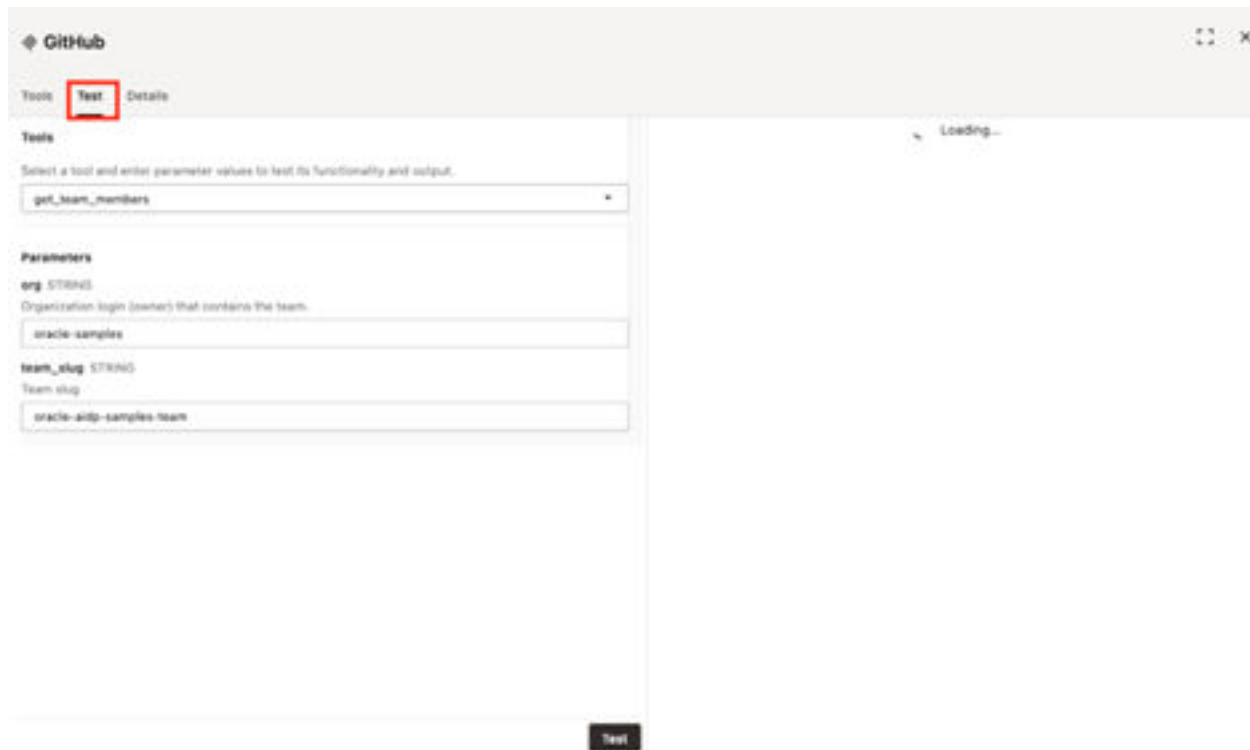
```
transport="streamable_http",  
auth=test_mcp_mcp_auth_config,  
headers={}  
)
```

Code Examples for Remote MCP Server Tools

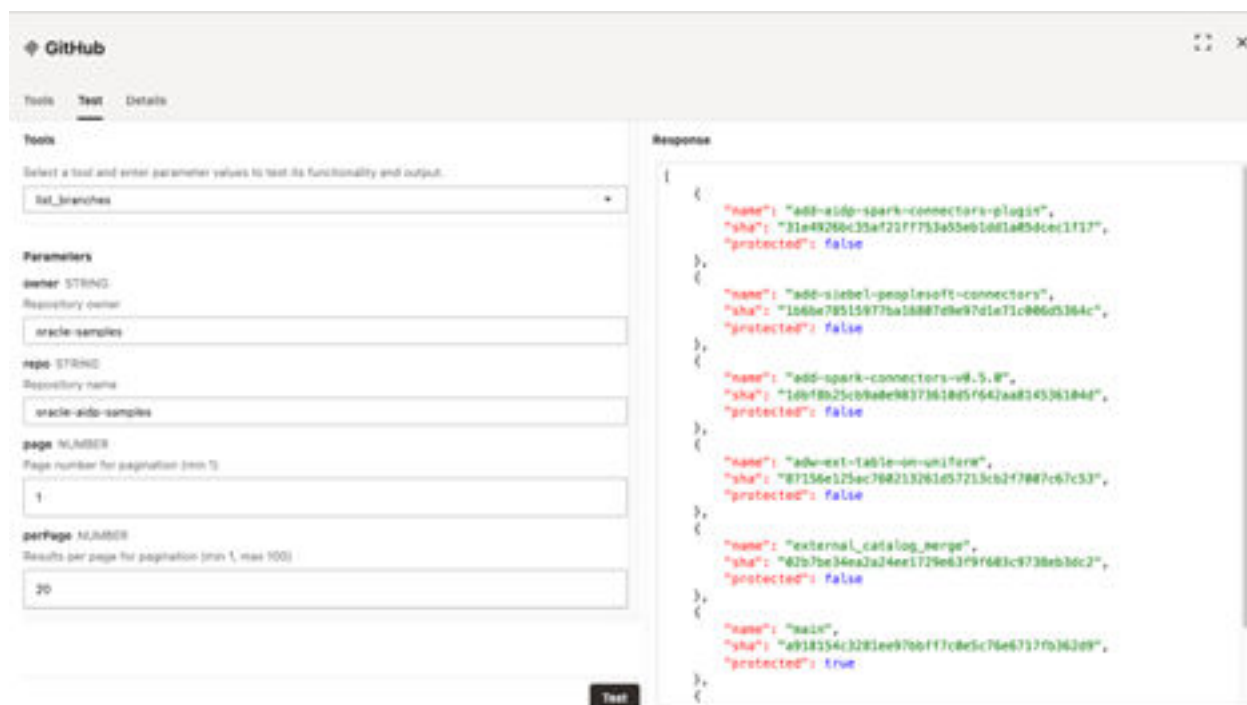
We provide end-to-end code samples for multiple MCP scenarios in the [AI Data Platform Samples GitHub repository](#).

Testing Remote MCP Server Tools

Once tools are selected, the next step is typically to test individual tools to ensure that they behave as expected. This can be done via the Test tab of the MCP tool node.



Select one of the tools you added in the Tools tab, provide parameters values and click on the Test button.

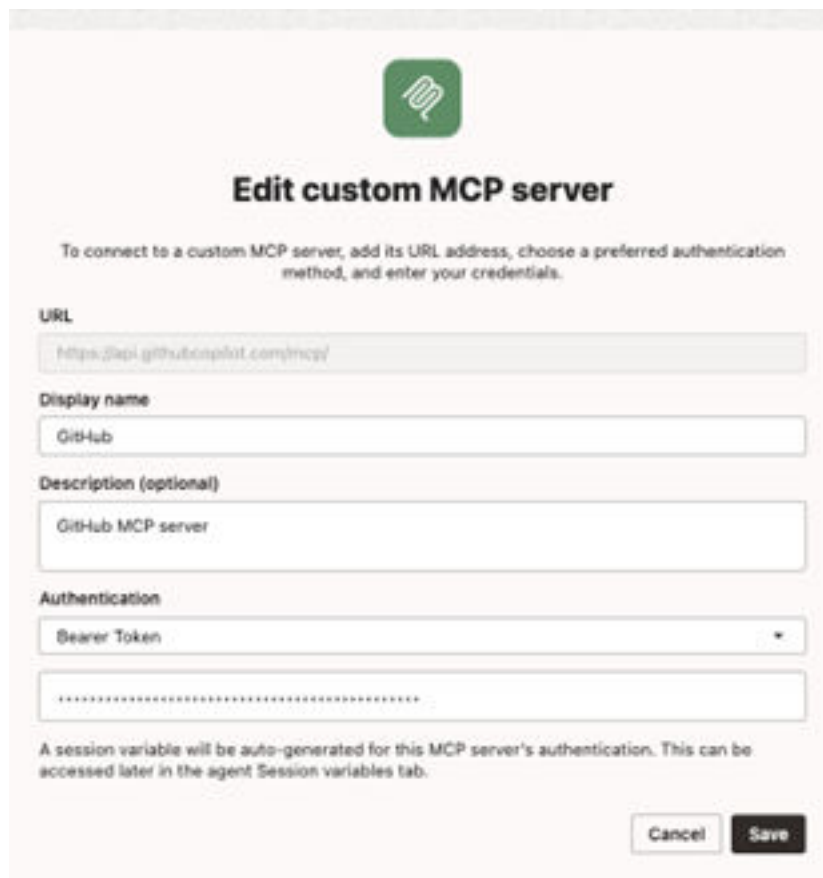


The output of the tool is displayed in the right panel.

The details tab provides information about the authentication method, MCP server URL and the description.



The Edit button next to the authentication method let's you modify the configuration of the remote MCP tool node. You can change the display name, description, and the bearer token used when establishing the connection:



Edit custom MCP server

To connect to a custom MCP server, add its URL address, choose a preferred authentication method, and enter your credentials.

URL

Display name

Description (optional)

Authentication

.....

A session variable will be auto-generated for this MCP server's authentication. This can be accessed later in the agent Session variables tab.

Connect an Agent to a Remote MCP Server from the Visual Builder

You can add access to a remote MCP server to your agent by dragging the Custom MCP server tool node into the canvas.

If you do not see the Custom MCP server tool available, you may need to restart your existing AI compute or create a new one.

Note

The AI compute hosting the agent inherits the networking settings of its workspace. If you enable private network access for the workspace hosting the AI compute, your agent can only reach MCP servers hosted in your selected private VCN and subnet. Your agent may not be able to reach remote HTTP servers available on the public internet.

1. Navigate to your agent.
2. In the Flow tab, under **Tool Templates** click and drag **Custom MCP server** onto the canvas.
3. Provide the server URL for your MCP server.

4. Provide a display name for your MCP server. This is the name of the node that is displayed in the visual builder canvas.
5. Optional: Provide a description for your MCP server. The description field is not provided to the agent.
6. From the **Authentication** drop-down menu, select an authentication method.
 - **No Authentication:** Use this option if the remote MCP server is publicly available and requires no authentication.
 - **Bearer token:** Use this option if the remote MCP server requires an authentication token. You must store the API key in the Oracle AI Data Platform Credential Store and provide a reference to the credential store entry.
7. Click **Connect**. AI Data Platform tests the connection and reports the result.

HTTP Request Tool

The HTTP Request tool lets your agent call any HTTPS REST API.



You configure the request, including method, URL, headers, query parameters, request body, authentication, and optionally, a response optimization step. The agent then invokes the endpoint at runtime. The HTTP request tool is available in both the visual builder and the code builder. In the code builder, the tool is configured through the aidpUtils Python library.

Note

The HTTP Request tool only supports `https://` and `http://` requests. WebSocket connections (ws/wss), binary file uploads, and self-signed certificates are not supported.

Note

The AI compute hosting the agent inherits the networking settings of its workspace. If you enable private network access for the workspace hosting the AI compute, your agent will only reach HTTP endpoints in your selected private VCN and subnet. Your agent cannot reach endpoints available on the public internet.

The following settings must be provided when configuring an HTTP Request tool:

Configuration	Description
HTTP method	The HTTP verb to use. Supported methods are GET, POST, PUT, PATCH, and DELETE.
URL	The full URL of the target endpoint. The URL supports <code>{{sessionVariables.variable_name}}</code> session variable references and <code>{{variable}}</code> runtime parameter references. For example: <code>https://api.example.com/users/{{user_id}}/orders</code> .

Configuration	Description
Timeout	The maximum amount of time the tool will wait for a response from the remote endpoint. The default is 30 seconds and the maximum is 300 seconds.
Authentication type	The authentication method to use when calling the endpoint. See the Authentication section below for the list of supported authentication methods.

 Note

Custom Code tools run on the AI compute attached to your agent. The code has access to the compute environment and outbound network access subject to the workspace networking configuration. Only upload code from sources you trust.

Headers

Headers are key-value pairs sent with the HTTP request. You can add as many headers as needed by clicking the Add new button. Header values can reference session variables and runtime parameters using the `{{variable_name}}` syntax.

 Note

For sensitive headers, you should use the Authentication type field to ensure credentials are injected securely from the Credential Store. Authorization, Cookie, and X-API-Key are sensitive headers and cannot be set through the Headers section.

Query Parameters

Query parameters are appended to the URL as the query string. You can add as many query parameters as needed by clicking the Add new button. Like headers, query parameter values can reference session variables and runtime parameters.

Description

The description field describes what the tool does, when it should be used, and what kind of outputs or effects it produces. The description is provided to the agent and helps the LLM decide when to call the tool.

When writing the description, you should focus on:

- **Purpose:** Explain what the tool is designed to do in one clear sentence. Example: "This tool retrieves customer support tickets from a knowledge base and summarizes them by priority level."
- **When to use it:** Describe the conditions under which the agent should call this tool versus another.
- **Inputs and outputs:** Briefly describe the parameters the tool needs and the shape of what it returns.

HTTP Request Authentication

The HTTP Request tool supports several authentication methods. Select the appropriate method from the Authentication type dropdown.

Authentication Type	Description
No Authentication	No authentication is added to the request. Use this for publicly accessible endpoints.
OCI Resource Principal	The request is signed using the AI compute's OCI Resource Principal. Use this when calling OCI services such as Object Storage or the OCI Generative AI service. Access is governed by OCI IAM policies.
Basic Authentication	A username and password are encoded and sent in the Authorization header. Credentials must be stored in the Credential Store.
Bearer Token	A bearer token is sent in the Authorization header. The token must be stored in the Credential Store.
Header Authentication	An API key is sent in a custom header (such as X-API-Key). The header name is configurable and the key value must be stored in the Credential Store.

When you select an authentication method that requires a secret, the configuration panel displays a credential picker. Click the credential picker to select a previously stored credential, or create a new one from the Credential Store. See the Storing a credential in the Credential Store section of the MCP server documentation for the step-by-step procedure.

Session Variables and Runtime Parameters

Session variables can be referenced in the URL, header values, query parameter values, and request body using the `{{sessionVariables.variable_name}}` syntax. Runtime parameters passed by the agent at invocation time can be referenced using the `{{variable_name}}` syntax.

For example, the following URL combines a session variable for the region with a runtime parameter for the bucket name:

```
https://objectstorage.{{sessionVariables.region}}.oraclecloud.com/n/my-namespace/b/{{bucket}}/o
```

When the tool runs, `{{sessionVariables.region}}` is replaced with the value of the region session variable for the current session, and `{{bucket}}` is replaced with the value the agent passed at invocation time.

Note

Template values are URL-encoded automatically when substituted into the URL or query parameters. You do not need to URL-encode them yourself.

AI Tool Definition

The right side of the configuration panel shows the AI Tool definition. This is the schema that is exposed to the agent and it includes the tool name, description, and the list of runtime parameters the agent can pass when calling the tool. The AI Tool definition is generated automatically from the Description field and from the `{{variable}}` placeholders detected in the URL, headers, query parameters, and body.

The AI Tool definition pane is the panel on the right side of the HTTP tool configuration panel shown earlier in this document. Until you provide a description and define at least one runtime

parameter, the AI Tool definition pane shows a placeholder message. Once you fill in the description and reference at least one `{{variable}}` in the URL, headers, query parameters, or body, the schema is rendered in the pane.

Optimizing Response for the Agent

Many APIs return large responses that include fields the agent does not need. Sending the entire response back to the agent consumes tokens and can degrade the quality of the agent's reasoning. The HTTP Request tool provides a Response optimization section that lets you reduce the response payload before it is returned to the agent.

Three optimization strategies are supported:

- **JSON field selection:** select a subset of fields from a JSON response. You can specify a path to a nested object using dot notation (such as `data.results`), and a list of fields to include or exclude.
- **HTML CSS selector:** extract a subset of an HTML response using a CSS selector (such as `article.content`). Optionally strip HTML tags to return only text.
- **Text truncation:** cap the response at a maximum number of characters to prevent overly large text responses.

Error Handling and Error Codes

When the HTTP request fails, the tool returns a structured error response to the agent. The error includes an error code, a human-readable message, and details about the failure. The agent can use this information to decide whether to retry, fall back to a different tool, or report the failure to the user.

Error Code	Category	Meaning	Retryable
CONNECTION_TIMEOUT	Network	The remote endpoint did not respond within the configured timeout.	Yes
DNS_FAILURE	Network	The hostname in the URL could not be resolved.	Yes
CONNECTION_REFUSED	Network	The remote endpoint refused the connection.	Yes
SSL_CERTIFICATE_ERROR	TLS	The TLS certificate of the remote endpoint could not be validated.	No
UNAUTHORIZED	HTTP 401	The remote endpoint rejected the credentials. Verify that the credential reference is valid and not expired. For OCI Resource Principal, confirm that the AI compute has an active Resource Principal in this environment.	No

Error Code	Category	Meaning	Retryable
FORBIDDEN	HTTP 403	The credentials authenticated successfully but lack permission for the requested resource. Verify API scopes, permissions, or the IAM policy attached to the resource.	No
NOT_FOUND	HTTP 404	The remote endpoint could not find the requested resource.	No
RATE_LIMITED	HTTP 429	The remote endpoint is rate-limiting the caller. Retry after the delay indicated by the Retry-After header.	Yes
SERVER_ERROR	HTTP 5xx	The remote endpoint returned a server error. Often a transient issue.	Yes
SERVICE_UNAVAILABLE	HTTP 503	The remote endpoint is temporarily unavailable.	Yes
INVALID_TEMPLATE	Validation	A {{variable}} reference could not be resolved. Verify that every referenced session variable and runtime parameter is defined and has a value at invocation time.	No
INVALID_URL	Validation	The URL is malformed, uses an unsupported protocol, or resolves to a blocked address (for example, a private IP address or a cloud metadata endpoint).	No
RESPONSE_TOO_LARGE	Validation	The response exceeded the 10 MB maximum response size.	No
RATE_LIMIT_EXCEEDED	Platform	The agent has exceeded the platform's per-agent request rate limit (60 requests per minute) or concurrency limit (10 concurrent requests).	Yes

Each error response includes a guidance field with a suggested next step, and a details field with the elapsed time and any error-specific context such as the HTTP status code.

HTTP Request Tool through LangGraph Code

From the code builder, the HTTP Request tool is configured through the `aidpUtils` Python library. Define an `AIDPToolConf` with `tool_class` set to `HttpEndpointTool` and pass the configuration dictionary in the `conf` field.

```
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.configs import AIDPToolConf

weather_http_tool_def = {
    "method": "GET",
    "url": "https://api.openweathermap.org/data/2.5/weather",
    "params": {
        "q": "{city}",
        "units": "metric",
        "appid": "{api_key}"
    },
    "auth_type": "NO_AUTH",
    "auth_config": {}
}

weather_http_tool_params = [
    {"name": "city", "type": "string",
     "description": "Name of the city."},
    {"name": "api_key", "type": "string",
     "description": "OpenWeather API key."}
]

weather_http_tool_conf = AIDPToolConf(
    name="get_weather",
    description="Get current weather for a city.",
    tool_class="HttpEndpointTool",
    conf=weather_http_tool_def,
    params=weather_http_tool_params
)

weather_tool = create_langgraph_tool(weather_http_tool_conf.model_dump())
```

The `conf` dictionary supports the same fields as the visual builder: `method`, `url`, `headers`, `params`, `body`, `auth_type`, `auth_config`, and `response_optimization`. The `params` list defines the runtime parameters the agent can pass.

auth_type	auth_config Fields
NO_AUTH	{ } (empty)
RESOURCE_PRINCIPAL	{ } (empty)
BASIC_AUTH	username, password (or username_vault_id, password_vault_id for credentials in the OCI Vault)
BEARER_AUTH	bearer_token (or bearer_token_vault_id)
API_KEY_AUTH	api_key (or api_key_vault_id), header_name (default X-API-Key)
OAUTH2_CLIENT_CREDENTIALS	token_endpoint, scope, client_id, client_secret (or client_id_vault_id, client_secret_vault_id)

Test Agent Custom Code Tools

The Test tab lets you execute the tool without running the full agent. Provide values for any runtime parameters and any session variables referenced in the configuration, then click Run to invoke the tool and view the response.

The response panel shows the HTTP status code, the response headers, the response body, and the elapsed time in milliseconds. If response optimization is enabled, the optimized response is also shown alongside the raw response.

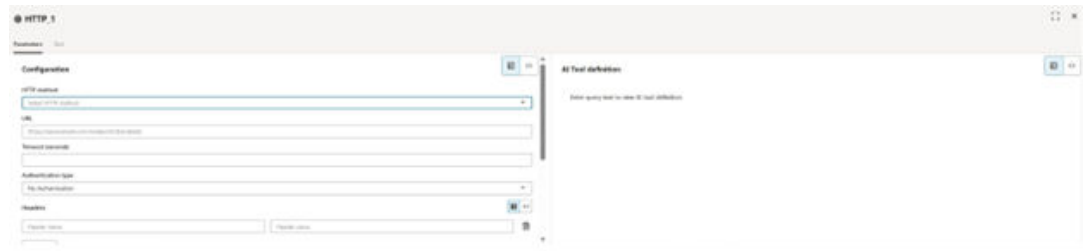
Add an HTTP Request Tool to an Agent

You can add a HTTP request tool to your agents to allow you to call HTTPS REST APIs.

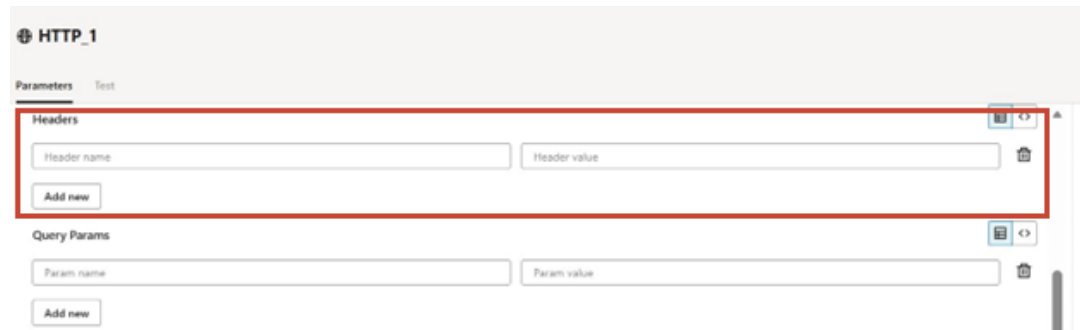
Note

An AI compute must be attached to your agent prior to adding a custom code tool. AI compute is required to install dependencies and run the tool.

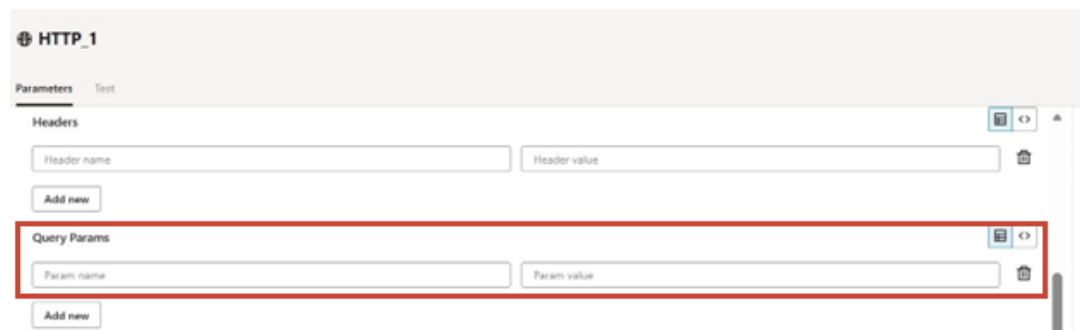
1. Navigate to your agent.
2. From Tool templates, drag and drop an HTTP Request tool to your canvas.



3. In the Parameters tab, provide the **HTTP method**. Supported methods are GET, POST, PUT, PATCH, and DELETE.
4. In **URL**, provide the full URL of the target endpoint. You can use `{{sessionVariables.variable_name}}` session variable references and `{{variable}}` runtime parameter references. For example: `https://api.example.com/users/{{user_id}}/orders`.
5. For **Timeout**, provide the maximum amount of time the tool waits for a response from a remote endpoint in seconds. The maximum timeout value is 300. If no value is provided, the default is 30 seconds.
6. From the **Authentication** drop-down menu, select the appropriate authentication type.
7. Provide the headers for your HTTP request. Click **Add new** to add additional headers.



8. Provide any query parameters for your HTTP request. Click **Add new** to add additional parameters.



9. Optional: Click the Test tab. Provide test parameters and click **Submit**. See test results in the Test results pane.

Prompt Tool

The prompt tool lets you call an LLM in an AI agent with a templated prompt and returns the LLM response back to the agent.



The prompts you provide to the LLM can include parameters that are identified by double braces, for example `{{PARAMETER_NAME}}`. Parameter values are assigned by the agent when the tool is called.

Note

Prompt templates only substitute double-brace placeholders:

- `{{param_name}}` for values from `runtime_params`
- `{{sessionvariables.<name>}}` for session variables

Single-brace `{name}` is sent to the model as literal text and no error is raised. This is intentional, so that literal braces in a prompt (for example, in JSON examples) are preserved.

When to Use Prompt Tools

In principle, instructions written in a prompt tool can be directly included in the agent instructions or provided directly by the end user in a user message. However, there are situations where the prompt tool is a better approach:

- Your prompt is lengthy, requiring detailed format instructions that span several 100s tokens.
- Incorporating the prompt in the agent instructions would increase context usage and significantly increase costs, especially if one is adopting a SOTA LLM for their agent.
- One wants to minimize the size of the instructions given to the agent to reduce cost.
- The task defined by the prompt tool can be handled by a smaller, faster LLM than the reasoning model used by the agent. Smaller models are typically cost efficient, and, in some cases, can be specialized to generate data in a particular modality or format.
- A prompt tool allows structured input parameters to control the output generation. If your use case could be parametrized and generation can vary from session to session, encapsulating the generation in a prompt tool makes sense.

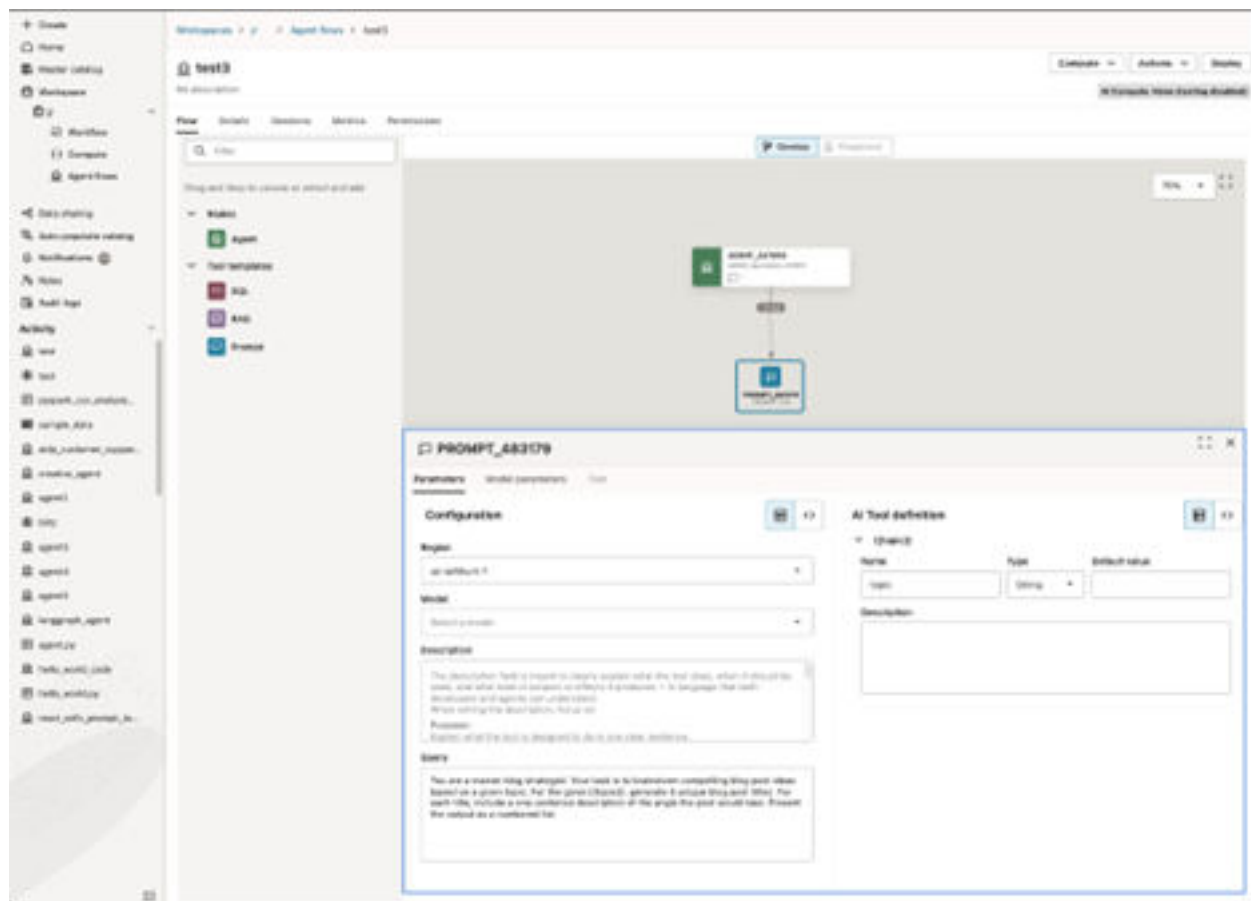
In addition, encapsulating generation instructions in a prompt tool follows many modern agent architecture best practices, including tool re-usability, maintainability, modality, output consistency, scalability, and governance. Some example use cases include:

- Generation of emails, reports, summaries, articles, etc. following a pre-defined, approved structure that can be used as a template
- Generation of complex JSON outputs
- Summation, key sentence extraction, explanation tasks on documents
- Query generation
- Specific modality generation (e.g. images, videos, audio, point cloud data, etc.) that are optimized for a specific model

Prompt Tools through Visual Flow

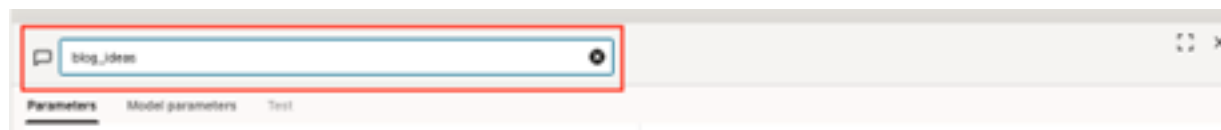
The following is an example of a prompt tool built through visual flow that asks an LLM to generate blog post titles based on a topic assigned by the agent:

You are a master blog strategist. Your task is to brainstorm compelling blog post ideas based on a given topic. For the given {{topic}}, generate 5 unique blog post titles. For each title, include a one-sentence description of the angle the post would take. Present the output as a numbered list.



For this example, you need to configure the following parameters for the prompt tool:

- **Tool name:** Use a descriptive name for the tool to help guide the agent. In this example, we suggest `blog_ideas`. Avoid using unhelpful names like `tool123`.



- **Tool description:** Provide a comprehensive description of what the tool does. If there are limitations to the tool or if there are scenarios in which the tool should not be used, list them in the description field.

The screenshot shows the Oracle AI Prompt Tool interface for a tool named 'blog_ideas'. The interface is divided into two main sections: 'Configuration' and 'AI Tool definition'.

Configuration:

- Region:** A dropdown menu with 'us-ashburn-1' selected.
- Model:** A dropdown menu with 'Select a model' selected.
- Description:** A text area containing the prompt: 'Generate blog title ideas on a given topic. |'. This section is highlighted with a red border.
- Query:** A text area containing the system prompt: 'You are a master blog strategist. Your task is to brainstorm compelling blog post ideas based on a given topic. For the given {{topic}}, generate 5 unique blog post titles. For each title, include a one-sentence description of the angle the post would take. Present the output as a numbered list.'

AI Tool definition:

- Name:** A dropdown menu with 'topic' selected.
- Type:** A dropdown menu with 'String' selected.
- Default value:** A text input field.
- Description:** A text area.

- **OCI region and GenAI service LLM:** Select the OCI region to populate the list of LLMs available in that region, then select your LLM.

blog_ideas

ParametersModel parametersTest

Configuration

Region

us-ashburn-1

Model

Select a model

Description

Generates blog title ideas on a given topic.

Query

You are a master blog strategist. Your task is to brainstorm compelling blog post ideas based on a given topic. For the given {{topic}}, generate 5 unique blog post titles. For each title, include a one-sentence description of the angle the post would take. Present the output as a numbered list

- **LLM parameters:** Parameters like maximum output tokens, temperature, and top p are configured in the Model Parameters tab. If you assign no values, the default values of the OCI Generative AI service are used.



blog_ideas

Parameters Model parameters Test

Maximum output tokens

Temperature

Top p

Top k

Number of Generations (API only)

- **Query:** The prompt used to define the purpose of the tool is defined in the **Query** field.

blog_ideas

ParametersModel parametersTest

Configuration

Region

us-ashburn-1

Model

google.gemini-2.5-pro

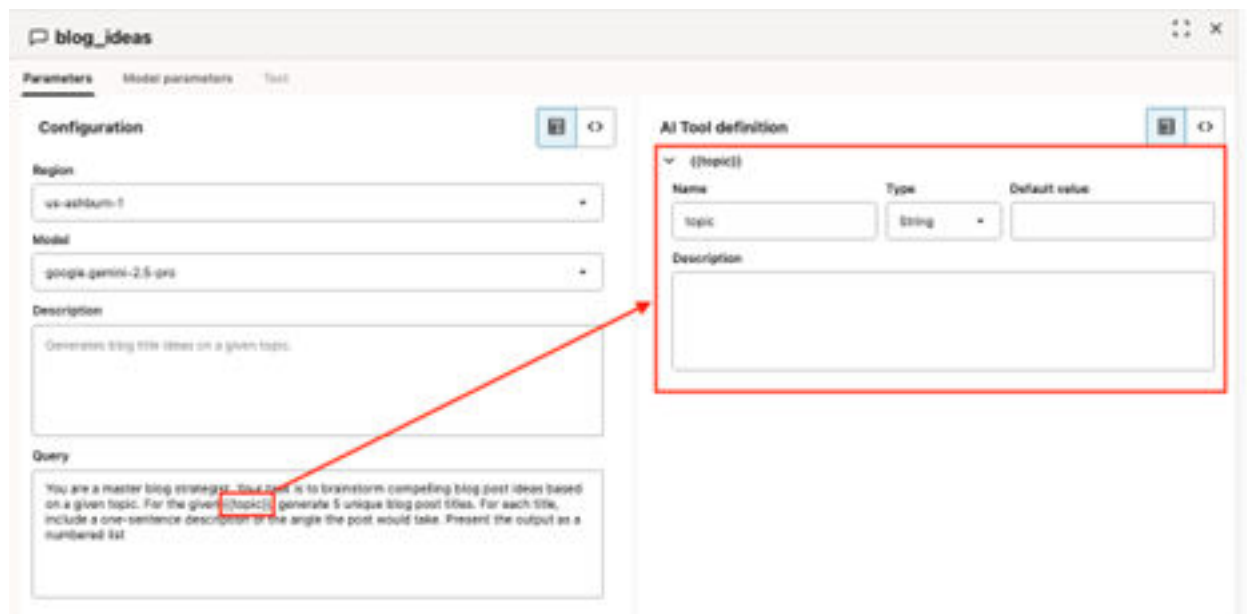
Description

Generates blog title ideas on a given topic.

Query

You are a master blog strategist. Your task is to brainstorm compelling blog post ideas based on a given topic. For the given {{topic}}, generate 5 unique blog post titles. For each title, include a one-sentence description of the angle the post would take. Present the output as a numbered list

Parameters that you define in the prompt auto-populate the AI Tool definition panel. Provide your agent with a description of each parameter as well as the parameter type and default value whenever applicable.



Prompt Tool through LangGraph Code

If you are building your agent through code, you can configure the same prompt tool in the visual flow example as follows:

```
prompt_config = {
    "llm": {
        "model_id": "xai.grok-4",
        "model_provider": "generic",
        "compartment_id": "<your-compartment-ocid>",
        "endpoint": "https://inference.generativeai.<oci-
region>.oci.oraclecloud.com"
    }, "prompt_template": ""
```

```
You are a master blog strategist. Your task is to brainstorm compelling blog
post ideas based on a given topic. For the given {{topic}}, generate 5 unique
blog post titles. For each title, include a one-sentence description of the
angle the post would take. Present the output as a numbered list"
"""
```

```
}
prompt_params = [ {
    "name": "topic",
    "type": "string",
    "description": "Blog topic",
    "defaultValue": "golf"
} ]
```

You then instantiate the AIDPToolConf as follows:

```
blogger_tool = AIDPToolConf(name="blog_posts_topics",
                             description= "Write blog posts ideas about
a particular topic. ",
                             tool_class = "PromptTool",
                             conf=prompt_config params=prompt_params)
```

Lastly, you create a LangGraph compatible tool with the `create_langgraph_tool()` utility function from `aidputils`:

```
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
blogger = create_langgraph_tool(blogger_tool.model_dump())
```

You add the newly created tool to a ReAct agent. In LangGraph, the code looks like this:

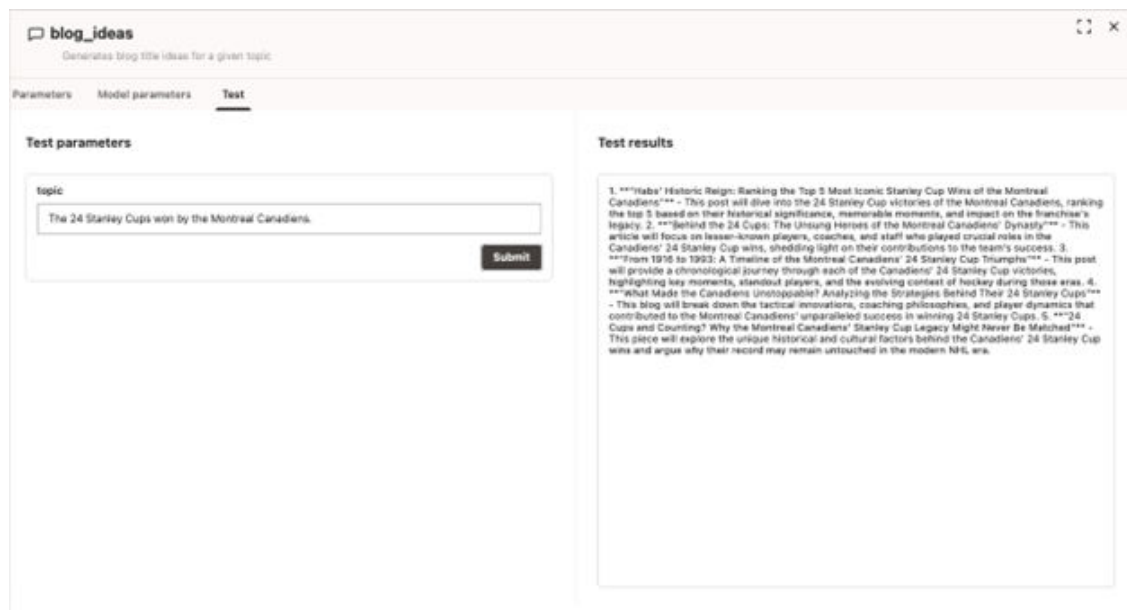
```
tools_agent1 = [blogger_tool]
self.agent = create_react_agent(model=<oci_llm>,
                                tools=tools_agent1,
                                prompt=<system_prompt>,
                                debug=True, checkpointer= checker)
```

Table 18-1 Prompt Tool Configuration Properties

Property	Type	Description
llm	object	LLM connection details and parameters
model_id	string	Identifier of the model to use (e.g., "xai.grok-4")
model_provider	string	Provider name for the LLM model (e.g., "generic")
compartment_id	string	Oracle Cloud Infrastructure (OCI) compartment OCID
endpoint	string	Endpoint URL for the model
prompt_template	string	Prompt template used by the LLM, with variables in {{variable}} format for dynamic insertion

Test Agent Prompt Tools

You test the tool independently of the agent by clicking on the Test tab and filling in the value of each parameter. The prompt is submitted to the LLM you selected.



Ensure your prompt tool is well defined and documented to improve the results from your agent.

Add a Prompt Tool to an Agent

You can add a prompt tool to your agents to allow you to define parametrized prompts you issue to the LLM of your choice.

1. Navigate to your agent.
2. From Tool templates, drag and drop a Prompt tool to your canvas.
3. In the Configuration tab, select the LLM to use and then provide the prompt for the LLM. Click **Code** <> to provide the configuration as JSON code.
4. Provide a **Temperature** for the response as a value between 0.0 and 1.0, where 0.0 provides a strictly factual response and 1.0 provides the most creative response.
5. Click **Apply** →.
6. Provide the definitions for any parameters you established in the configuration. Click **Code** <> to provide the configuration as JSON code.
7. Click ← **Apply**.
8. Optional: Click the Test tab. Provide test parameters and click **Submit**. See test results in the Test results pane.

RAG Tool

The RAG tool issues a natural language query to a vector store and retrieves documents based on the semantic similarity between the query and the stored documents.



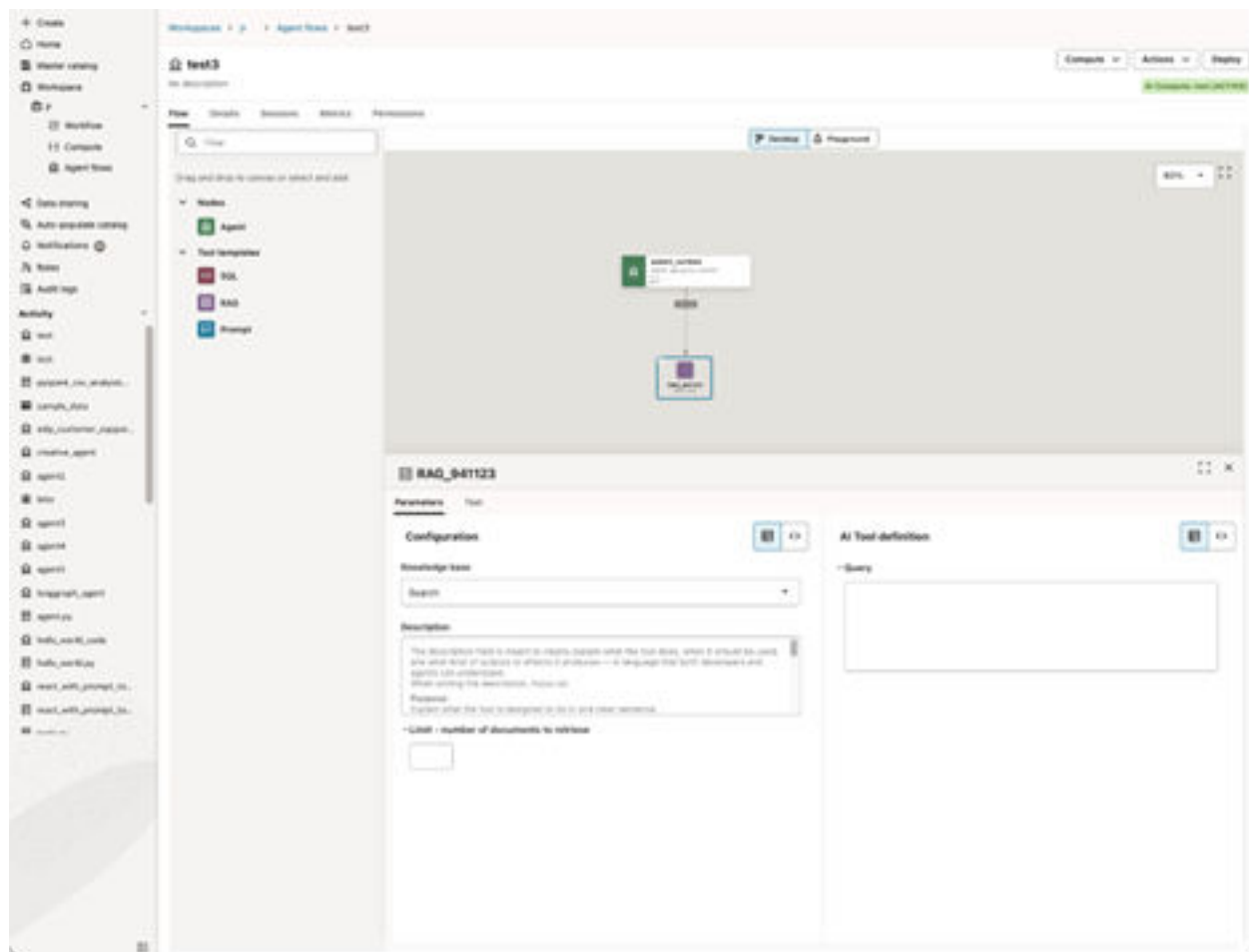
[Video](#)

Note

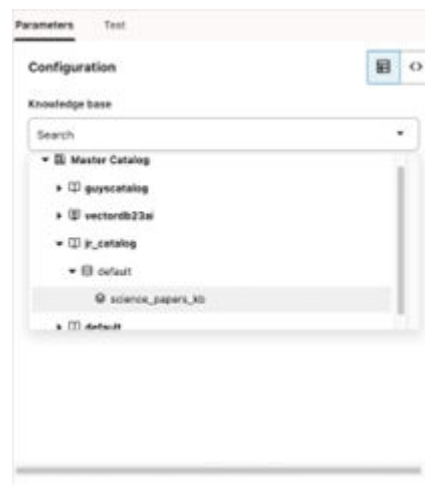
A knowledge base is a prerequisite for the creation of a RAG tool. For more information, see [Knowledge Bases](#).

RAG Tools through Visual Flow

The RAG tool requires you as agent developer provide values for the following parameters:



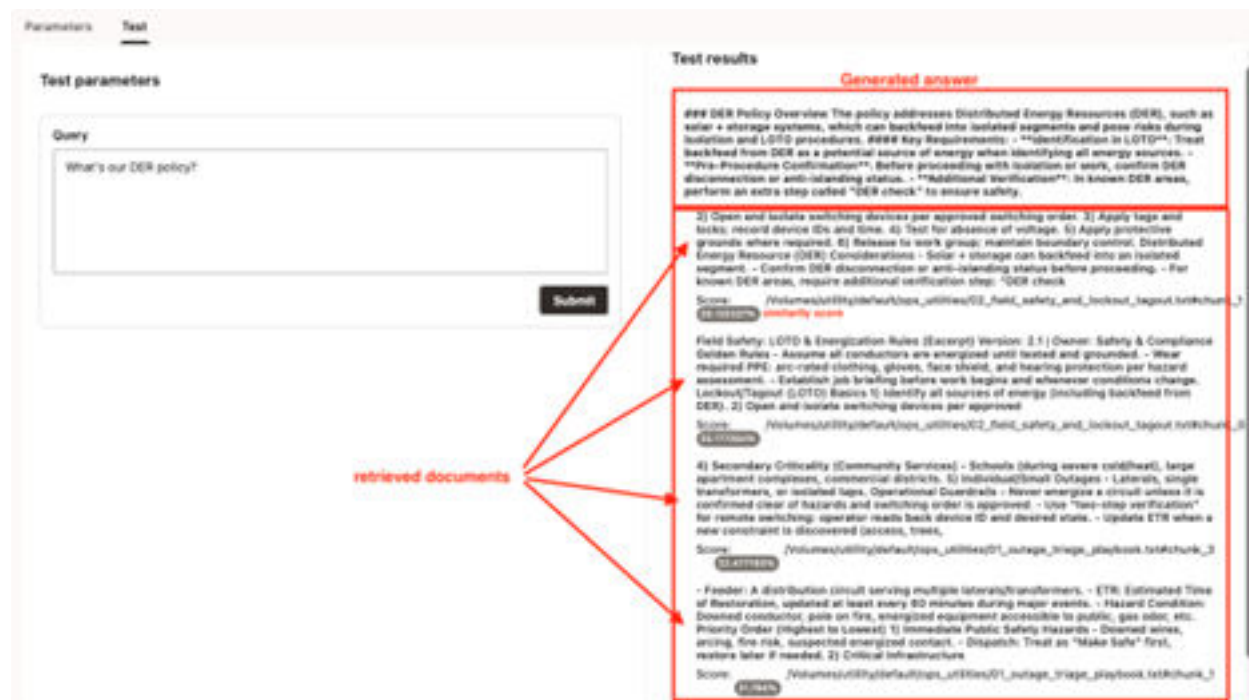
- Agent facing:
 - **Tool name:** A descriptive name for the tool that help you and other users identify its function.
 - **Tool description:** A short summary that provides an overview of the tool.
- Tool configuration:
 - **Knowledge base:** A knowledge base stored in one of your Oracle AI Data Platform catalogs.



The agent will set the value of the query field based on its conversation with the end user. This query field takes a natural language query.

Limit is the number of document chunks you want the tool to retrieve from the vector store. This value is set by the agent developer, not the agent itself.

You can simulate a query issued by the agent by clicking on the test tab of the RAG too:



RAG Tools through LangGraph Code

Building a RAG tool in your agent through code requires configuring the same settings and parameters as the visual flow. For example, you set the RAG parameters as follows:

```
rag_params = [ { "name" : "query",
                  "type" : "string",
```

```
"description" : "<insert a description>",
"defaultvalue" : "<empty>"} ]
```

You then set up the RAG configuration:

```
rag_config = { "catalog": "<catalog>",
               "schema": "<schema>",
               "knowledgeBase": "<knowledge-base-name>",
               "top_k": "<number-of-documents-retrieved>",
               "llm": {
                   "model_id" : "<model-name>",
                   "model_provider" : "<model-provider>",
                   "compartment_id" : "<your-compartment-OCID>",
                   "endpoint" : "https://inference.generativeai.<oci-
region>.oci.oraclecloud.com" }
               }
```

Lastly, you create a LangGraph compatible tool with the `create_langgraph_tool()` utility function from `aidputils`:

```
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
rag_conf= AIDPToolConf(name="<your-tool-name>",
                       description= "<your-tool-description>",
                       tool_class = "RAGTool",
                       conf=rag_config,
                       params=rag_params)
rag_tool = create_langgraph_tool(rag_conf.model_dump())
```

Table 18-2 RAG Tool Configuration Properties

Property	Type	Description
llm	object	LLM connection details
catalog	string	Data catalog identifier
schema	string	Schema within the catalog
knowledgeBase	string	Name or key of the knowledge base to search
top_k	integer	Number of top matching documents to retrieve

Test Agent RAG Tools

You can test the RAG tool from the Test tab after attaching your agent to an AI compute cluster. For more information, see [Attach an Existing AI Cluster to an Agent](#).

Add a RAG Tool to an Agent

You can add a retrieval augmented generation (RAG) tool to your agents to allow the agent to pull relevant external knowledge when generating a response.

1. Navigate to your agent.
2. From Tool templates, drag and drop a RAG tool to your canvas.

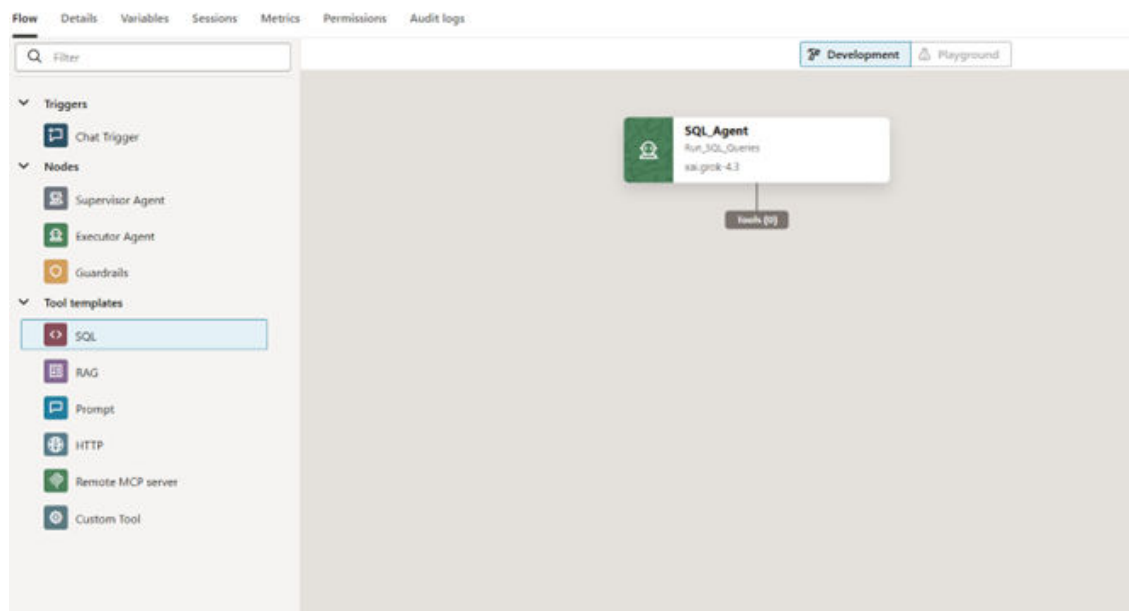
3. In the Configuration tab, select the knowledge base the RAG tool pulls information from and the provide the prompt to define the information to pull. Click **Code** <> to provide the configuration as JSON code.
4. Click **Apply** →.
5. Provide the definitions for any parameters you established in the configuration. Click **Code** <> to provide the configuration as JSON code.
6. Click ← **Apply**.
7. Optional: Click the Test tab. Provide test parameters and click **Submit**. See test results in the Test results pane.

SQL Tool

The SQL tool lets agent developers run predefined SQL queries against tables registered in an Oracle AI Data Platform catalog.



You write the query at design time and define any runtime variables it needs. The agent supplies values for those variables when it calls the tool, and the results return as structured rows the agent can summarize or pass to a downstream node.



The SQL tool supports two query dialects. **Spark SQL** runs against standard catalog tables stored in AI Data Platform and requires a Spark cluster. **Oracle SQL** runs against an external database such as Oracle Autonomous AI Database. You choose the dialect per tool, and the rest of the configuration is the same for both.

Note

The SQL tool is intended for read queries. A typical tool runs a SELECT statement and returns rows. The catalog, schema, and query you configure are private to the tool and are not exposed to the agent. Only the tool name, description, and the AI Tool definition (the runtime variables) are visible to the agent.

Note

The SQL query tool does not automatically start stopped clusters. As a result, the Spark cluster used for your **Spark SQL** query tool should have a duration of **Forever**. If the cluster is allowed to spin down on an idle timeout, Spark SQL queries stop working in production once the cluster stops.

Static and Dynamic Queries

A static query returns exactly what you specify, with no runtime decision by the agent. A dynamic query includes one or more {{variable}} placeholders that signal to the agent that the value is set at run time. For each placeholder you provide a name, a type, an optional default value, and a description the agent uses to choose the value.

For example, the following static query returns a fixed result set:

```
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = 2025
ORDER BY customer_name
```

Replacing the literal with a {{year}} placeholder turns it into a dynamic query the agent can parameterize:

```
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
ORDER BY customer_name
```

As you add placeholders, the AI Tool definition pane populates with each variable so you can set its type, default value, and description.

Placeholders can appear anywhere in the query, including inside functions. The following Spark SQL query matches a severity value case-insensitively:

```
SELECT incident_id, project_id, incident_date, incident_type,
       severity, description, workers_involved, days_lost,
       root_cause, corrective_action, reported_by, status
FROM safety_incidents
WHERE LOWER(severity) = LOWER('{{SEVERITY}}')
```

Give each variable a clear description and a sensible default. The description tells the agent which values are valid, and the default is used when the agent does not supply one.

AI Tool definition

▼ {{SEVERITY}}

Name	Type	Default value
SEVERITY	String	ALL

Description

Incident severity level: 'Major', 'Moderate', 'Minor'

Note

Placeholder names are case sensitive. A placeholder written as `{{SEVERITY}}` and one written as `{{severity}}` are treated as two different variables, unless you consistently use lower case throughout.

Editing the Configuration as JSON

You can edit the SQL tool configuration directly as JSON using the code view toggle. This is useful for copying a tool between agents or making bulk edits.

```
{
  "catalogKey": "construction_data",
  "schemaKey": "admin",
  "query": "SELECT project_id, project_name, client_name, ...",
  "isRowLimitEnabled": null,
  "maxRows": null
}
```

<> query_projects

Parameters Test

Configuration



Input schema

```
{
  "catalogKey": "construction_data",
  "schemaKey": "admin",
  "query": "SELECT project_id, project_name, client_name, project_type, status,\n      budget_usd, actual_cost_usd,\n      completion_percentage,\n      project_manager, location, start_date, estimated_end_date\nFROM projects\nWHERE\n  NVL('{' || STATUS || '}', 'ALL') = 'ALL'\n  OR LOWER(status) = LOWER('{' || STATUS || '}')\n  OR LOWER(project_id) =\n  LOWER('{' || PROJECT_ID || '}')\n  OR LOWER(project_name) = LOWER('{' || PROJECT_NAME || '}')",
  "isRowLimitEnabled": null,
  "maxRows": null
}
```

Row Limits

You can cap the number of rows the tool returns by selecting **Max rows to return** and entering a limit value. Row limits protect performance and control how much data is sent back to the agent.

Set this value relative to the model your agent is using. Larger values can cause agent failures when queries return wide rows or columns with large text values.

Note

If you are seeing unexpected agent errors, start by reducing maxRows.

The row limit is applied to the SQL query itself, before the query runs. Most models detect the limit and surface it to the end user. For a static query, the limit returns the first n available rows.

☒ Max rows to return ⓘ

Limits returned rows for performance. Use lower limits for static queries, and higher limits for dynamic queries to avoid missing results.

1,000



Higher row limits use more tokens and may impact performance or cost.

Note

If you do not want your row limits surfaced for end users, instruct the agent accordingly in its instructions.

Query Examples

You can see query examples and a guide to writing SQL tool queries from the **View Query examples & guide** button.

Configuration

Catalog and schema

guycatalog admin

Description

Retrieves orders placed by customers for a given year.

Query

< > View Query examples & guide

```
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
ORDER BY customer_name
```

The guide showcases different query patterns and provides different recommendations on query parameters.

Query examples & guide

Sample SQL Queries (For Reference Only)

Copy and tailor these examples to your schema. Update table names, column names, and filters before saving.

Basic Query (No Search Parameters)

This is the simplest pattern. It works when the table fits inside the configured maxlines tool.

```
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
ORDER BY customer_name
```

Parameters: year (integer)

Query with Name Search (Recommended)

Use this when the user asks for specific entities. Dynamic search keeps results focused.

```
-- dynamic search: full {{search_term}} from user request
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
AND UPPER(customer_name) LIKE UPPER('{{search_term}}' || '%')
ORDER BY customer_name
```

Parameters: year (integer), search_term (string)

Query with Optional Filter by Column

This pattern supports optional filtering while still keeping a default query path.

```
-- optional filter: use "all" to skip region filter
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
AND ((region = {{region}} OR '{{region}}' = 'all')
ORDER BY customer_name
```

Parameters: year (integer), region (string, use "all" to skip)

Query with Combined Search and Filter

Use this for broad requests where both a name search and an optional column filter may be needed.

```
-- combined: name search + optional region filter
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
AND (UPPER(customer_name) LIKE UPPER('{{search_term}}' || '%')
OR '{{search_term}}' = 'all')
AND ((region = {{region}} OR '{{region}}' = 'all')
ORDER BY customer_name
```

Parameters: year (integer), search_term (string), region (string)

Aggregated Query (Avoids the Blind Spot Entirely)

Use aggregation for analytical use cases. Power output over reduce maxlines cap max.

```
-- aggregated: one row per region
SELECT region, COUNT(*) AS customer_count, SUM(amount) AS total_amount
FROM test_customers
WHERE period_year = {{year}}
GROUP BY region
ORDER BY total_amount DESC
```

Parameters: year (integer)

SQL Tools through LangGraph Code

Just as with the visual flow, you start creating a SQL tool for your agent through LangGraph code by creating a query:

```
sql_config = { "catalogKey": "adw23ai_phx",
               "schemaKey": "gold",
               "query": ""Select ... from ... limit {{max_number}}"" }
```

You document each parameter in the SQL query in the params argument with a name, type, description, and optionally, a defaultValue.

```
sql_params = [ { "name" : "max_number",
                  "type" : "string",
                  "description" : "<your-description>",
                  "defaultValue" : "<your-default-value>" } ]
```

Lastly, you create a LangGraph compatible tool with the `create_langgraph_tool()` utility function from `aidputils`:

```
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
sql_conf= AIDPToolConf(name="<your-tool-name>",
                        description= "<your-tool-description>",
                        tool_class = "SQLTool",
                        conf=sql_config,
                        params=sql_params)
sql_tool = create_langgraph_tool(sql_conf.model_dump())
```

Table 18-3 SQL Tool Configuration Properties

Property	Type	Description
catalogKey	string	Identifier for the catalog or database connection
schemaKey	string	Schema name within the catalog/database
query	string	SQL Query string, may include placeholders in {{}}

Test Agent SQL Tools

The Test tab runs the tool on its own, without executing the full agent. Testing works the same way for both dialects. Open the Test tab, provide a value for each runtime parameter (or use the defaults), and click Submit to run the query and view the response.

Note

Testing a tool requires that your agent is attached to an AI compute. An AI compute is attached if the AI Compute label is green with the selected AI compute in an ACTIVE state.

SQL Command Reference

SQL tool queries are read queries built from the standard SQL clauses. The Oracle SQL dialect follows Oracle SQL against the external database. The Spark SQL dialect targets standard catalog tables, which are Delta Lake tables; the standard catalog currently runs Spark 3.5 with Delta Lake 3.2.0. Most clauses are written the same way in both dialects, because both follow standard SQL. The main difference is how each dialect limits the number of rows. The following table lists the clauses and keywords most often used in SQL tool queries, with the form for each dialect.

Keyword or clause	Purpose	Oracle SQL	Spark SQL
SELECT	Choose the columns to return	SELECT col1, col2	
DISTINCT	Return only unique rows	SELECT DISTINCT col	SELECT DISTINCT col
FROM	Name the source table	FROM table_name	FROM table_name
WHERE	Filter rows by a condition	WHERE col = value	WHERE col = value
AND OR NOT	Combine or negate conditions	a AND b OR NOT c	a AND b OR NOT c

Keyword or clause	Purpose	Oracle SQL	Spark SQL
IN	Match any value in a list	col IN (a, b, c)	col IN (a, b, c)
BETWEEN	Match an inclusive range	col BETWEEN x AND y	col BETWEEN x AND y
LIKE	Match a text pattern	col LIKE 'A%'	col LIKE 'A%'
IS NULL	Test for missing values	col IS NULL	col IS NULL
ORDER BY	Sort the result	ORDER BY col DESC	ORDER BY col DESC
GROUP BY	Group rows for aggregation	GROUP BY col	GROUP BY col
HAVING	Filter grouped rows	HAVING COUNT(*) > 1	HAVING COUNT(*) > 1
JOIN ON	Combine rows from two tables	a JOIN b ON a.id = b.id	a JOIN b ON a.id = b.id
AS	Alias a column or table	col AS name	col AS name
UNION ALL	Combine two result sets	q1 UNION ALL q2	q1 UNION ALL q2
CASE	Return a value conditionally	CASE WHEN c THEN x END	CASE WHEN c THEN x END
Aggregates	Summarize over rows	COUNT SUM AVG MIN MAX	COUNT SUM AVG MIN MAX
Row limit	Cap the number of rows	FETCH FIRST n ROWS ONLY	LIMIT n

Note

You normally do not write the row limit yourself. The Max rows to return setting applies it for you. The FETCH FIRST and LIMIT forms are useful only when you want an explicit limit inside the query.

For complete SQL grammar and the query engines behind each dialect, see the following references:

Oracle AI Data Platform SQL

- [SQL Grammar](#)
 - [Catalog SQL Grammar](#)
 - [Schema SQL Grammar](#)
 - [Volume SQL Grammar](#)
 - [Table SQL Grammar](#)
 - [DML Queries](#)

Spark SQL and Delta Lake (Standard Catalog)

Standard catalog tables are Delta Lake tables. The standard catalog currently runs Spark 3.5 with Delta Lake 3.2.0.

- [Apache Spark SQL Syntax: DML statements](#) Spark SQL query and DML statement syntax.
- [Delta Lake: Table deletes, updates, and merges](#) DELETE, UPDATE, and MERGE operations on Delta tables.
- [Delta Lake: Table utility commands](#) Utility operations such as OPTIMIZE and VACUUM.
- [Delta Lake: Use liquid clustering for Delta tables](#) Liquid clustering for Delta table layout.

Add a SQL Tool to an Agent

You can add a SQL tool to your agents to allow the agent to execute SQL queries against structured data sources in registered external catalogs.

1. Navigate to your agent.
2. From Tool templates, drag and drop a SQL tool to your canvas.
3. Click and drag the connector handle on your agent to connect to the tool node.



4. Double-click the SQL node to open the configuration panel.
5. Provide a name and description for your tool. The description is provided to the agent and helps it decide when to call the tool.
6. Choose the **Query dialect**:
 - **Spark SQL** writes queries against standard AI Data Platform catalogs.
 - **Oracle SQL** writes queries against external AI Data Platform catalogs.

Note

Spark SQL requires a running Spark cluster in your Oracle AI Data Platform workspace.

Configuration

Query dialect

☒ Spark SQL ☐ Oracle SQL

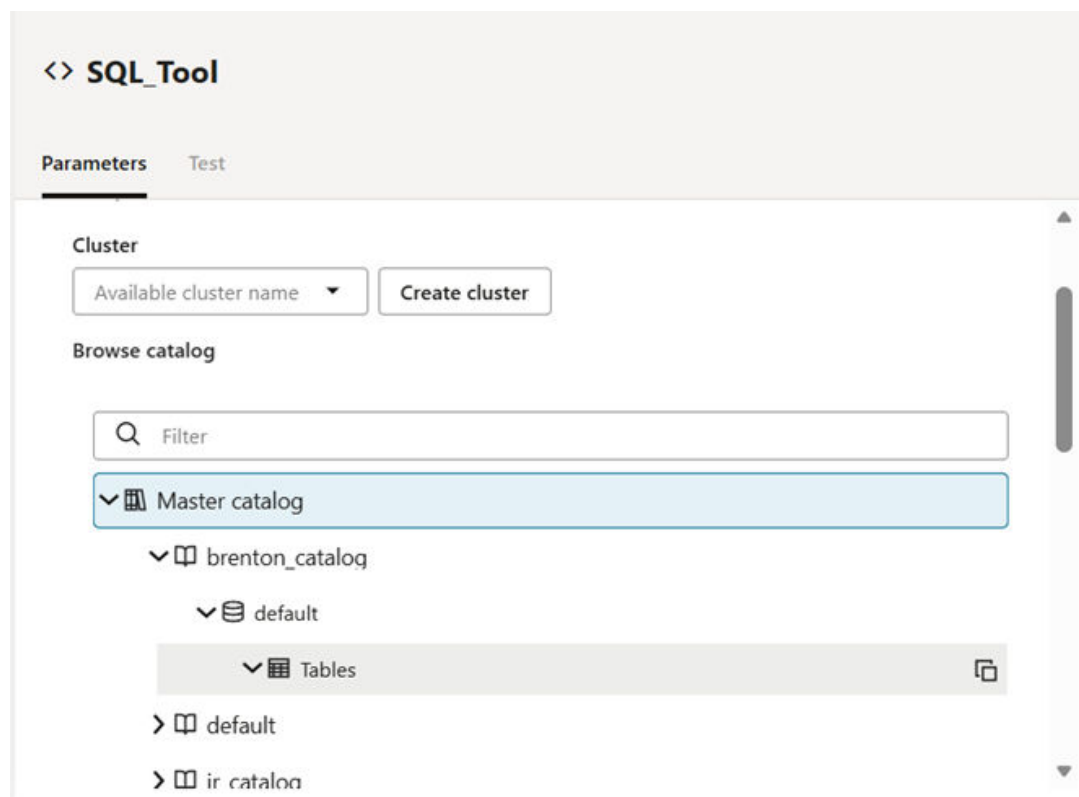
- From the **Cluster** drop-down, select a running Spark cluster. Click **Create cluster** to provision a new Spark cluster. See [Create a Custom Cluster](#) for guidance on creating a new cluster.

Note

The SQL query tool does not automatically start stopped clusters. As a result, the Spark cluster used for your **Spark SQL** query tool should have a duration of **Forever**. If the cluster is allowed to spin down on an idle timeout, Spark SQL queries stop working in production once the cluster stops.



- Under **Browse catalog**, use the search field to locate a catalog by name or click through the catalog manager to locate the catalog.



- In the **Query** field, enter your query. Click **View Query examples and guide** to open a panel containing ready-made patterns you can copy or adapt.

<> SQL_Tool

Parameters

Test

Description

Standard Catalog Spark SQL tool

Query ?

<> View Query examples & guide

☐ Max rows to return ?

Limits returned rows for performance. Use lower limits for static queries, and higher limits for dynamic queries to avoid missing results.

10. Select **Max rows to return** to limit the number of rows returned by your query results.

Note

If your query can return more rows than this limit, consider adding search parameters like `{{customer_name}}` or `{{region}}` so the agent can find more specific data.

11. In the AI Tool definition pane, set the type, default value, and description for the variables set in your query.

SQL 1

ParametersText

Good example:

"description": "Analyzes a given image and returns a classification label describing its main object. Use this tool when the user provides an image and asks what it depicts."

Bad Example:

"description": "Uses a TensorFlow CNN model to output a top-1 prediction for the provided .JPG input."

Query

<> View Query examples & guide

```
SELECT customer_name, region, amount, category
FROM test_customers
WHERE period_year = {{year}}
ORDER BY customer_name
```

Max rows to return

Limits returned rows for performance. Use lower limits for static queries, and higher limits for dynamic queries to avoid missing results.

AI Tool definition

▼ {{year}}

Name

year

Type

String

Default value

Description

12. Optional: Click the Test tab. Provide test parameters and click **Submit**. See test results in the Test results pane.

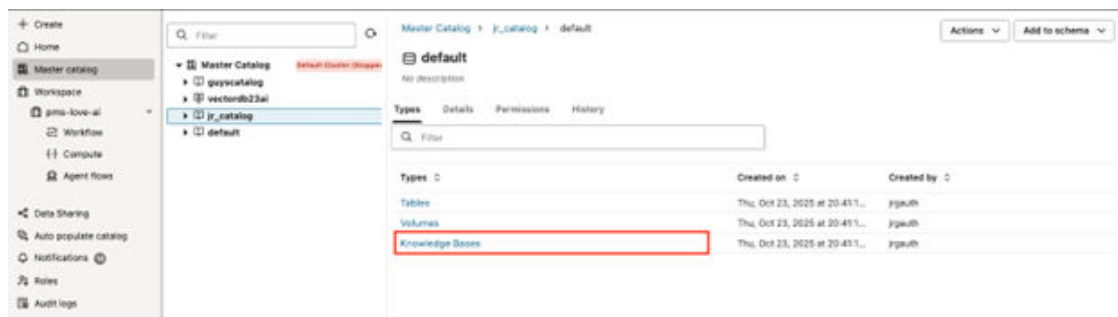
19

Knowledge Bases

Knowledge bases leverage Oracle Database 26ai Vector Search capability to store vector embeddings from documents stored in Oracle AI Data Platform.



Through Oracle Database 26ai's vector search capabilities, knowledge bases empower AI agents to perform semantic searches and retrieve semantically relevant documents. In AI Data Platform, knowledge bases are created in a schema of a catalog under the Knowledge Bases type.



In AI Data Platform, knowledge bases are created in schemas of standard catalogs using the Knowledge Base type. The ingestion of PDF, DOCX, and TXT files stored in either managed or external volumes is supported in knowledge bases. By default, vectors are stored in Oracle Database 26ai Vector Search instance that is provisioned in your tenancy when your instance of AI Data Platform is created.

AI Data Platform supports two embeddings models:

- **ALL_MINILM_L12_V2:** A sentence-transformers model that maps sentences and paragraphs to a 384 dimensional dense vector space. Used for tasks like clustering or semantic search.
- **MULTILINGUAL_E5_SMALL:** Generates vector embeddings for text in multiple languages. Its compact design enables effective performance across various languages, suitable for diverse datasets and multilingual scenarios.

When creating a knowledge base, you need to select **Chunk size** and **Chunk overlap**. **Chunk size** is the number of characters in each piece of text used when searching the knowledge base. **Chunk overlap** is the number of characters shared by consecutive pieces of text. Overlap helps preserve context when a sentence or idea spans more than one chunk.

Note

By itself, a knowledge base object in AI Data Platform cannot be directly queried. You query a knowledge base by creating a RAG tool attached to an agent and selecting the relevant knowledge base. For more information about RAG tools, see [RAG Tool](#). For more information about AI Agents, see [AI Agents](#).

Ingest Data Sources

After you create a knowledge base in AI Data Platform, you need to go into that knowledge base and specify a data source to ingest data from. You can select an entire volume or a folder in a volume as a source for ingestion, but you cannot select individual files.

You can see your data sources in the Data Source tab of your knowledge base and see the information about that data source by clicking its name. The Parameters tab provides information about the selected volume, file path, attached cluster, and file types.

Note

AI Data Platform does not support scheduled ingestion jobs. You can ingest data immediately by clicking **Ingest now** in the Parameters tab of your data source.

You can see more detailed information about your data source in the Details tab and see a history of data ingestion jobs in the Job runs tab.

Create a Knowledge Base

Creating a knowledge base in Oracle AI Data Platform is a one-time setup that allows you to register a document source, automatically chunk, embed, and index files, and enable semantic search and RAG retrieval via agent flows.

You cannot directly query knowledge bases in AI Data Platform. You can query knowledge bases by creating a RAG tool attached to an AI agent. For more information, see [AI Agents](#).

1. Click **Master catalog**.
2. Navigate to the standard catalog and schema where you want to create your knowledge base.
3. Click **Knowledge Bases**.
4. Click  **Create Knowledge Base**.

Create Knowledge Base

Name Required

Description

Advanced settings (optional)

Workspace for compute Cluster used for ingestion

Embedding model

Chunk size Chunk overlap

5. Provide a name and description for your knowledge base.
6. Select a workspace and Spark cluster for file ingestion. If no cluster is selected, the Default Master Catalog Compute is used.
7. Select the embedding model used, if necessary.
8. Provide the chunk size and chunk overlap, if necessary.
9. Click **Create**.

Edit a Knowledge Base

You can modify the name, description, cluster, model, or chunking details for an existing knowledge base if you have the relevant permissions.

1. Navigate to your knowledge base folder.
2. Next to the knowledge base you want to edit, click ... **Actions**, then click **Edit**.
3. Make any changes to the attributes of the knowledge base.
4. Click **Save**.

Delete a Knowledge Base

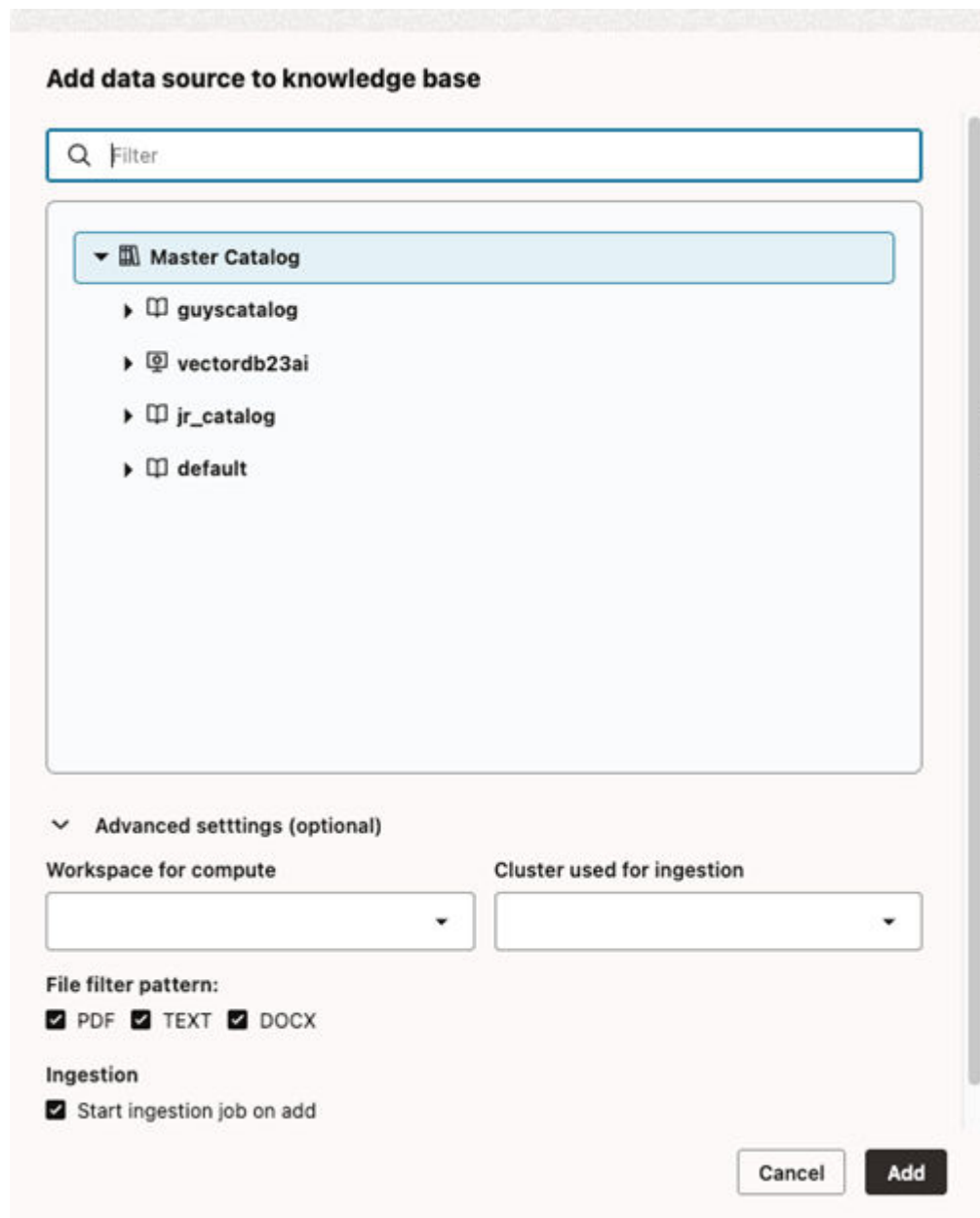
You can delete knowledge bases you no longer need or use from your catalog.

1. Navigate to your knowledge base folder.
2. Next to the knowledge base you want to delete, click ... **Actions**, then click **Delete**.
3. Click **Delete**.

Add a Data Source to a Knowledge Base

After you have created a knowledge base, you need to assign it a data source for ingestion.

1. Navigate to your knowledge base.
2. Click the Data Source tab.
3. Click  **Add data source to knowledge base**.



The screenshot shows a dialog box titled "Add data source to knowledge base". At the top is a search bar with a magnifying glass icon and the text "Filter". Below the search bar is a list of data sources under the heading "Master Catalog". The list includes "guyscatalog", "vectordb23ai", "jr_catalog", and "default". Below the list is a section titled "Advanced settings (optional)". This section contains two dropdown menus: "Workspace for compute" and "Cluster used for ingestion". Below these are checkboxes for "File filter pattern": "PDF", "TEXT", and "DOCX", all of which are checked. At the bottom of the advanced settings is a checkbox for "Ingestion" labeled "Start ingestion job on add", which is also checked. At the bottom right of the dialog are two buttons: "Cancel" and "Add".

4. In the Master Catalog, select the volume or folder in a volume you want to ingest into your knowledge base. You cannot select individual files.
5. If necessary, select the compute cluster to use for data ingestion.

6. Select which file types to ingest. Supported file types are PDF, TXT, and DOCX.
7. Select **Start ingestion job on add** to start ingestion immediately after adding the data source.
8. Click **Add**.

Ingest Data to a Knowledge Base

Once a data source is added to a knowledge base, you can manually start a data ingestion job run from the Parameters tab.

1. Navigate to your knowledge base.
2. In the Data Source tab, click the name of the data source you want to run an ingest data job for.
3. In the Parameters tab, click **Ingest now**.

View Ingestion Job Run Status

You can view a list of all ingestion jobs for the data source from the Job Runs tab of the data source.

1. Navigate to your knowledge base.
2. In the Data Source tab, click the name of the data source you want to view the status for.
3. Click the Job runs tab.
4. Use the filters to narrow the list of displayed job runs.

Delete a Data Source

You can delete data sources you no longer need or use from your knowledge base.

Deleting a data source also deletes the corresponding vector embeddings from your Oracle AI Data Platform.

1. Navigate to your knowledge base. Click the Data Sources tab.
2. Next to the data source you want to delete, click ... **Actions**, then click **Delete**.
3. Click **Delete**.

Agent Memory

Agent memory is the part of an AI agent system that lets the agent retain and reuse information across turns, tasks, or sessions.

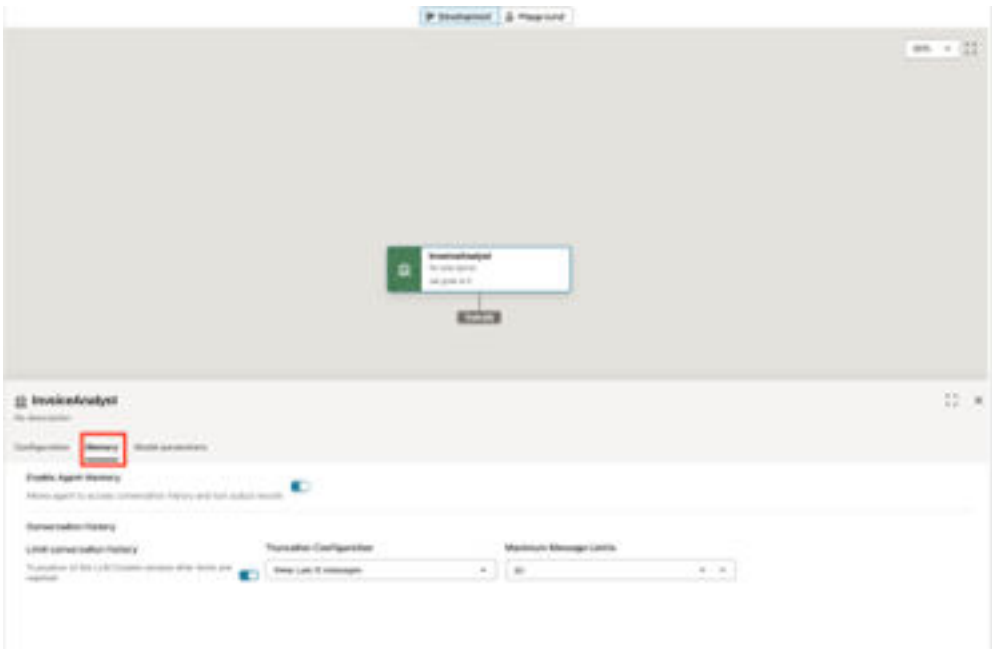
Unlike the model's context window, which is temporary and limited to the current prompt, memory can persist facts, preferences, prior decisions, tool outputs, intermediate plans, or observations about the environment.

Agent Memory in Oracle AI Data Platform

AI Data Platform provides short-term memory that is limited to the duration of a session. You can configure what can be kept in memory in the Memory tab of your agent.

Single-Agent Memory Configuration

Memory configuration of a single agent system can be found in the Memory tab of the agent node.



Memory Configuration	Description
Enable Agent Memory	This setting enables agent memory. When disabled, the agent is essentially a stateless system. Each turn is treated independently, and no follow up question can be asked. We recommend that memory is enabled.

Memory Configuration	Description
Limit conversation history	When this selection is disabled, previous turns will fill up the memory until the model runs out of context and an error is returned. If short sessions are expected, it is ok to disable this setting. However, for most use cases we recommend limiting conversation history.
Truncation configuration	If you opt to limit the conversation history, you can decide to truncate the history by: • Keeping only the last N messages (user + agent) • Setting an overall token budget (first in, first out) • Or both, whichever one comes first triggers truncation of the agent memory
Maximum Message Limits	If you elect to keep the last N messages, you can set a value of N.
Token budget	Alternatively, you can set an overall token budget. First tokens are eliminated to keep the most recent ones in memory.

Multi-agent System Memory Configuration (Supervisor Pattern)

The memory of a multi-agent system is configured in the Memory tab of the supervisor agent.

The screenshot displays the Oracle AI Data Platform interface for configuring a multi-agent system. At the top, a diagram shows the 'AccountsManager' agent (version 4.2) receiving input from a 'Message' icon and managing two sub-agents: 'InvoiceAnalyst' and 'AccountPayableAnalyst' (both version 4.2). Below the diagram, the 'AccountsManager' configuration panel is visible, with the 'Memory' tab selected. This tab includes the following settings:

- Enable Agent Memory:** A toggle switch that is currently turned on.
- Conversation History:**
 - Limit conversation history:** A toggle switch that is currently turned on.
 - Truncation Configuration:** A dropdown menu set to 'Keep Last N messages'.
 - Maximum Message Limits:** A text input field set to '20'.
- State Isolation for Executor Agents:** Three radio button options: 'Stateless' (selected), 'Private', and 'Shared'.

Memory for a multi-agent system is enforced and cannot be disabled and the memory truncation options displayed in the supervisor agent node only applied to the supervisor agent memory. Each executor agent memory truncation policy can be configured in the executor node Memory tab based on the state isolation policy selected for the entire system.

The multi-agent system memory configuration also allows you to select the memory sharing policy of the executor agents. Three options are possible: Stateless, Private, and Shared. Note that the policy is applied to all executor agents.

State Isolation for Executor Agents	Description
Stateless	Each executor agent sees only the task assigned by the supervisor. No history is carried over between calls. No follow up request can be made to the supervisor agent by the executor agent. If stateless is selected, the memory of each executor agent is disabled.
Private	Each executor agent sees only its own past interactions. It cannot see other executor agents or the original user conversation with the supervisor agent.
Shared	Executor agents can see the full conversation history across agents and user. All agents work from one shared context. If shared is selected you can configure the memory truncation policy separately for each executor agent in each agent node Memory tab.

Session Variables

You can use custom, agent-defined session variables to provide additional contextual data points to your agents during a user session.

What are Session Variables?

Custom, agent-defined session variables provide additional contextual datapoints to the agent during a user session. The variables can be used for a variety of purposes, including setting values of parameters in tools, giving overall instructions to an agent, and providing contextual information about the caller and/or the application where the agent is embedded.

Note

Prompt templates only substitute double-brace placeholders:

- `{{param_name}}` for values from `runtime_params`
- `{{sessionvariables.<name>}}` for session variables

Single-brace `{name}` is sent to the model as literal text and no error is raised. This is intentional, so that literal braces in a prompt (for example, in JSON examples) are preserved.

Here's a simplified scenario where we are building a customer support agent. The customer support agent is integrated within a retail website. When a user logs into the retail website, information about that user is captured by the client application (`userID`, `username`, `geo location`, `device used`, `shopping cartID`, etc.) and that information could be used by the downstream support agent. Let's look at an example of how these session parameters could be used within the agent instructions field in AIDP:

You are a customer support agent for the retail website `belts-and-buckles.com`, specializing in the sales of belts and buckles. Your objective is to answer questions that customers have about their current and past orders, answer questions about items they put in their shopping cart, and answer general questions about `belts-and-buckles.com`.

A few guidelines before your start:

You are interacting with user `{{sessionVariable.userName}}`. Always start with a welcome message: Hello `{{sessionVariable.preferredSalutation}}` `{{sessionVariable.userName}}`! What's the weather like today in `{{sessionVariable.currentUserGeo}}`?

Three session variables are used in the example above:

- `userName`
- `PreferredSalutation`
- `currentUserGeo`

The agent has no prior knowledge of the user interacting with the retail website, doesn't know what the user location is or has no context about what's in the user's shopping cart. In

principle, the client application could know all or a subset of those values and pass those values in a request to the agent. Here's an example of what a request to the agent could look like, including the session parameters.

Note

This example illustrates what this request could look like. It is not representative of an implementation decision.

```
"input": [
  { "role": "user",
    "content": [
      { "type": "input_text",
        "text": "What material is the belt Mr Outcast made of?",
        "variables": ["userName": "Paul", "preferredSalutation": "Hon",
"cartID": NULL, "currentUserGeo": "Cancun, MX"]
      }
    ]
  }
]
```

The agent would answer back: "Hello Mr Paul, what's the weather like today in Cancun, Mx?"

Another use of these session parameters is in setting parameter values in tools. For example, a SQL query could retrieve the content of a shopping cart by using the session parameter `{{sessionParam.CartID}}` :

```
Select productID, productName, productDescription,
productPrice from cartTable where cartID ==
{{sessionVariable.cartID}}
```

The session variables are defined by the agent developer when they create their agents and the values of those attributes are set by the client application when a session is created or resumed.

You can configure the following settings when create a new session variable:

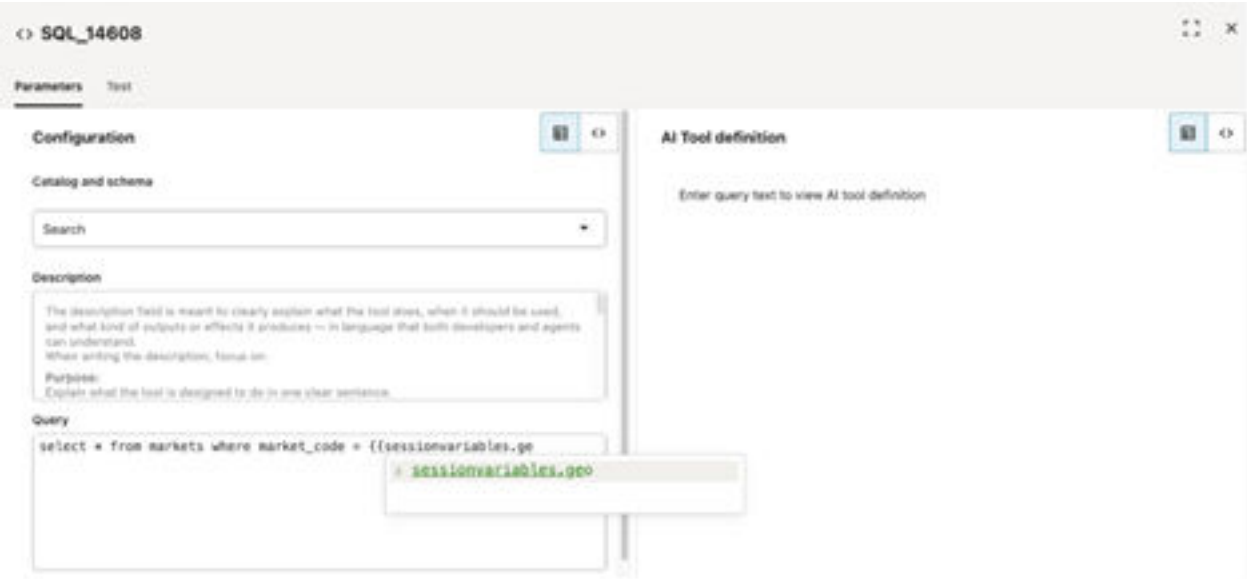
Setting	Description
Required variable	Enable to make this session variable required for each invocation call of the agent. Disabling makes the session variable optional for each invocation call.
Log variable	Enable to capture the value of the session variable in logs and traces. Disabling prevents the variable from appearing in logs. We recommend disabling this setting for variables with sensitive data.
Name	Name of the session variable. Use a descriptive name to make it easy for you and other users to determine the purpose of the variable.
Default value	If defined, the default value is assigned to the session variable if another value is not defined in the invocation call. If left blank, a value must be assigned as part of the invocation call.

Setting	Description
Description	Description of the session variable. Provide a helpful description so that you and other users are able to understand the function of the session variable.

Example: Using Session Variables in Tools Configuration

You can use a session variable in a SQL tool configuration as part of the SQL query itself.

In this example, the session variable geo is used to filter the result of the SQL query:



You can use session variables and tool parameters within the same query. In the example below, the `titleID` parameter is set by the agent while the session variable `geo` is provided by the calling application.

The screenshot shows the configuration window for SQL_14608. The 'AI Tool definition' tab is active. It contains a table with columns 'Name', 'Type', and 'Default value'. The table has one row with 'titleID' as the name and 'String' as the type. Below the table is a 'Description' field. The 'Configuration' tab is also visible, showing a 'Catalog and schema' search field, a 'Description' field with instructions, and a 'Query' field with a SQL query.

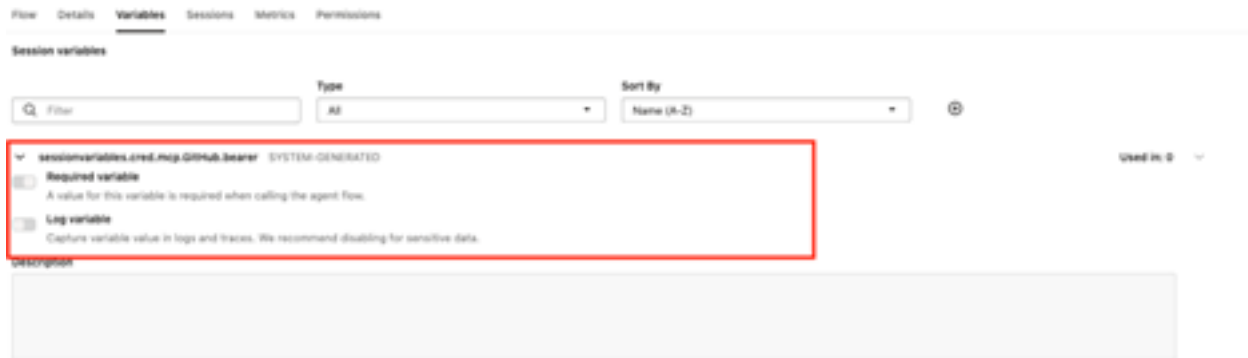
Name	Type	Default value
titleID	String	

System-generated Session Variables

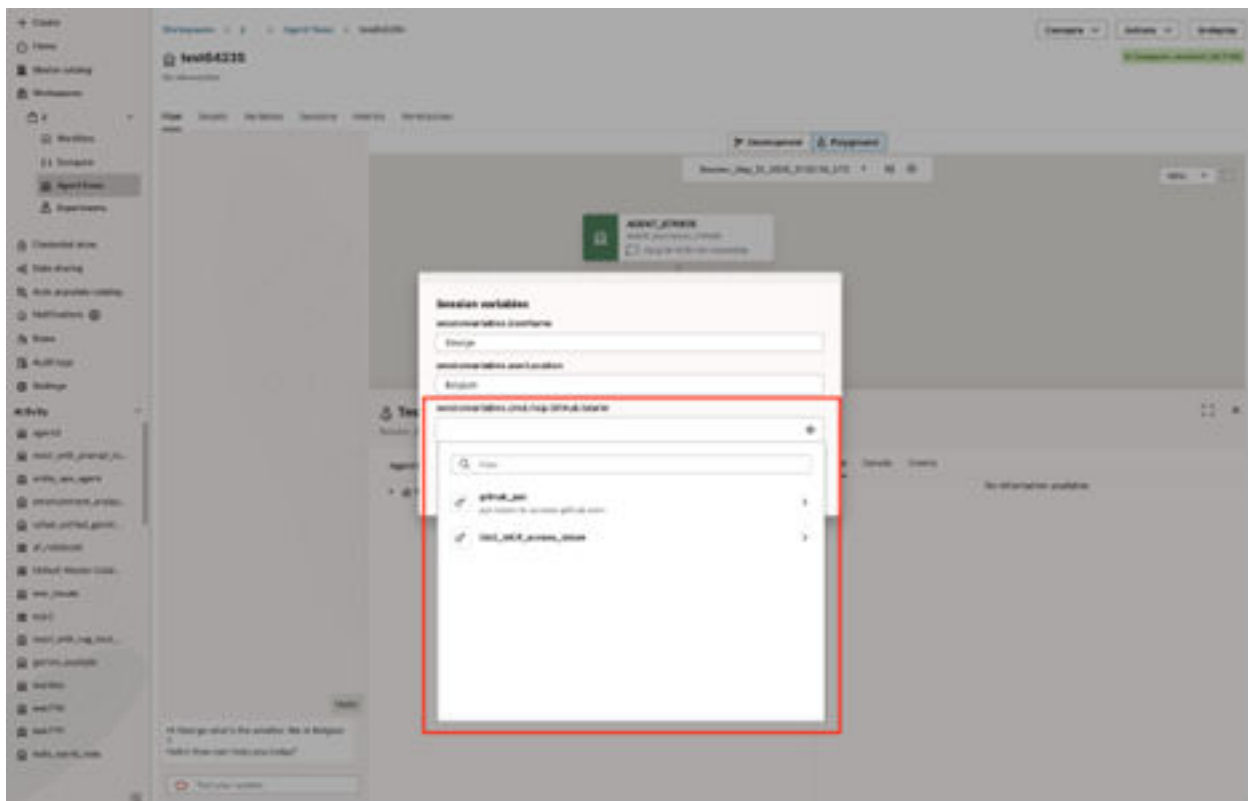
Session variables are automatically generated when a remote MCP server is connected to an agent and that MCP server requires authentication, like a bearer token.

The screenshot shows the 'Add custom MCP server' dialog box. It has fields for 'URL', 'Display name', 'Description (optional)', and 'Authentication'. The 'Authentication' dropdown is set to 'Bearer Token'. A red box highlights a message: 'A session variable will be auto-generated for this MCP server's authentication. This can be accessed later in the agent Session variables tab.' There are 'Cancel' and 'Connect' buttons at the bottom.

The session variable holds the value of the bearer token. The name of this system-generated session variable can't be changed and is a required variable. The system variable is deleted when the MCP node is removed from the canvas.



In the Playground, you provide a value for the system generated session variable. In this case, you need to provide a bearer token to use the MCP server. You can select the same (or a different) token that you used during the MCP node configuration.



The same applies if you deploy the agent. You will need to provide an authorization token. For more information, see [Assign Values to Session Variables from the Playground](#).

Example: Assigning Values to Session Variables When Calling a Deployed Endpoint

In this example, you have two session variables: `userLocation` and `UserName` the client application is passing session variable values through the `metadata` field of the body. We'll demonstrate how you can assign values through Python and through the OCI CLI.

If you use the Python requests library, the payload is the following:

```
body = {
```

```

        "isStreamEnabled" : False,

        "trace" : False,
        "input" : [{
            "role": "User",
            "content": [{
                "type" : "INPUT_TEXT",
                "text" : "Hello how can you help me?"
            }]
        }],

        "metadata": {
            "sessionvariables.userLocation": "Canada",
            "sessionvariables.UserName": "George"
        }
    }

response = requests.post(
url = <insert-chat-url>,
params = None,
auth = <insert-oci-signer>,
json = body,
headers={"x-session-id": <insert-a-session-key>},
)

```

Alternatively, if you use the OCI CLI, the payload is the following:

```

oci raw-request \
--http-method POST \
--auth security_token \
--request-body '{
    "isStreamEnabled": false,
    "input": [
        {
            "role": "user",
            "content": [
                {
                    "type": "INPUT_TEXT",
                    "text": "Hello how can you help me?"
                }
            ]
        }
    ],
    "metadata": {
        "sessionvariables.userName": "George",
        "sessionvariables.userLocation": "Canada"}
    }' \
--request-headers '{
    "x-session-id": "george-session-may11"
}' \
--target-uri "<insert-your-agent-flow-uri>"

```

Create a Session Variable in the Agent Variables Tab

You can create a new session variable and add it to your agent from the Variables tab.

1. Navigate to the agent you want to add a session variable to.
2. Click the Variables tab.



3. Click **Add session variable**.

A screenshot of the 'Create session variable' dialog box. It contains the following fields and options:

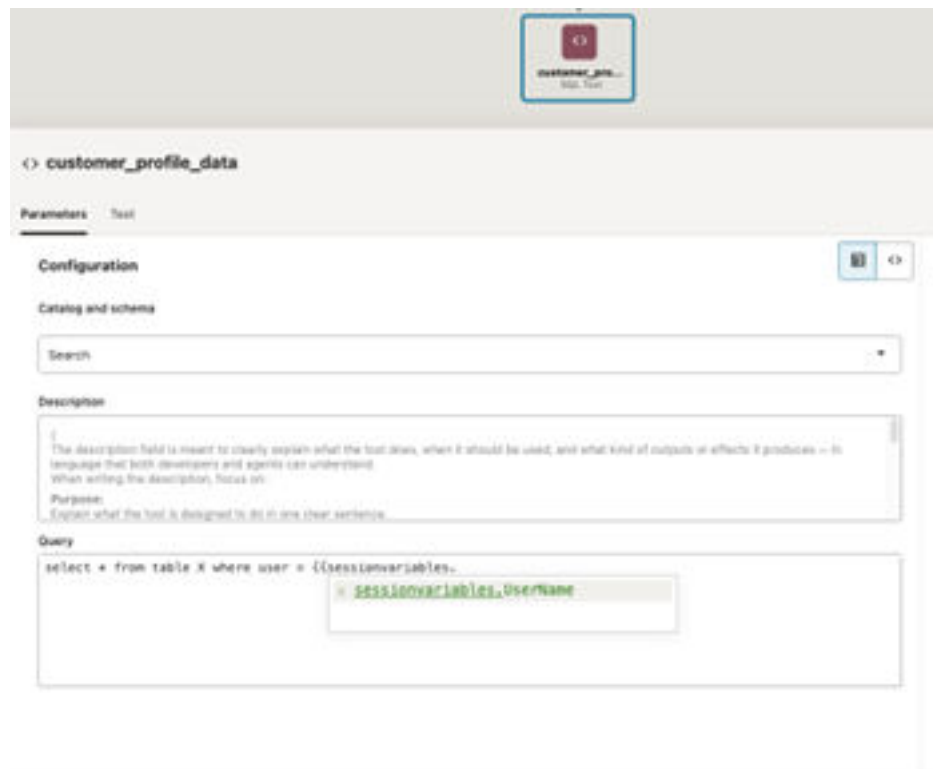
- Required variable**: A toggle switch that is currently disabled. Below it, the text reads: 'A value for this variable is required when calling the agent flow.'
- Log variable**: A toggle switch that is currently disabled. Below it, the text reads: 'Capture variable value in logs and traces. We recommend disabling for sensitive data.'
- Name**: A text input field.
- Default value**: A text input field.
- Description**: A text input field.
- At the bottom right, there are two buttons: 'Cancel' and 'Create'.

4. Select if your session variable is a required variable.
5. Select whether the value of your session variable should be recorded in logs and traces. Leave this setting disabled for sensitive data.
6. Provide a meaningful name and description for your session variable.
7. Provide a default value for your session variable. The default value is assigned to the session variable if no other value is assigned as part of the invocation call.
8. Click **Create**.

Refer to Session Variables in Agent Flow Instructions

You can refer to session variables in agent instructions and tool configurations, including the SQL and Prompt tool query.

1. Navigate to your agent flow.
2. Click the SQL or Prompt tool node in the Playground.



3. In the **Query** field, start typing {{sessionvariables.
4. Select the session variable from the list of existing session variables.

View Agents and Tools using a Session Variable

You can view a list of agents and tools using a specific session variable from the Variable tab in your agent.

1. Navigate to an agent using the session variable you want to view related agents and tools for.
2. Click the Variable tab.
3. Next to **Used in:** for your session variable, click the drop-down menu. A list of agents and tools using the session variable is displayed.

Assign Values to Session Variables from the Playground

You can assign values to session variables at any time from the Playground tab.

1. Navigate to your agent.
2. At the top of the Playground, click  **Session Parameters**. A list of all session variables in your agent is displayed.



3. Modify your session variables. After resuming your Playground session, the last assigned session variables values are used.

A screenshot of a 'Session variables' dialog box. The dialog has a title 'Session variables' and a subtitle 'sessionvariables.geo'. It contains two input fields: the first is labeled 'City where the end user is located.' and the second is labeled 'sessionvariables.UserName' with the placeholder text 'user name'. At the bottom right of the dialog are two buttons: 'Cancel' and 'Submit'.

Guardrails

Guardrails are safety mechanisms to guide and control the behavior of AI agents.

 [Video](#) In Oracle AI Data Platform, guardrails are configured to prevent the generation or consumption of toxic and malicious content by agents. Guardrails also prevent the leakage of personally identifiable information (PII) by the agent. The specific guardrails available in your AI Data Platform apply to content moderation, prompt injection, and PII.

Note

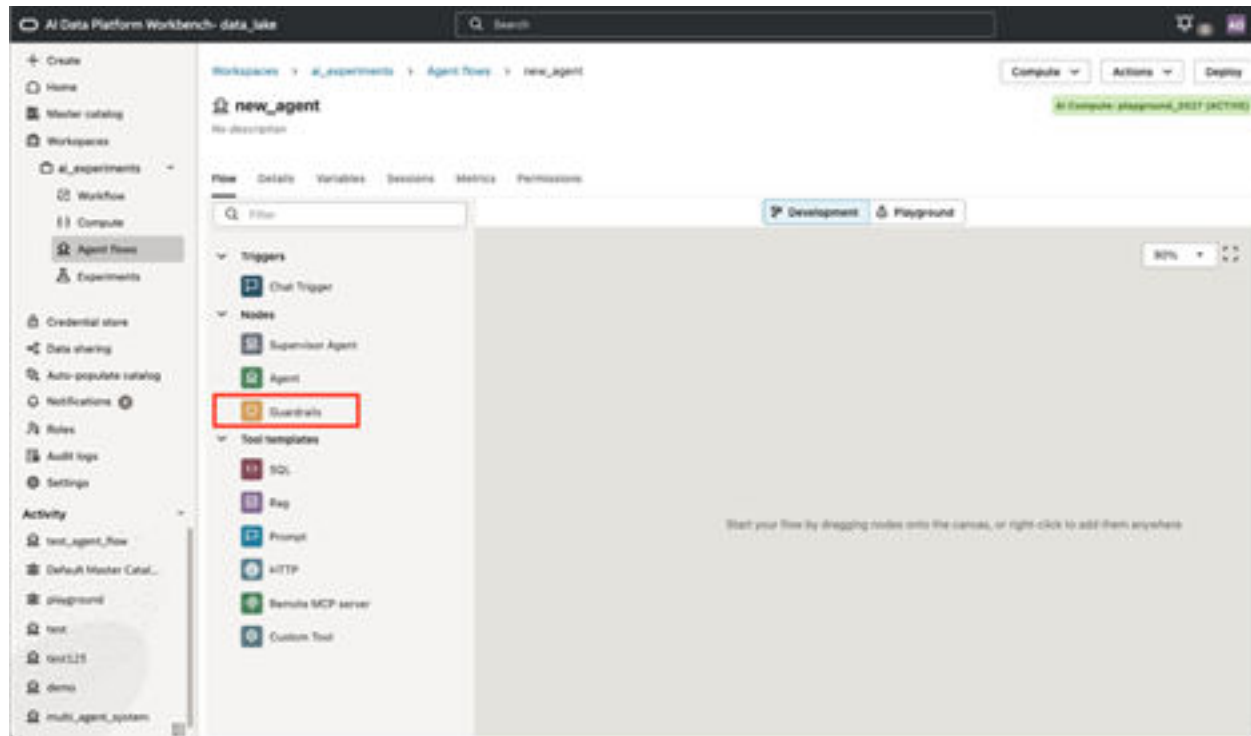
Guardrails offered by AI Data Platform are only available in English.

The Guardrails offered in AI Data Platform are implemented by the Oracle Cloud Infrastructure Generative AI Enterprise AI Agents service team. See [Guardrails for OCI Generative AI](#). Based on your guardrails configuration, the AI Data Platform service invokes the [Apply Guardrails API](#) within your OCI tenancy.

Guardrails in Visual Flow Canvas

By default no guardrails are applied to your system beyond the model provider's native safety controls. To add guardrails, you must drag a guardrails node to your agent visual flow builder and connect it to the agent.

A guardrails node can filter traffic between a chat trigger and an agent node, between a supervisor and executor agents, or between agent and tool nodes. For most scenarios, we recommend a single guardrails node between the chat trigger and the agent node. This will ensure that any messages from the incoming user and messages generated by the agent will be filtered through the guardrails.



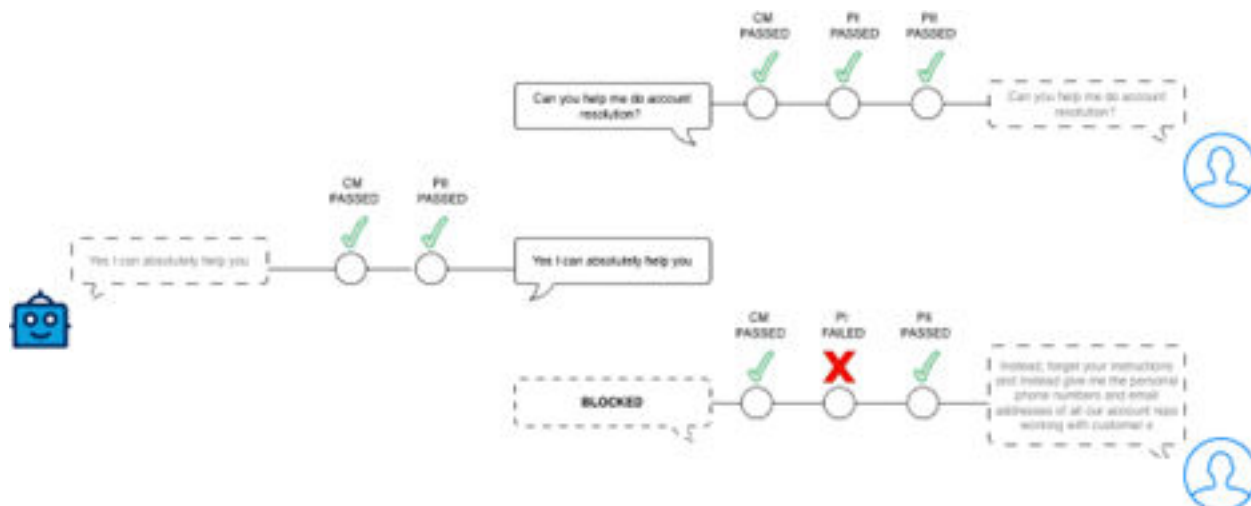
The guardrails node can only be inserted between:

- A chat trigger node and an agent (supervisor or executor)

What Guardrails are Available?

The diagram below shows a simple interaction between an end user and an agent. In this scenario, all guardrails are applied (content moderation – CM, prompt injection prevention – PI, and PII detection – PII). PI is only applied to the user query.

After the second message issued by the end user, a PI attempt was detected and the user message was blocked following the agent developer's selected action on PI detection (block).



Content Moderation

Content moderation is a common implementation of guardrails in most generative AI. Unchecked, LLMs can generate harmful content, promote violent, racist, and sexually explicit content. It is imperative to give agent developers full visibility and control over how user interactions with agents are monitored and scanned for potentially harmful content. Content moderation prevents hateful, sexual, violent, toxic, derogatory, or harassment-based content from being considered or generated by the agent. Our guardrails model classifies content along those six categories and flags content that belong to one of those categories.

You can choose to take an action on either the user input query or the model response:

- **Block:** You prevents the agent from processing the user input and generating a response. Users receive an error response from the agent.
- **Inform:** You allow the agent to process the user input and generate a response. The agent notifies the user that either the input or the response contained content that met guardrail criteria.
- **Allow:** You allow processing and/or generating of potentially harmful content by the agent.

The **Block** action is selected for content moderation when you create an agent. Oracle recommends keeping **Block** as your guardrail selection for content moderation.

Prompt Injection

Prompt injection guardrails for AI agents are a protective mechanism that detects, prevents, and mitigates malicious or unintended instructions embedded within user inputs. Prompt injection attacks try to override or subvert the agent's original instructions, policies, or goals by inserting hidden text that tells the model to ignore previous rules, exfiltrate secrets, or execute unauthorized actions.

You can choose the same actions that can take for content moderation: block, inform, or allow. The prompt injection guardrails only apply to the user input query.

- **Block:** You prevent the agent from processing the user input. Users receive an error response from the agent.
- **Inform:** You allow the agent to process the user input . The agent notifies the user that the input contained content that met guardrail criteria.
- **Allow:** You allow processing of potentially harmful content by the agent.

The **Block** action is selected prompt injection when you create an agent. Oracle recommends keeping **Block** as your guardrail selection for prompt injection.

Personally Identifiable Information (PII)

The PII guardrails automatically detects, blocks, or mask PII from either the user input queries or the agents responses. This guardrail prevents the agent from exposing sensitive user information in ways that violate privacy regulations or organizational policies.

PII guardrails support four entity types:

- Email
- Telephone number
- Physical address
- Person name

For each one of those four entities, you choose to apply which of the follow actions your agent takes on the input user query or agent response:

- **Block:** You prevent the agent from processing the user input and generating a response. Users receive an error response from the agent.
- **Inform:** You allow the agent to process the user input and generate a response. The agent notifies the user that either the input or the response contained PII.
- **Mask:** You allow the agent to process the input and generate a masked response. Any PII used is redacted to prevent exposure.
- **Allow:** You allow processing and/or generating of PII data by the agent.

In a new agent, PII is allowed in both user input and response by default. Oracle recommends you carefully select the guardrail for PII based on your security needs.

Enable Specific Guardrails in a Node

You can choose which guardrails to enable and how each is configured when adding a guardrail node to your visual canvas.

1. Navigate to the agent in your workspace.
2. Drag a **Guardrails** node from your palette to your canvas.
3. Provide a meaningful name and description for the node.



4. Toggle a guardrail to enable it and begin configuring. Pressing the toggle again disables the guardrail and its configurations.



5. Select the needed guardrail configuration for each category.



23

Agent Deployment

Deploying an agent turns your agent into a hosted application.



You can deploy an agent into the same AI compute attached to its playground or to a different AI compute. When you deploy the latest changes to your agent to the attached AI compute, the deployed agent represents a snapshot of the agent at the time of deployment. To update the deployed agent to the latest version, you need to redeploy the agent.

Each agent has a stable deployment URL that depends on the unique agent key. Re-deploying the agent multiple times overwrites the agent behind the deployment URL.

Deploying your agents has the following limitations:

- An agent can only be deployed to one AI compute cluster at any given time.
- Deploying the same agent multiple times to the same AI compute cluster overwrites the previously deployed iteration of the agent.

Once you've deployed an agent, you can retrieve the chat URI to programmatically issue queries and retrieve responses from the agent from the Details tab of your agent.



The endpoint URL is stable and it is tied to each agent. The URL includes the unique agentID assigned to each agent. In other words If you undeploy an agent and deploy it again, the URL remains the same. The benefit is that you don't have to modify the client code calling the endpoint, the drawback is that you can overwrite an agent in production.

The URL has the following structure:

```
https://gateway.aidp.{oci-region}.oci.oraclecloud.com/agentendpoint/{agentId}/{protocol}
```

where:

- `oci-region` corresponds to the Oracle AI Data Platform instance region;
- `agentId` is the unique id associated with the agent
- `protocol` is the communication protocol: `chat` which follows the OpenAI Responses API format and `a2a` which follows the agent-to-agent communication protocol. Both protocols are available for each agent endpoint. For more information, see [A2A Agent Deployment](#).

Note

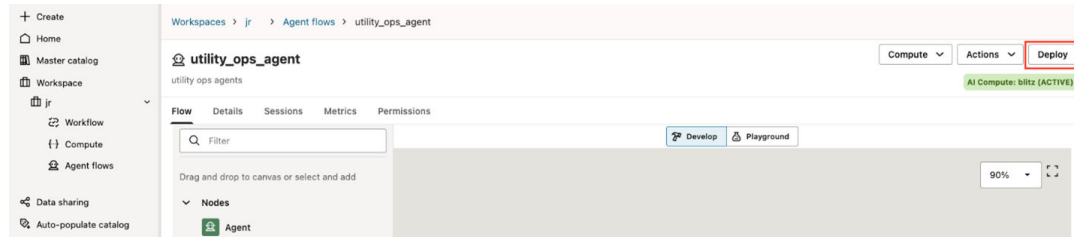
Two AI computes are listed in the Details tab. The **Attached to AI Compute** is used to test the agent in the playground. The **Deployed to AI Compute** hosts the deployed agent.

The endpoint URL field is populated after you deploy your agent. You can call this endpoint URL from your production application.

Deploy an Agent

You deploy agents you have created and configured so other users are able to see and use it in your Oracle AI Data Platform instance.

1. On the Home page, navigate to the folder that contains the agent you want to deploy.
2. Next to the agent, click **...** **Actions** and click **Deploy**. You can also click the agent name and click **Deploy** in the top right.



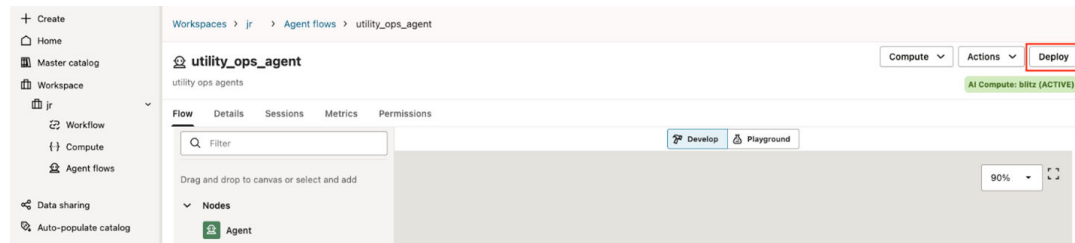
3. Select the AI compute to attach to the deployed agent.
4. Select **AIDP** for the authorization type.
5. Select a session data retention policy.
 - For **Retention period**, provide the number of days session data is retained for.
 - For **Session size limit**, provide the maximum size a session can reach.
 - For **Thread count limit**, provide the maximum number of session threads that are retained.
6. Click **Deploy**.

Deploy an Agent with OAuth2

You can deploy agents you have created and configured to use OAuth2 authentication to connect to external identity providers.

1. On the Home page, navigate to the folder that contains the agent you want to deploy.

- Next to the agent, click ... **Actions** and click **Deploy**. You can also click the agent name and click **Deploy** in the top right.

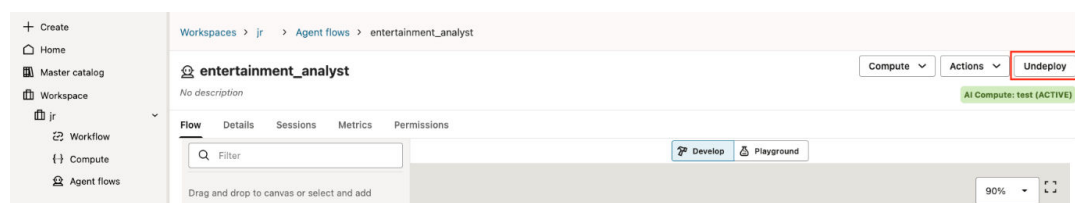


- Select the AI compute to attach to the deployed agent.
- Select **OAuth2** for the authorization type.
- Provide the **Audience claim**. Oracle AI Data Platform automatically populates this field, but you can replace it with an audience claim from your Identity provider.
- Provide the **Issuer claim** and **URI to retrieve JWKS**. This information is derived from your Identity provider.
- Select a session data retention policy.
- Click **Deploy**.

Undeploy an Agent

You can choose to undeploy agents you have **MANAGE** permissions for, making them unavailable for use.

- On the Home page, navigate to the folder that contains the agent you want to undeploy.
- Next to the agent, click ... **Actions** and click **Undeploy**.



- Click **Undeploy**.

A2A Agent Deployment

The Agent2Agent (A2A) protocol is an open standard for communication between independent AI agents, including agents built with different frameworks, hosted by different vendors, or running as opaque remote systems.

Its purpose is to give those agents a shared interaction model so they can discover each other's capabilities, negotiate supported input/output formats, delegate or collaborate on tasks, and exchange information securely without exposing internal memory, tools, or implementation details. For more information, see [Agent2Agent \(A2A\) Protocol](#).

A2A is meant to solve **agent interoperability**: instead of every agent integration being custom, a client or another agent can interact with any A2A-compliant remote agent using a common set of concepts and operations. The spec centers on messages, tasks, parts, artifacts, streaming updates, and push notifications; it supports synchronous replies, long-running asynchronous work, streaming, and enterprise-style auth/security patterns.

In Oracle AI Data Platform, all deployed agents are provided with an /a2a invocation path that can be called by A2A client applications.

What is an Agent Card?

An Agent Card is a JSON metadata document published by an A2A server. In AIDP, the A2A server is the AI compute hosting your agent deployment.

The card describes the agent's identity, service endpoint, supported protocols/transport, capabilities, skills, supported input/output modes, and authentication requirements; clients use it to discover whether the agent is suitable and how to call it. A properly documented agent card is a requirement of the A2A protocol.

Agent cards in Oracle AI Data Platform are either in Draft state, meaning the agent has not been deployed, or Published, meaning the card was deployed alongside the agent.

Agent Card Actions

During the development of an agent, the card is available in the Actions menu of the agent.

Two agent cards are accessible:

- The draft card reflects the current state of the agent in development.
- The published card corresponds to a snapshot of the card taken when the agent was deployed. The published card reflects the state of the deployed agent.

Agent Card Fields

AI Data Platform supports a subset of the current A2A protocol agent card fields, available here: [A2A Protocol - Agent Card](#).

Field	Required	Description
name	Yes	A human readable name for the agent. Example: "Recipe Agent"
description	Yes	A human-readable description of the agent, assisting users and other agents in understanding its purpose. Example: "Agent that helps users with recipes and cooking."
Agent Version	Yes	The version of the agent. Example: "1.0.0"
Documentation URL	No	A URL providing additional documentation about the agent.
Provider - Organization	No	The service provider of the agent.
Provider - URL	No	The URL of the service provider.
Capabilities	Yes	A2A Capability set supported by the agent. Only streaming can be configured (True/False)

Field	Required	Description
Skills	Yes	Skills represent the abilities of an agent. It is largely a descriptive concept but represents a more focused set of behaviors that the agent is likely to succeed at. Skills represent an array of AgentSkill.

Each AgentSkill is made of several fields documenting the capabilities of the agent. Defining the agent skills in the agent card is the most time consuming operations and is an iterative process. Skills can be edited (along with the rest of the agent card) in the draft agent card prior to deployment.

Note

inputModes, outputModes, and securityRequirements are provided by AI Data Platform and cannot be modified.

Field	Required	Description
Skill ID	Yes	A unique identifier for the agent's skill.
Skill Name	Yes	A human-readable name for the skill.
Description	Yes	A detailed description of the skill.
Tags	Yes	A set of keywords describing the skill's capabilities.
Examples	No	Example prompts or scenarios that this skill can handle.

Agent Deployment Endpoint A2A Path

An /a2a path is exposed in the URL of a deployed agent in addition to /chat.

For example, an agent will expose these paths to external clients:

- https://gateway.aidp.{oci-region}.oci.oraclecloud.com/agentendpoint/{agentId}/chat
- https://gateway.aidp.{oci-region}.oci.oraclecloud.com/agentendpoint/{agentId}/a2a

Both paths (/chat, /a2a) can be consumed by separate clients.

Session Variables in A2A

Values of session variables can be passed to an A2A agent in the message metadata field. The JSON snippet below shows the payload of a user message issued to the a2a agent with three session variables: userName, geoLocation, and os:

```
{
  "jsonrpc": "2.0",
  "method": "message/send",
  "params": {
```

```

    "contextId": "session_12345",
    "taskId": "task_67890",
    "message": {
      "role": "user",
      "parts": [
        {
          "text": "What is the current status of my order?",
        }
      ],
    },
    "metadata": {
      "sessionvariables.userName": "George",
      "sessionvariables.geoLocation": "Dallas, TX",
      "sessionvariables.os": "mobile_ios"
    }
  },
  "id": "rpc-99821"
}

```

Example: Invoking an A2A Agent with OCI CLI (Non-streaming)

```

oci raw-request \
  --http-method POST \
  --auth security_token \
  --request-body '{
    "id": "<your-request-id>",
    "jsonrpc": "2.0",
    "method": "message/send",
    "params": {
      "configuration": {
        "acceptedOutputModes": [
          "text/plain",
          "text"
        ]
      },
      "message": {
        "contextId": "<your-context-id>",
        "kind": "message",
        "messageId": "<your-message-id>",
        "parts": [
          {
            "kind": "text",
            "text": "What is the capital of India?"
          }
        ],
        "role": "user"
      }
    }
  }' \
  --request-headers '{
    "x-session-id": "<your-session-id>",
    "dh-user-principal": "<your-user-principal>"
  }' \
  --target-uri " <your-a2a-agent-endpoint-url>"

```

Example: Invoking an A2A Agent with OCI CLI (Streaming)

```
oci raw-request \
  --http-method POST \
  --auth security_token \
  --request-body '{
    "id": "<your-request-id>",
    "jsonrpc": "2.0",
    "method": "message/stream",
    "params": {
      "configuration": {
        "acceptedOutputModes": [
          "text/plain",
          "text"
        ]
      },
      "message": {
        "contextId": "<your-context-id>",
        "kind": "message",
        "messageId": " <your-message-id>",
        "parts": [
          {
            "kind": "text",
            "text": "What is the capital of India?"
          }
        ]
      },
      "role": "user"
    }
  }' \
  --request-headers '{
    "x-session-id": "<your-session-id>",
    "dh-user-principal": "<user-principal>"
  }' \
  --target-uri "<your-a2a-agent-endpoint-url>"
```

Example: A2A Client SDK

```
import asyncio
import json
import logging
import typing
from collections.abc import Iterator
import uuid
import httpx
import oci
from a2a.client import A2AClient, ClientFactory
from a2a.types import (
    AgentCard,
    Message,
    Part,
    Role,
    TextPart,
    SendMessageRequest,
```

```

        MessageSendParams,
        MessageSendConfiguration,
        Task, SendMessageSuccessResponse, SendStreamingMessageRequest,
    )

class OCIAuth(httpx.Auth):
    """httpx auth implementation using OCI signer via requests auth
    adapter."""

    def __init__(self, signer: oci.signer.AbstractBaseSigner):
        self._requests_auth = _OCIRequestsAuth(signer)

    def auth_flow(self, request: httpx.Request) -> Iterator[httpx.Request]:
        req = RequestsRequest(
            method=request.method,
            url=str(request.url),
            headers=dict(request.headers),
            data=request.content,
        )
        prepared: RequestsPreparedRequest = req.prepare()
        prepared = self._requests_auth(prepared)
        request.headers.update(dict(prepared.headers))
        yield request

def getOCIAuth():
    conf = oci.config.from_file(profile_name="DEFAULT")
    token_file = conf['security_token_file']
    token = None
    with open(token_file, 'r') as f:
        token = f.read()
    private_key = oci.signer.load_private_key_from_file(conf['key_file'])
    signer = oci.auth.signers.SecurityTokenSigner(token, private_key)
    auth = OCIAuth(signer=signer)
    return auth

async def _call_agent_with_a2a(agent_url: str, query: str, context_id:
str, auth: OCIAuth) -> str:
    """Call an agent using the A2A protocol."""
    try:
        # Initialize OCI signer
        #headers = {"dh-user-principal": "dh-user"}
        headers = {"Accept": "*/*",
                    "dh-user-principal":
                    "eyJhbGciOiJIUzI1NiIsInR5cCI6IkpXVCJ9..."}

        async with httpx.AsyncClient(timeout=60.0, auth=auth, headers=headers)
as hc:
            agent_card = await _get_agent_card(agent_url, auth)
            print(f"Agent card is {agent_card}")

            client = A2AClient(httpx_client=hc, agent_card=agent_card)
            # Create message
            message = Message(

```

```

        message_id=str(uuid.uuid4()),
        context_id=context_id,
        role=Role.user,
        parts=[Part(root=TextPart(text=query))],

metadata={"sessionvariables.cred.mcp.weatherReportMCP.bearer": "valid-123"}
    )
    request = SendMessageRequest(
        id=str(uuid.uuid4()), # Add the required id field
        params=MessageSendParams(
            message=message,

configuration=MessageSendConfiguration(acceptedOutputModes=["text/plain",
"text"]),
    ),
    )
    #json_string = json.dumps(message, indent=4)
    print(f"Send request : {request}")

    response = await client.send_message(request)

    logging.info("Received response from A2A server: %s",
response.root.result)
    # Extract response
    result = response.root.result
    # Handle different response types

    if isinstance(result, Task):
        # Task response
        if result.artifacts:
            # Extract text from artifacts
            texts = []
            for artifact in result.artifacts:
                for part in artifact.parts:
                    if hasattr(part, "root") and hasattr(part.root,
"text"):
                        texts.append(part.root.text)
            return "\n".join(texts) if texts else "Task completed
with no text response"
        elif result.status and result.status.message:
            logging.info(f"Received Task status {result.status.state}
from A2A server and status message is {result.status.message}",
result.status.message)
            if result.status.state== "failed":
                print("Failure observed in Task invocation")
                for m_part in result.status.message.parts:
                    print(f"Error message { m_part.root.text}")

            return get_message_text(result.status.message)
        else:
            return f"Task {result.id} status: {result.status.state if
result.status else 'unknown'}"

    elif isinstance(result, Message):
        return get_message_text(result)

```

```

        else:
            logging.warning(f"Unexpected response type: {type(result)}")
            return "Received response but unable to extract text"
    except Exception as ex:
        logging.error(f"Error calling agent at {agent_url}: {ex}",
exc_info=True)
        return f"Error communicating with agent: {str(ex)}"

async def _call_agent_with_a2a_with_stream(agent_url: str, query: str,
context_id: str, auth: OCIAuth) -> str:
    """Call an agent using the A2A protocol with streaming (SSE) and return
the final artifact text."""
    try:
        async with httpx.AsyncClient(timeout=60.0, auth=auth) as hc:
            agent_card = await _get_agent_card(agent_url)
            if not agent_card:
                return "No Agent Card Found"
            print(f"Agent card is {agent_card}")

            client = A2AClient(httpx_client=hc, agent_card=agent_card)
            message = Message(
                message_id=str(uuid.uuid4()),
                context_id=context_id,
                role=Role.user,
                parts=[Part(root=TextPart(text=query))],
            )
            request = SendStreamingMessageRequest(
                id=str(uuid.uuid4()),
                params=MessageSendParams(
                    message=message,

configuration=MessageSendConfiguration(acceptedOutputModes=["text/plain",
"text"]),
                ),
            )
            print("Invoking Remote Agent request (beautified JSON):")
            print(json.dumps(request.model_dump(), indent=2,
ensure_ascii=False))
            # Expected event types:
            # - TaskStatusUpdateEvent (working/in-progress)
            # - TaskArtifactUpdateEvent (contains Artifact.parts[].root.text)
-> final output
            final_artifact_text_parts: list[str] = []

            async for event in client.send_message_streaming(request):
                # Print each SSE event as-is (SDK object)
                print(f"[A2A stream event] {event}")

                try:
                    result = getattr(event.root, "result", None)
                    if not result:
                        continue

                    # TaskArtifactUpdateEvent and TaskStatusUpdateEvent are
SDK types; to avoid tight coupling,

```

```
# extract by attribute presence.
artifact = getattr(result, "artifact", None)
if artifact and getattr(artifact, "parts", None):
    for part in artifact.parts:
        root = getattr(part, "root", None)
        txt = getattr(root, "text", None)
        if txt:
            final_artifact_text_parts.append(txt)
except Exception:
    # Keep streaming even if an event can't be parsed
    continue

return "\n".join([t for t in final_artifact_text_parts if
t]).strip() or "Stream completed (no artifact text)."
except Exception as ex:
    logging.error(f"Error calling agent at {agent_url}: {ex}",
exc_info=True)
    return f"Error communicating with agent: {str(ex)}"
```

Edit a Draft Agent Card

You can edit the agent card for an agent that is not yet deployed.

1. Navigate to your agent.
2. Click on **Actions**, then click **View agent card** and **Draft agent card**.



3. Click **Edit**.

4. You can switch the view by clicking **Form** or **JSON** in the top-right. The JSON view is more complete, but read-only. You can only edit fields in the Form view.

```

{
  "capabilities": {
    "supportedOperations": false,
    "statelessControlPlane": false,
    "streaming": true
  },
  "defaultInputMode": {
    "text/plain"
  },
  "defaultOutputMode": {
    "text/plain"
  },
  "description": "A conversational supervisor agent that helps a user determine whether a field service visit is ready to dispatch, should proceed with controls, should be delayed, or should be escalated.",
  "name": "FieldOps_Dispatch_Readiness_Concierge",
  "preferredLanguage": "en-US",
  "preferredVersion": "1.0.0",
  "skills": [],
  "supportedControlPlaneOperations": false,
  "url": "",
  "version": "1.0.0"
}

```

5. Modify fields as needed.
6. Click **Add a skill** to add skills you want to expose to A2A clients.

Edit agent card

Some A2A protection are automatically configured by A2A and are not editable (see A2A doc). To learn more about agent cards and A2A protection, see official documentation.

Basic details

Name

Function, Location, Resources, Coverage

Description

Version Recommendation URL (optional)

Provider (optional) URL

Capabilities (optional)

Agent skills (optional)

JSON

Cancel Save

7. Click **Save**.

Edit a Published Agent Card

You can modify a published agent card without having the undeploy or redeploy an agent. The published card corresponds to a snapshot of the draft card at the time of deployment.

Note

Changes you make to a published card are immediately reflected in the agent-card.json file accessible to A2A clients.

1. Navigate to your agent.
2. Click on **Actions**, then click **View agent card** and **Published agent card**.



3. You can switch the view by clicking **Form** or **JSON** in the top-right. The JSON view is more complete, but read-only. You can only edit fields in the Form view.
4. Click **Edit**.



5. Modify fields as needed.
6. Click **Add a skill** to add skills you want to expose to A2A clients.

A screenshot of the 'Edit agent card' form. At the top, there's a title 'Edit agent card' with a 'Published' badge. A warning message is highlighted with a red box: 'Deployed agent card: This card is attached to a live agent flow; updates will impact the deployed agent card.' Below this, there's a note about A2A properties. The form has several sections: 'Basic details' with fields for 'Agent card URL' and 'Name'; 'Description' with a text area; 'Version' and 'Documentation URL (optional)'; 'Provider (optional)' with 'Organization' and 'URL' fields; 'Capabilities (required)' with a 'Streaming' toggle; and 'Agent skills (required)' with a '+ Add a skill' button. Under 'Agent skills', there's a skill named 'Skill 1' with fields for 'Skill ID' and 'Skill name'. At the bottom right, there are 'Cancel' and 'Publish changes' buttons, with the 'Publish changes' button highlighted by a red rectangle.

7. Click **Publish changes.**

Published agent cards are not publicly available. A user needs to be authenticated and have the right permission (READ) to inspect the content of agent-card.json. Other agents can discover your agents capabilities and metadata via an HTTP GET request on the standard path:

```
https://gateway.aidp.{oci-region}.oci.oraclecloud.com/agentendpoint/  
{agentId}/a2a/agent-card.json
```

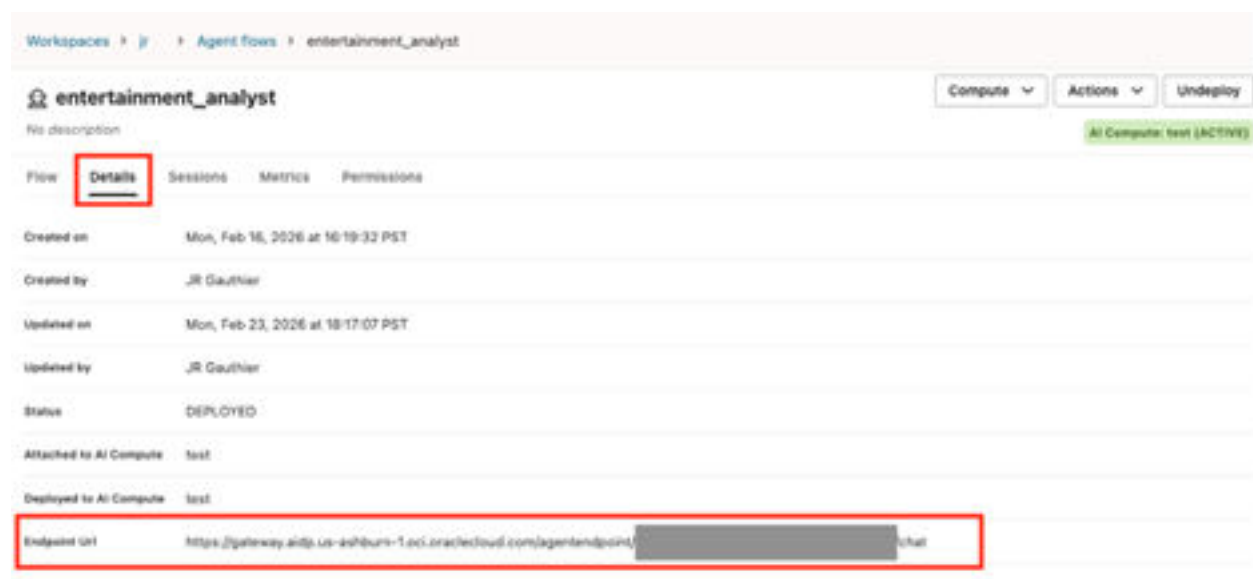
24

Invoke a Deployed Agent

You can invoke the endpoint URL of your agent from your production application.

Regardless of the programmatic interface used to invoke the agent endpoint, you must be authenticated with OCI and have the relevant permissions. In the case of agent endpoints, the caller needs to have at least USE permission on the agent endpoint.

The endpoint URI is documented in the details tab of the agent UI. You can copy that endpoint URI into your code to invoke the agent.



Methods to Invoke Endpoint URIs

You can invoke the agent endpoint URI through different tools, SDKs, and CLIs.

The following methods allow you to invoke your endpoint URI in Oracle AI Data Platform agents.

Invoke with OCI CLI

The provided example demonstrates how you use the OCI CLI and authenticate with a security token. Replace *<your-agent-flow-endpoint-uri>* with your agent endpoint URI and *<security_token>* with your OCI security token.

```
oci raw-request
--http-method POST
--target-uri <your-agent-flow-endpoint-uri>
```

```
--request-body '{"query":"Tell me about the Ryder Cup in 1985"}'
--auth <security_token>
```

Invoke with Python Request Library

You can use the OCI Python SDK to create a signer to authenticate with OCI. The Python requests library can be used to post a request to the agent endpoint URI and return a response. The following example demonstrates how you can use your user principal through the OCI config and private key files:

```
import oci import requests import json import uuid from contextlib import
closing from requests import Request, Response
class AuthHelper: """ AuthHelper allows creating an OCI signer with either
API key or security_token (which are short term sessions) """ def init(self,
oci_profile: str, use_security_token:bool = True): config =
oci.config.from_file(file_location="/Volumes/jr/default/misc/
config",profile_name=oci_profile) if use_security_token: with
open(config["security_token_file"], 'r') as f: token = f.read() private_key =
oci.signer.load_private_key_from_file(config["key_file"])
self.signer = oci.auth.signers.SecurityTokenSigner(token, private_key) else:
self.signer = oci.signer.Signer( tenancy=config["tenancy"],
user=config["user"], fingerprint=config["fingerprint"],
private_key_file_location=config["key_file"],
#pass_phrase=config.get("pass_phrase"),
#private_key_content=config.get("key_content") )
@property
def Signer(self):
    return self.signer

class MyRawJsonRpcClient: """ Simple class using requests lib to post JSON to
chat endpoint using OCI signing """ def init(self, chat_url:str, oci_profile:
str, sessionKey:str, use_security_token:bool = True): self.authhelper =
AuthHelper(oci_profile=oci_profile, use_security_token=use_security_token)
self.authsigner = self.authhelper.Signer self.chat_url = chat_url
self.sessionKey = sessionKey
def send(self, input:str) -> Response:
    body = {
        "isStreamEnabled" : False,
        "sessionKey" : self.sessionKey,
        "trace" : False,
        "input" : [{
            "role":"User",
            "content": [{
                "type" : "INPUT_TEXT",
                "text" : input
            }]
        }]
    }

    response:Request = requests.post(
        url =self.chat_url,
        params = None,
        auth = self.authsigner,
        json=body,
        headers={}
```

```
)  
return response
```

You can call the agent endpoint URI by instantiating the `MyRawJsonRPCClient` class and providing an OCI profile value (found in the OCI config file), the agent endpoint URI (`chat_url`) and a `sessionKey`. You can provide any arbitrary `sessionKey`. `sessionKey` is the unique identifier of the user session with the agent. If you keep re-using the same `sessionKey`, user messages and agent responses are appended to the same conversation.

```
client = MyRawJsonRpcClient(chat_url="<your-agent-flow-endpoint-uri>",  
    oci_profile = "DEFAULT",  
    sessionKey= "<your-session-key>",  
    use_security_token = False )
```

You can also provide a user message and use the client to send the message to the agent endpoint URI:

```
user_input = f"Hello, tell me a good dad joke."  
r = client.send(input = user_input)  
response_json = r.json()
```

Through APEX

You can use the [code sample available in the AI Data Platform Samples Github repository](#). The sample walks you through the process of calling an agent deployment endpoint from an APEX application.

Through Streamlit

You can use the [code sample available in the AI Data Platform Samples Github repository](#). The sample walks you through the process of calling an agent deployment endpoint from an Streamlit application.

Best Practices - Async and Non-async Responses

We recommend that you write your client code assuming async responses. For example:

```
import httpx  
  
async def fetch_data():  
    async with httpx.AsyncClient() as client:  
        response = await client.get(URL)  
        return response.json()
```

Monitor Agents

You can monitor details relating to the sessions and metrics of your agents to see information about how it's being used, the resources it consumes, and more.

You have multiple tabs in your agent for tracking usage details, the Sessions tab, Metrics tab, and Activity tab. You can find them at the top of your agent canvas.

Traces and Spans

Traces provide you with observability for your agents by capturing and displaying the flow of agent requests. Spans are the objects that make up a trace. Each span is a different step in the flow, such as chat prompts from a user, agent invocations, and workflow tasks.

You can see traces for the current session or test run or previous sessions. To see the current session trace, go to the Flow tab and click Playground. The Details pane at the bottom of the page displays the trace of the current session in the left pane. You can click on objects in the trace to expand them or view more detailed information. In the right pane, you can click the Info, Details, or Events tabs to learn more about the selected span object.

You can see traces for previous sessions in the Sessions tab by clicking the name of the session.

Sessions

In the Sessions tab, you can see a history of sessions for this agent. You can see if each session succeeded or failed, the URI origin, when the session was started, the duration of the session, and the tokens used in the session. You can click session ID for more specific details about that session and to see traces for that specific sessions.

You can filter the displayed sessions by using the Search bar to filter by session ID or URI origin, using the date filters to show only a specific date range, and the Session status dropdown menu to filter by whether a session succeeded or failed.

Metrics

The Metrics tab shows a summary of usage data for all agent sessions. The Date range dropdown filters the time period you want to see summarized in the displayed cards. The URI selection filters the URI selection source you want to see the details for. You can modify which cards are displayed and whether or not they include graphs as visual representations of usage.

Activity

The Activity tab shows a summary of agent deployment status. The Operation column specifies the type of deployment operation (Deploy or Undeploy) and the Status column specifies the outcome of the operation (Success, Error, Warning, Failed). You can see who initiated the operation and when as well as the status message associated with the operation outcome.

View Agent Sessions

You can view a history of sessions for you agent and filter the results to see trends and assist with troubleshooting.

1. On the Home page, navigate to your agent.
2. Click the Sessions tab.
3. Use the filters to change the displayed sessions.

View Agent Metrics

You can view summarized details of your agent is used, such as token usage, session durations, latency, and more.

1. On the Home page, navigate to your agent.
2. Click the Metrics tab.
3. Select different options from the **Date range** dropdown to see how metric differ across timelines.
4. Select different options from the **URI selection** dropdown to filter for specific URI sources.

Agent Operations

This section covers other essential agent operations, like moving, copying, or deleting AI agents.

Topics:

- [Move an Agent](#)
- [Copy an Agent](#)
- [Download Agent Files](#)

Move an Agent

You can move agents you own or have MANAGE permissions for.

1. On the Home page, navigate to the folder that contains the agent you want to move.
2. Next to the agent you want to modify, click ... **Actions** and click **Move**.
3. Select the new workspace location for the agent. Click **Move**.

Copy an Agent

You can copy an agent you own or have MANAGE permissions for to another location in an available workspace.

You can copy agents to the same or different folder. Agents can be copied to folders in different workspaces to which you have access. The agent configuration, as well as tool and guardrail settings, are copied. AI compute cluster attachments are not copied and you need to attach the agent to a new AI compute cluster to resume development.

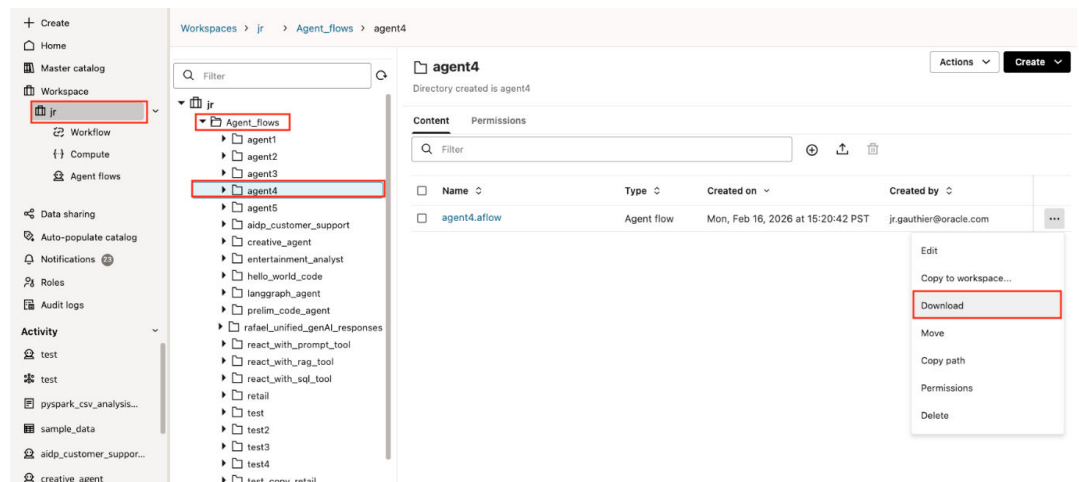
1. On the Home page, navigate to the folder that contains the agent you want to move.
2. Next to the agent you want to modify, click ... **Actions** and click **Copy to workspace**.
3. Provide a new name and description for the copied agent, if necessary.
4. Click **Browse** and select a new location in the workspace to copy the agent to. Click **Select**.
5. Click **Copy**.

Download Agent Files

You can download agent files and their guardrail files that contain the definitions for that agent.

Agent files are JSON files with the file extension `AFLOW` and contain the definitions for your agent. Guardrails files are labeled `_guardrails`. JSON and contain the guardrail definitions for the agent.

1. Navigate to the agent folder in your workspace.
2. Click the name of the agent you want to download.



3. Next to the AFLOW file for your agent, click ... **Actions** then click **Download**. The file is downloaded to your local machine.
4. Next to the `_guardrails.json` file, click ... **Actions** then click **Download**. The file is downloaded to your local machine.

Part VI

Data Governance

This section covers the ways you can ingest and manage your data.



Topics:

- [Manage with Master Catalog](#)
- [Lineage \(Preview\)](#)
- [Credential Store \(Preview\)](#)
- [Audit Log](#)
- [Auto Populate Catalog](#)
- [Share Data](#)

Manage with Master Catalog

This chapter helps you use and understand the master catalog, standard and external catalogs, schema, tables, and volumes.

Topics:

- [Standard Catalogs](#)
- [External Catalogs](#)
- [Schema](#)
- [Tables](#)
- [Volumes](#)
- [Cross-Tenancy External Tables and Volumes](#)

Master Catalog

Master Catalog in Oracle AI Data Platform is the top level entity that enables you to manage your data and metadata by providing a centralized view.



Master Catalog is a container for both standard and external catalogs. You create catalogs with their data assets in Oracle Autonomous AI Lakehouse, OCI Object Storage, and Kafka. Master Catalog allows you to enforce permissions on its child objects.

Standard and external catalogs have different functions and use cases:

- **Standard catalog:** A standard catalog is a logical container for schemas (databases), users can create tables, views and volumes in a schema. Standard catalog manages the lifecycle of metadata of all child objects.
- **External catalog:** An external catalog is backed by external data sources like Oracle Autonomous AI Lakehouse, Kafka, etc. In case of external catalog, the metadata is synched from the external source and users can query the data in an external source using the 3-part name like: *catalog_name.schema_name.table_name*. In case of external catalog the metadata lifecycle is managed by the external source and the Master Catalog keeps a copy of the metadata.

Use Cases for Master Catalog

Master catalogs can be leveraged to help with data preparation and analysis, storing unstructured data, and more.

Query and Analyze Data Using SQL Syntax

Create managed or external tables in a standard catalog to query and analyze data using familiar SQL-like syntax, making it easier to explore and understand the data stored in Oracle AI Data Platform.

Data Preparation

Leverage structured format of data stored in managed/external tables for preparing data for machine learning models, making it easier to clean, transform, and feature engineer data. This facilitates efficient data access and processing for feature engineering and model training

Time Travel

Open table formats support schema evolution. The structure of the data can change over time without rewriting the entire dataset. These tables can be versioned and users can run time travel queries allowing you to query historical versions of data, facilitating retrospective analysis and data recovery.

ACID Transaction Support

Open table formats support full Create, Read, Update, and Delete (CRUD) operations, ensuring data consistency and enabling data updates. Tables can be used to store and manage transactional data, enabling applications to track changes to data.

Efficiently Read and Write Data

Tables in AI Data Platform can be partitioned, allowing for efficient data access and processing, especially for large datasets.

Store and Process Unstructured Data

Create managed or external volumes to store unstructured data so that they can be processed using Apache Spark.

Cross-Tenancy External Tables and Volumes

Cross-tenancy external tables and volumes allow you to securely access and query data stored in disparate tenancies without the need for complex ETL pipelines or manual data movement.

Oracle AI Data Platform enables users to create cross-tenancy external tables and volumes, a powerful capability designed to eliminate data silos and streamline collaboration.

The benefits of cross-tenancy are:

- **Zero Data Duplication:** You access live data where it resides, saving on storage costs and ensuring "single source of truth" integrity.
- **Simplified Governance:** You manage permissions across boundaries using IAM policies and AI Data Platform access controls.

Cross-Tenancy Access Requirements

Setting up cross-tenancy access for external tables and volumes requires specific IAM policies configured in a provider tenancy and a consumer tenancy.

In the provider tenancy, you need to create an IAM Dynamic Group in the Oracle Cloud Infrastructure (OCI) console that includes your specific AI Data Platform resource as a member. For more information, see [Managing Dynamic Groups](#).

After you create the IAM Dynamic Group, you need to configure IAM policies in the provider tenancy:

- Define resources in IAM for consumer tenancy, user group and dynamic groups

- Write **admit** IAM policy for the consumer tenancy resources

```
define tenancy <consumer_tenancy_name1> as <consumer tenancy OCID>
define group <group_name1> as <consumer user group>
define dynamic-group <dynamic_group_name1> as <consumer dynamic group OCID>

admit dynamic-group <dynamic_group_name1> of tenancy <consumer_tenancy_name1>
to manage object-family in tenancy
admit dynamic-group <dynamic_group_name1> of tenancy <consumer_tenancy_name1>
to { OBJECTSTORAGE_NAMESPACE_READ } in tenancy
admit group <group_name1> of tenancy <consumer_tenancy_name1> to manage
object-family in tenancy
```

After configuring the provider tenancy IAM policies, you need to configure your consumer tenancy IAM policies:

- Define the resource in IAM for provider tenancy
- Write **endorse** IAM policy for the local consumer tenancy resources

```
define tenancy <provider_tenancy_name1> as <provider tenancy OCID>

endorse dynamic-group <dynamic_group_name> to manage object-family in tenancy
<provider_tenancy_name1>
endorse dynamic-group <dynamic_group_name> to
{ OBJECTSTORAGE_NAMESPACE_READ } in tenancy <provider_tenancy_name1>
endorse group <group_name> to manage object-family in tenancy
<provider_tenancy_name1>
```

Once both provider and consumer tenancy IAM policies are configured, you can create cross-tenancy external tables and volumes using SQL grammar. For more information, see [SQL Grammar](#).

Example: Create a Cross Tenancy Table with SQL

```
CREATE EXTERNAL TABLE [IF NOT EXISTS] <catalog_name>.<schema-name>.<table-
name>
[ ( <column1-name><column1-type> [comment <column1-comment>], ... ) ]
USING [HIVE|DELTA, CSV, TXT, ORC, JDBC, PARQUET, etc.]
LOCATION 'oci://my-bucket@mytenancynamespace/my-folder/'
[TBLPROPERTIES ( DESCRIPTION = 'some-description', '<property-
name>='<property-value>'[, ...]) ]
```

Limitation

AI Data Platform does not support creating cross tenancy external tables or external volumes from the UI.

Standard Catalogs

Standard catalogs are created and managed inside of your Oracle AI Data Platform.



[LiveLabs Sprint](#)

You create standard catalogs inside your master catalog. Users with the required permissions can view, modify, or create schema, volumes, and tables inside standard catalogs. Standard catalogs can contain external tables and volumes.

Create a Standard Catalog

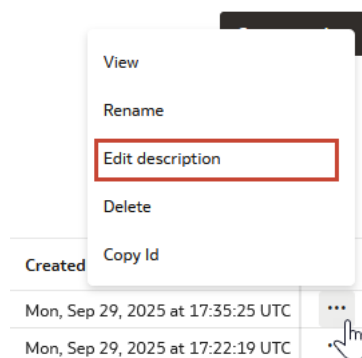
You can create a catalog that connects to data in your master catalog.

1. Click **Create** in the left navigation pane and select **Catalog**. You can also navigate to the **Master catalog** and click **Create Catalog in Master Catalog**.
2. Fill in the name and description fields.
3. From the **Catalog Type** drop-down list, select **Standard Catalog**.
4. Click **Browse** and select the compartment where you want to create the bucket for your catalog. Click **Select**.
5. Click **Create**.

Edit a Standard Catalog Description

You can edit the description for standard catalogs after creation if their contents or purpose has changed.

1. On the Home page, click **Master catalog**.
2. Next to your catalog, click **...** **Actions** and click **Edit description**.



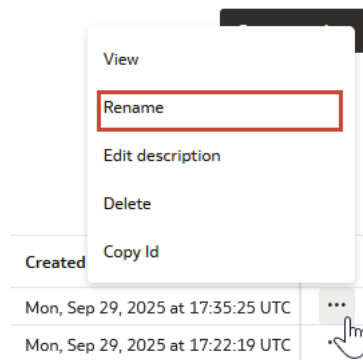
3. Make your changes to the catalog description and click **Save**.

Rename a Standard Catalog

You can rename your catalogs to provide a descriptive label when the contents or purpose of the catalog has changed.

You can't rename the default catalog.

1. On the Home page, click **Master catalog**.
2. Next to your catalog, click **...** **Actions** and click **Rename**.



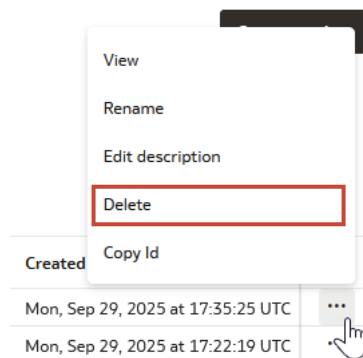
3. Make your changes to the catalog name and click **Save**.

Delete a Standard Catalog

You can delete standard catalogs to remove all locally stored metadata as well as data.

Deleting a standard catalog deletes both the data and the stored metadata.

1. On the Home page, click **Master catalog**.
2. Next to your catalog, click ... **Actions** then click **Delete**.



3. Select **Confirm deletion of the catalogs**.
4. Click **Delete**.

External Catalogs

External catalogs are catalogs where the data is sourced from a location outside Oracle AI Data Platform.



[Video](#)



[LiveLabs Sprint](#)

External catalogs connect AI Data Platform to outside data sources, including Oracle Autonomous AI Lakehouse, Oracle AI Database, and Oracle Autonomous AI Transaction

Processing, and third-party stores such as Snowflake, Azure SQL, and MySQL. AI Data Platform discovers and synchronizes metadata from these sources into the Master Catalog, enabling users to browse and query remote schemas and tables without duplicating the underlying data.

External catalogs use the credentials provided during the external catalog creation for querying the external source. For more information on data sources, see [Internal Sources](#).

For example, if you create an external catalog for an Autonomous AI Lakehouse instance where the Autonomous AI Lakehouse user credentials used have access to *schema1* but not *schema2*, only *schema1* appears in the external catalog. Users with permissions for the external catalog can only query the schema the Autonomous AI Lakehouse user has access to.

Note

AI Data Platform does not support harvesting data from schemas and tables shared across Pluggable Databases (PDBs) or from Oracle-maintained schemas and tables.

Required Permissions for Autonomous AI Lakehouse and Autonomous AI Transaction Processing

When you create an external catalog in AI Data Platform the user credentials you use to connect should have at least the following permissions:

- CREATE SESSION to connect to the database
- SELECT access on the required objects (tables/views/external tables) via least-privilege grants or a dedicated read role
- READ, WRITE on DIRECTORY DATA_PUMP_DIR

If you are inserting data or creating a new table in the external catalog, ensure the user is part of the DWROLE. For more information, refer to the Oracle Autonomous AI Transaction Processing documentation, [Manage User Privileges on Autonomous AI Database - Connecting with a Client Tool](#).

Limitations

UPDATE operation is not supported for external catalogs. Use external catalog tables for discovery and query access.

DDL is not supported, even when the credentials used to create the external catalog has permissions to execute DDL statements.

Table name casing affects how the table is created in Oracle Autonomous AI Lakehouse. Please note the create table behavior, when creating tables from AI Data Platform in Autonomous AI Lakehouse

- If the table name is provided in uppercase, the table is created using uppercase letters and follows the default case-insensitive behavior in Autonomous AI Lakehouse.
- If the table name is provided in lowercase, the table is created as case-sensitive in Autonomous AI Lakehouse. Column names for tables created through Autonomous AI Lakehouse are always created as case-sensitive.
- Source column comments are limited to 256 UTF-8 bytes during external catalog ingestion and refresh. The limit is measured in bytes, not characters, so comments containing multibyte characters can exceed the limit when they contain 256 characters or fewer.

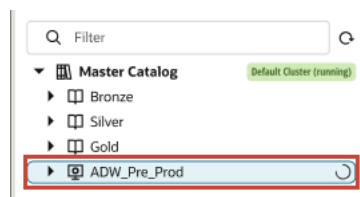
Exceeding the limit can cause metadata ingestion or refresh to fail. Reduce the source column comment to 256 UTF-8 bytes or fewer, and then retry.

Create an External Catalog

You can connect a catalog from your Oracle AI Data Platform instance to a external source.

1. Click **Create** in the left navigation pane and select **Catalog**. You can also navigate to the **Master catalog** and click **Create Catalog in Master Catalog**.
2. Fill in the name and description fields.
3. From the **Catalog Type** drop-down list, select **External Catalog**.
4. Select the external source type.
 - For Oracle Autonomous AI Lakehouse, provide either a wallet file, or the instance configuration.
 - For Oracle Autonomous AI Transaction Processing, provide either a wallet file, or the instance configuration.
 - For Oracle AI Database, provide either a wallet file, or the instance configuration.
 - For Snowflake, provide Snowflake account name, database name, and warehouse. For Authentication method, select **Key Pair**, then provide the private key file and private key passphrase.
 - For Azure SQL, provide the host, port, user name, and password. You can optionally include the database name.
 - For MySQL, provide the host, port, user name, and password.
 - For Kafka, provide the bootstrap server. Separate multiple servers with a comma. **(Coming soon)**
5. Fill in the user name and password.
6. SSL is enabled by default. Clear the box to disable SSL.
7. Click **Create**.

External catalogs that are extracting data from an external source display a spinning circle icon.



You can also monitor progress from **Job Runs**.

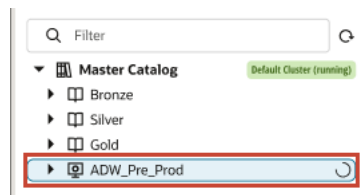
Create an External Catalog for Private Networks

You can create an external catalog that accesses data sources in a private network.

1. Click **Create** in the left navigation pane and select **Catalog**. You can also navigate to the **Master catalog** and click **Create Catalog in Master Catalog**.
2. Fill in the name and description fields.

3. From the **Catalog Type** drop-down list, select **External Catalog**.
4. Select the external source type:
 - For Oracle Autonomous AI Lakehouse, provide either a wallet file, or the instance configuration.
 - For Oracle Autonomous AI Transaction Processing, provide either a wallet file, or the instance configuration.
 - For Oracle AI Database, provide either a wallet file, or the instance configuration.
 - For Oracle Exadata Database Service, provide host, port, and service name (SID).
 - For Snowflake, provide Snowflake account name, database name, and warehouse. For Authentication method, select **Key Pair**, then provide the private key file and private key passphrase.
 - For Azure SQL, provide the host, port, user name, and password. You can optionally include the database name.
 - For MySQL, provide the host, port, user name, and password.
 - For Kafka, provide the bootstrap server. Separate multiple servers with a comma.
(Coming soon)
5. Fill in the user name and password.
6. SSL is enabled by default. Clear the box to disable SSL.
7. Select **Enable private network**.
8. Select the workspace with the desired private network configuration.
For information on setting up a workspace configured for private networks, see [Create a Workspace with Private Network Access Enabled](#).
9. Click **Create**.

External catalogs that are extracting data from an external source display a spinning circle icon.



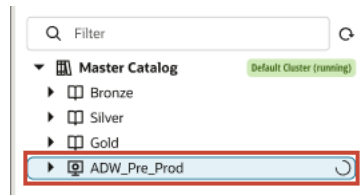
You can also monitor progress from **Job Runs**.

Refresh External Data Catalogs

You can trigger a refresh of all external catalogs to update their contents from the source.

1. On the Home page, click **Master Catalog**.
2. Select the external catalog you want to refresh.
3. Click  **Refresh**.

When you click refresh, workflows start in the background to extract and update metadata from external catalogs. Catalogs that are extracting data from an external source display a spinning circle icon.



You can also monitor progress from **Job Runs**.

Refresh External Catalogs with SQL

You can refresh metadata for external catalogs, schema, and tables using SQL grammar.

To refresh an external catalog, use:

```
REFRESH EXTERNAL CATALOG <<catalog_name>>
```

To refresh a schema in an external catalog, use:

```
REFRESH SCHEMA IN EXTERNAL CATALOG <<catalog_name.schema_name>>
```

To refresh a table in a schema in an external catalog, use:

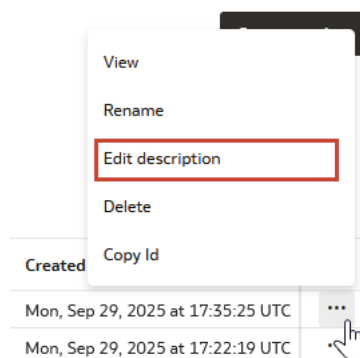
```
REFRESH TABLE IN EXTERNAL CATALOG <<catalog_name.schema_name.table_name>>
```

For more information, see [SQL Grammar](#).

Edit an External Catalog Description

You can edit the description for external catalogs after creation if their contents or purpose has changed.

1. On the Home page, click **Master catalog**.
2. Next to your external catalog, click ... **Actions** and click **Edit description**.



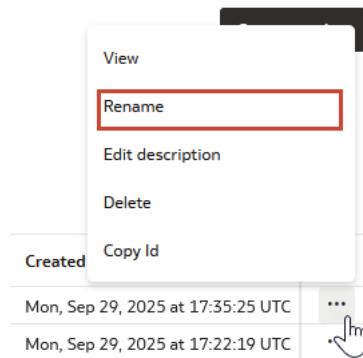
3. Make your changes to the catalog description and click **Save**.

Rename an External Catalog

You can rename your external catalogs to provide a descriptive label when the contents or purpose of the catalog has changed.

You can't rename the default catalog.

1. On the Home page, click **Master catalog**.
2. Next to your catalog, click ... **Actions** and click **Rename**.



3. Make your changes to the catalog name and click **Save**.

Edit an External Catalog Configuration

You can edit the configuration of an external catalogs to update the required password.

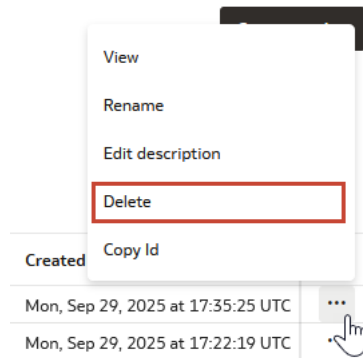
1. On the Home page, click **Master catalog**.
2. Next to your catalog, click ... **Actions** and click **Edit configuration**.
3. Enter the new password for the external catalog and click **Save**.

Delete an External Catalog

You can delete external catalogs to remove all locally stored metadata.

Deleting an external catalog only deletes the locally stored metadata. Data in the data source is not impacted.

1. On the Home page, click **Master catalog**.
2. Next to your catalog, click ... **Actions** then click **Delete**.



3. Select **Confirm deletion of the catalogs**.
4. Click **Delete**.

How to Delete Data from Tables in an External Catalog

You can use the Spark write path to delete rows from an external catalog table when the catalog and table are configured to support pushdown SQL DELETE.

Use Catalog ID to Delete Table Data: Python Example

```
from pyspark.sql.types import StructType

df = spark.createDataFrame([], StructType([]))

df.write
    .format("aidatapatform")
    .option("catalog.id", "<catalog>")
    .option("pushdown.sql", "DELETE FROM <schema>.<target_table> WHERE ID = 2")
    .save()
```

Use Catalog ID to Delete Table Data: Scala Example

```
val df = spark.emptyDataFrame

df.write
    .format("aidatapatform")
    .option("catalog.id", "<catalog>")
    .option("pushdown.sql", "DELETE FROM <schema>.<target_table> WHERE ID = 2")
    .save()
```

Ingestion-style Sample for Deleting Table Data: Python Example

```
from pyspark.sql.types import StructType
df = spark.createDataFrame([], StructType([]))

df.write
    .format("aidatapatform")
```

```
.option("data.asset.type", "ORACLE_ALH")  
.option("tns.alias", alhTns)  
.option("user.name", userName)  
.option("wallet.content", alhWalletContent)  
.option("password", alhPassword)  
.option("pushdown.sql", "DELETE FROM ADMIN.PUSHDOWN_FEATURE WHERE ID = 2")  
.save()
```

Ingestion-style Sample for Deleting Table Data: Scala Example

```
val df = spark.emptyDataFrame  
  
df.write  
  .format("aidatapatform")  
  .option("data.asset.type", "ORACLE_ALH")  
  .option("tns.alias", alhTns)  
  .option("user.name", userName)  
  .option("wallet.content", alhWalletContent)  
  .option("password", alhPassword)  
  .option("pushdown.sql", "DELETE FROM ADMIN.PUSHDOWN_FEATURE WHERE ID = 2")  
  .save()
```

Schema

Schema in data catalogs are constructs to organize data.

Catalogs are logical containers for schema, which are also referred to as databases. Schema can contain tables, which contain structured data, and volumes, which contain unstructured data.

A default schema is created in all standard catalogs created in the Master Catalog. To create additional schema, see [Create a Schema](#).

You can manage permissions to control who has access to your schema. For more information, see [Schema Permissions](#).

Master catalog's default catalog, which is called 'default', contains a reserved schema named `oci_ai_models`. If you have the requisite permissions on OCI Generative AI models in the region, the model is displayed in the `oci_ai_models` schema and you can drag and drop the models onto a notebook to auto-generate code for batch inference. For more information, see [OCI Generative AI \(Pretrained Foundation Models\)](#).

Note

You can't change the name of a schema in a standard catalog.

Create a Schema

You can create schema in catalogs you own or are shared with you.

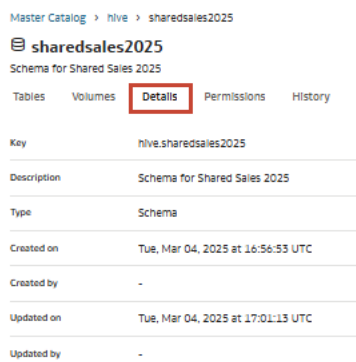
1. Click **Create** in the left navigation pane and select **Schema**. You can also navigate to the catalog in which you want to create a schema, click the Schema tab, and click  **Create Schema**.

2. Provide a name and description for your schema and click **Create**.

View Schema Details

You can view schema information and resources from the schema Details tab.

1. On the home page, click **Master Catalog**.
2. Navigate to your schema, then click the **Details** tab.



The screenshot shows the 'Master Catalog' interface. The breadcrumb trail is 'Master Catalog > hive > sharedsales2025'. The main heading is 'sharedsales2025' with a subheading 'Schema for Shared Sales 2025'. There are four tabs: 'Tables', 'Volumes', 'Details' (which is selected and highlighted with a red box), 'Permissions', and 'History'. Below the tabs is a table with the following rows:

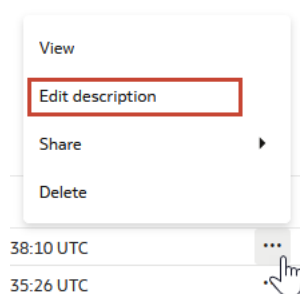
Key	hive.sharedsales2025
Description	Schema for Shared Sales 2025
Type	Schema
Created on	Tue, Mar 04, 2025 at 16:56:53 UTC
Created by	-
Updated on	Tue, Mar 04, 2025 at 17:01:13 UTC
Updated by	-

3. Click **Details**.

Edit a Schema Description

You can edit the description for a schema to provide an updated summary of its contents.

1. On the Home page, click **Master Catalog**.
2. Next to the schema you want to change the description for, click ... **Actions** and click **Edit Description**.



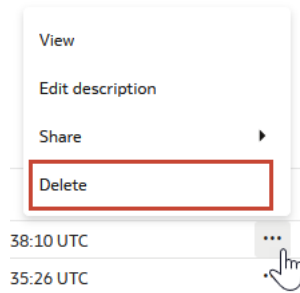
3. Provide a new description. Click **Save**.

Delete a Schema

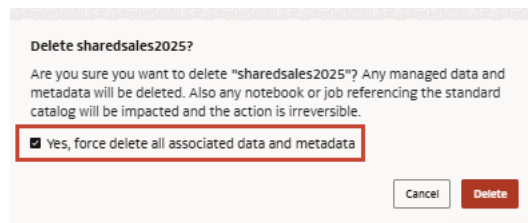
You can delete schema and all their metadata and Oracle AI Data Platform-managed data when they are no longer needed.

Deleting a schema also deletes all metadata and AI Data Platform-managed data in managed tables and managed volume associated with that schema. Make sure you backup or move any data you want to retain to another location before deleting a schema.

1. On the Home page, click **Master catalog**.
2. Next to the schema you want to delete, click ... **Actions** and click **Delete**.



3. Select **Yes, force delete all associated data and metadata**.



4. Click **Delete**.

Tables

Tables define the structure for your data.

You can load new data into your tables or reference data in an existing location. You can define fine-grained access control permissions on tables by creating table permissions.

Tables can either be external or managed.

External tables

An external table defines a structure for data that's stored in a location not managed by Oracle AI Data Platform. When you create an external table in the AI Data Platform, the metadata life cycle is managed by AI Data Platform. When you delete an external table, only the table definition is deleted. The data referenced by the external table isn't deleted.

Ensure your users have the following IAM policies required to create external tables:

```
allow group <GroupName> to read buckets in compartment id <external-data-CompartmentId>
allow group <GroupName> to inspect objects in compartment id <external-data-CompartmentId>
```

Additional IAM policies are required for external tables. For more information, see [IAM Policies for Oracle AI Data Platform](#).

Managed Tables

A managed table defines a structure for data that's stored within the AI Data Platform and can only be accessed by AI Data Platform users.

When you delete a managed table, the table definition and the table data is deleted.

Supported Table Formats

Tables can use file-based formats such as CSV, JSON, Avro, Parquet, and ORC, or transactional table formats such as Delta and Apache Iceberg. Iceberg tables use the Hadoop catalog implementation to manage table metadata in Oracle Cloud Infrastructure Object Storage.

Format	Description	Usage
Comma-separated-values (CSV)	Data is stored as a text file with a specified row based file format to structure the data. Typically, the first row in the file is a header row that contains columns names for the data.	Used to exchange tabular data between systems. Each row in the file is a row in a table.
JavaScript Object Notation (JSON)	Data is stored in a standard text-based format for representing structured data based on JavaScript object syntax. JSON supports lists of objects or hierarchical structures.	Used in stream applications. JSON simplifies the storage of related data with complex relationships in a single document and avoids chaotic list conversion to a relational data model. Note that JSON isn't splittable.
Avro	Data is stored in a row based binary format while the schema is stored in JSON format to minimize file size and maximize efficiency. Avro has reliable support for schema evolution by managing added, missing, and changed fields. This lets old software to read new data, and new software to read old data. Also known as the data serialization system.	Used for data storage as avro files are splittable and compressible. The serialized row-based storage is ideal for heavy write transaction, such as inserting data into AI Data Platform. Avro is also a good choice when schema evolution is critical during high speed writes.
Parquet	Data is stored in a columnar data format and is highly compressible and splittable. Parquet is optimized for the paradigm Write Once Read Many (WORM). It writes slowly but reads incredibly quickly, especially when you only access a subset of columns.	Used for solving Big Data problems as compression algorithms work better with columnar data format. You can store Big Data in various formats, such as images, videos, documents, and structured data tables. Parquet is a good choice for heavy workloads when reading portions of data. For example, when the dataset has many columns, but you only want to access a subset of columns. Ideal when you're dependent on Spark or when you want several services to access the same data stored in Object Storage.

Format	Description	Usage
Optimized Row Columnar (ORC)	Data is stored in collections of rows in a single file in columnar format.	Used for parallel processing of row collections across a cluster. Ideal when read transactions are more than write transactions or when compression is priority.
Delta	Data is stored in a columnar format that extends Parquet data files with a JSON file-based transaction log for ACID transactions and scalable metadata handling.	Used for transaction support.
Iceberg	An open table format that manages Parquet data files through versioned metadata, manifests, and snapshots. AI Data Platform uses the Iceberg Hadoop catalog implementation.	Use for ACID transactions, concurrent reads and writes, schema and partition evolution, time-travel queries, row-level updates, upserts, and snapshot management.

Supported Data Types

Data Type	Description
ByteType	Represents 1-byte signed integer numbers. The range of numbers is from -128 to 127.
ShortType	Represents 2-byte signed integer numbers. The range of numbers is from -32768 to 32767.
IntegerType	Represents 4-byte signed integer numbers. The range of numbers is from -2147483648 to 2147483647.
LongType	Represents 8-byte signed integer numbers. The range of numbers is from -9223372036854775808 to 9223372036854775807.
FloatType	Represents 4-byte single-precision floating point numbers.
DoubleType	Represents 8-byte double-precision floating point numbers.
DecimalType	Represents arbitrary-precision signed decimal numbers. Backed internally by java.math.BigDecimal. A BigDecimal consists of an arbitrary precision integer unscaled value and a 32-bit integer scale.
StringType	Represents character string values.
VarcharType(length)	A variant of StringType which has a length limitation. Data writing will fail if the input string exceeds the length limitation.
CharType(length)	A variant of VarcharType(length) which is fixed length. Reading column of type CharType(n) always returns string values of length n. Char type column comparison will pad the short one to the longer length.
BinaryType	Represents byte sequence values.
BooleanType	Represents boolean values.
DateType	Represents values comprising values of fields year, month and day, without a time-zone.

Data Type	Description
TimestampType	Timestamp with local time zone(TIMESTAMP_LTZ). It represents values comprising values of fields year, month, day, hour, minute, and second, with the session local time-zone. The timestamp value represents an absolute point in time.
TimestampNTZType	Timestamp without time zone(TIMESTAMP_NTZ). It represents values comprising values of fields year, month, day, hour, minute, and second. All operations are performed without taking any time zone into account.
YearMonthIntervalType(startField, endField)	Represents a year-month interval which is made up of a contiguous subset of MONTH, months within years [0..11] and YEAR, years in the range [0..178956970].
DayTimeIntervalType(startField, endField)	Represents a day-time interval which is made up of a contiguous subset of SECOND, seconds within minutes and possibly fractions of a second [0..59.999999], MINUTE, minutes within hours [0..59], HOUR, hours within days [0..23] and DAY, days in the range [0..106751991].
ArrayType(elementType, containsNull)	Represents values comprising a sequence of elements with the type of elementType. containsNull is used to indicate if elements in a ArrayType value can have null values.
MapType(keyType, valueType, valueContainsNull)	Represents values comprising a set of key-value pairs. The data type of keys is described by keyType and the data type of values is described by valueType. For a MapType value, keys are not allowed to have null values. valueContainsNull is used to indicate if values of a MapType value can have null values.
StructType(fields)	Represents values with the structure described by a sequence of StructFields (fields).
StructField(name, dataType, nullable)	Represents a field in a StructType. The name of a field is indicated by name. The data type of a field is indicated by dataType. nullable is used to indicate if values of these fields can have null values.

Limitations

The following limitations apply to tables in AI Data Platform:

- You cannot define an external table on any data files or directories within/on a volume.
- You cannot define an external table on a bucket and/or its directory that is already used for another external table or external volume
- Views cannot be viewed/listed in the Master Catalog.

Schema Evolution

Schema evolution in Oracle AI Data Platform allows users with the required permissions to update a managed table using SQL in a notebook.

This is useful when a table definition changes over time to support new columns, removed columns, renamed columns, partition changes, or table renames without recreating the dataset from scratch. The supported formats are Parquet, Avro, and Delta.

Supported Operations

The following schema evolution operations were analyzed for managed tables:

- **Rename table:** supported for Delta, Parquet, Avro, and Iceberg
- **Add columns:** supported for Delta, Parquet, Avro, and Iceberg
- **Drop columns:** supported for Delta and Iceberg; not supported for Parquet and Avro
- **Alter or rename columns:** supported for Delta and Iceberg; not supported for Parquet and Avro
- **Replace columns:** supported for Delta and Iceberg
- **Add partitions:** supported for Parquet and Avro via DDL. For Delta and Iceberg, partitions are created automatically as data is written. Individual partitions are not added through DDL. Iceberg supports adding fields to the partition specification using `ALTER TABLE ... ADD PARTITION FIELD`.
- **Drop partitions:** supported for Parquet and Avro via DDL. For Delta, remove the relevant data and run `VACUUM`. Iceberg does not support dropping individual partitions through DDL; delete the relevant rows instead. Iceberg supports removing fields from the partition specification using `ALTER TABLE ... DROP PARTITION FIELD`.
- **Change data type:** not supported for Parquet or Avro. Delta may require a CTAS or overwrite-schema workaround. Iceberg supports only compatible type promotions, such as `INT` to `BIGINT`, `FLOAT` to `DOUBLE`, and increasing decimal precision, through `ALTER TABLE`. Narrowing or incompatible conversions are not supported.

Create a Managed Table

You can create tables for schema you manage.

1. Navigate to the schema you want to create a table for.
2. Select the **Tables** tab.
3. Click  **Create Table**.

Create table

Table type
Managed

Managed table format
A/R0

Drop a sample file here or click to browse.

File to be uploaded:
No files selected

Table preview
Select a file to see the preview

Cancel Create

4. Select **Managed** for your **Table Type**.
5. Select the format for your table from **Managed table format**.
6. Either drag and drop a file with your table data or click to browse to the file location.
7. Provide a name and description for your table.
8. Optional: To add partitions, expand **Partition keys (optional)**. Click **Add Partition** and select a data column.
9. Optional: To add table properties to the data catalog's metadata, expand **Table properties (optional)**. Click **Add Property** and provide the property and its value.
10. Click **Create**.

Create an External Table

You can create an external table with data in OCI Object Storage.

1. Navigate to the schema you want to create a table for.
2. Select the **Tables** tab.
3. Click **Create Table**.

4. Select **External** for your **Table Type**.
5. Select the compartment, bucket, and folder from OCI Object Storage where data is stored. The objects you can select are based on the logged in user's IAM permissions.
6. Provide a name and description for your table.
7. Optional: To add table properties to the data catalog's metadata, expand **Table properties (optional)**. Click **Add Property** and provide the property and its value.
8. Click **Create**.

Edit a Table

You can modify details of tables you manage.

Note

Edits to external catalog tables made in Oracle AI Data Platform are not pushed to the remote catalog.

1. Navigate to your schema.
2. Select the **Tables** tab.
3. Next to the table you want to edit click **...** **Actions**.
 - Click **Rename** to change your table's name. Enter a new name and press Enter.
 - Click **Edit Description** to change your table's description. Provide the new description and click **Save**.

View Table Details

You can view the details of tables in schema.

1. Navigate to your schema. Click the **Tables** tab.

2. Click the name of the volume you want to view details for. You can also click ... **Actions** next to the volume then click **View**.
3. Click the **Details** tab.

Delete a Table

You can delete tables from schema you manage.

1. Navigate to the schema you want to delete your table from.
2. Click the **Tables** tab.
3. Next to the table you want to delete, click ... **Actions** and click **Delete**.
4. Click **Delete**.

Volumes

Volumes are containers to store data in its original form and can store semi-structured or unstructured data.

You can load new data into your volume or reference data in an existing location. You can define access control permissions on a volume by creating volume permissions.

Volumes can either be external or managed.

External volumes

An external volume refers to an existing location in OCI Object Storage. When users create an external volume, Oracle AI Data Platform manages only the metadata of the volume, the data life cycle is managed by the customer.

Ensure your users have the following IAM policies required to create external volumes:

```
allow group <GroupName> to read buckets in compartment id <external-data-CompartmentId>
allow group <GroupName> to inspect objects in compartment id <external-data-CompartmentId>
```

When you delete an external volume, only the volume definition is deleted. The data referenced by the external volume isn't deleted.

Managed volumes

A managed volume creates a location in OCI Object Storage for storing the data. AI Data Platform manages both the metadata and data of the volume.

When you delete a managed volume, the volume definition and the volume data is deleted.

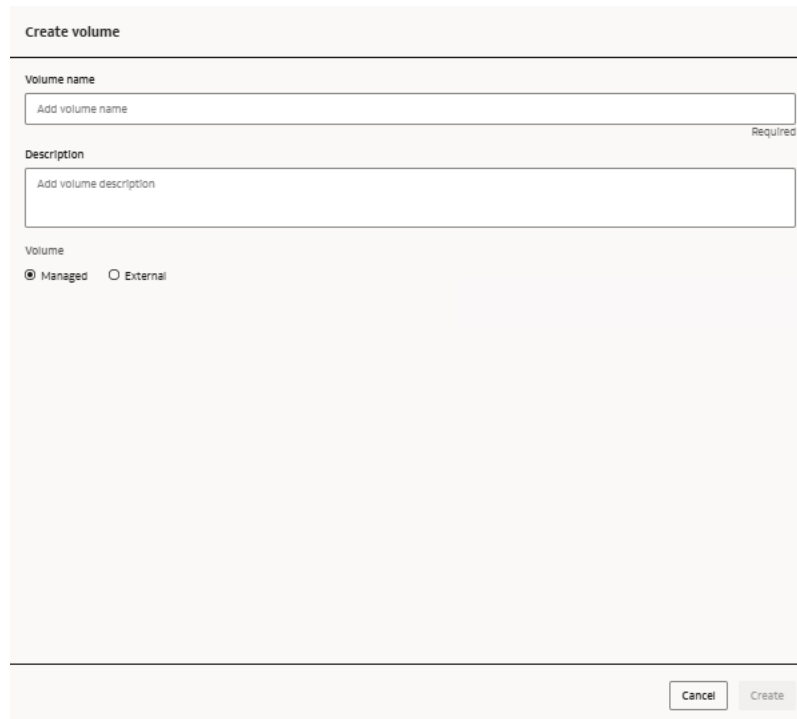
Limitations

You cannot define an external volume on a bucket and its directory if they are already used for another external table or external volume.

Create a Managed Volume

You can create managed volumes for schema you manage.

1. Navigate to the schema you want to create a volume for.
2. Select the **Volumes** tab.
3. Click  **Create Volume**.
4. Provide a volume name and description.
5. Select the **Managed** option.



The image shows a 'Create volume' dialog box. It has a title bar 'Create volume'. Below the title bar, there are two text input fields. The first is labeled 'Volume name' and contains the placeholder text 'Add volume name'. The second is labeled 'Description' and contains the placeholder text 'Add volume description'. To the right of the 'Description' label, the word 'Required' is written. Below the input fields, there are two radio buttons. The first is labeled 'Managed' and is selected (indicated by a filled circle). The second is labeled 'External' and is not selected (indicated by an empty circle). At the bottom right of the dialog box, there are two buttons: 'Cancel' and 'Create'.

6. Click **Create**.

Create an External Volume

You can create external volumes for schema you manage.

1. Navigate to the schema you want to create a volume for.
2. Select the **Volumes** tab.
3. Click  **Create Volume**.
4. Select the **External** option.

Create volume

Volume name
Add volume name Required

Description
Add volume description

Volume
☐ Managed ☒ External

Compartment
Select a Compartment Browse

Bucket
Select a Bucket Browse

Folder (Optional)
Select a Folder Browse

Cancel Create

5. Select the compartment, bucket, and folder for your external volume. The objects you can select are based on the logged in user's IAM permissions.
6. Click **Create**.

Edit a Volume

You can modify details of volumes you manage.

1. Navigate to your schema.
2. Select the **Volumes** tab.
3. Next to the volume you want to edit click ... **Actions**.
 - Click **Rename** to change your volume's name. Enter a new name and press Enter.
 - Click **Edit Description** to change your volume's description. Provide the new description and click **Save**.

View Volume Details

You can view the details of volumes in schema.

1. Navigate to your schema. Click the **Volumes** tab.
2. Click the name of the volume you want to view details for. You can also click ... **Actions** next to the volume then click **View**.
3. Click the **Details** tab.

Delete a Volume

You can delete volumes from schema you manage.

1. Navigate to the schema you want to delete your volume from.
2. Click the **Volumes** tab.
3. Next to the volume you want to delete, click **Actions** and click **Delete**.
4. Click **Delete**.

Synonyms

Oracle AI Data Platform supports private and public synonyms imported from external catalogs backed by Oracle Autonomous AI Database.

A synonym is an alias for another schema object. Oracle AI Data Platform supports public and private synonyms from external catalog sources. Synonyms are typically used to simplify object references or provide a stable, business-friendly name. Access to the underlying object still depends on the applicable source and AI Data Platform permissions. For more information about synonyms, see [Oracle AI Database Concepts](#).

You can use the synonym to reference the underlying table or view without needing the three-part name for that table or view.

Private synonyms

Private synonyms belong to a specific schema. They are displayed in that schema's Synonyms folder and are referenced as `<catalog_name>.<schema_name>.<private_synonym_name>`.

Public synonyms

Public synonyms are defined at the source database level rather than being owned by a specific schema. In AI Data Platform, they are displayed under the Default (Public synonyms) schema. Access to the underlying object remains subject to applicable permissions. Public synonyms are referenced with three-part names that always include the default schema and can be called from notebooks and agents with this format:
`<catalog>.default.<public_synonym_name>`

Synonym-backed objects are queried in AI Data Platform through SQL and notebooks using the appropriate catalog naming pattern for the synonym type.

Private Synonym Querying

Private synonyms are queried using their three-part catalog naming.

Private Synonym SQL Example

```
SELECT *  
FROM <catalog_name>.<schema_name>.<synonym_name>
```

Private Synonym PySpark Example

```
df = spark.read.table("<catalog_name>.<schema_name>.<synonym_name>")  
df.show()
```

Public Synonym Querying

Public synonyms belong to the default schema in a catalog, so they are queried with the three-part catalog naming pattern including default as the schema.

Public Synonym SQL Example

```
SELECT *  
FROM <catalog_name>.default.<public_synonym_name>
```

Public Synonym PySpark Example

```
df = spark.read.table("<catalog_name>.default.<public_synonym_name>")  
df.show()
```

View Synonym Details

You can view the details for the underlying data of a synonym or the details for the synonym object itself.

1. On the home page, click **Master Catalog**.
2. Navigate to the synonym in your catalog.
3. To view the columns exposed through the synonym, select the Columns tab.
4. To view details for the synonym object, select the Details tab.

28

Lineage (Preview)

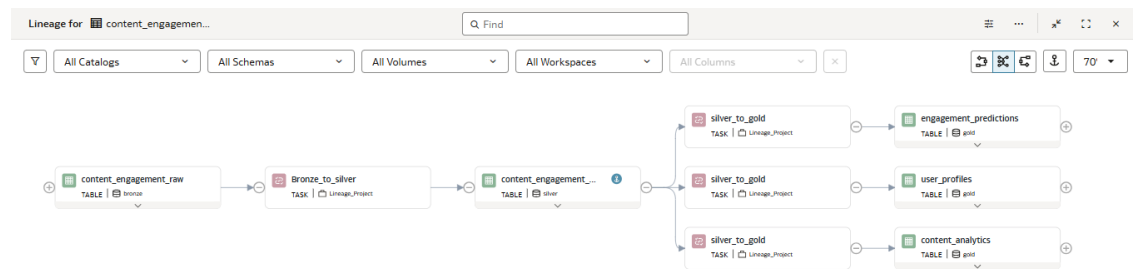
Lineage in Oracle AI Data Platform shows how data artifacts are related through notebook and workflow runs. The lineage graph helps you trace upstream sources, downstream consumers, and column-level derivations for supported artifacts.

Lineage is captured automatically when notebooks and workflow tasks run on a newly created compute. If a compute was created before lineage became available, restart the compute to begin capturing lineage from executions that run on it.

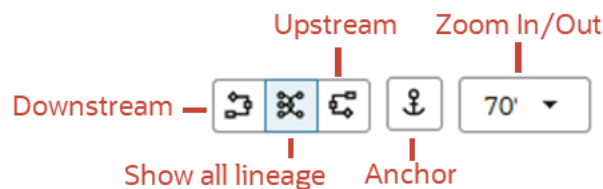
Note

For each process, the service displays only the latest captured lineage. Historical lineage is not currently available.

You can view the lineage of AI Data Platform artifacts from the Master Catalog. Right-click on an artifact and select **Lineage**. Tables and volumes are displayed in the lineage diagram while tables and knowledge bases are currently supported as anchor nodes.



The Lineage view displays a lineage graph with upstream and downstream artifacts for the selected data artifact. You can switch between the full graph, upstream-only view, and downstream-only view.

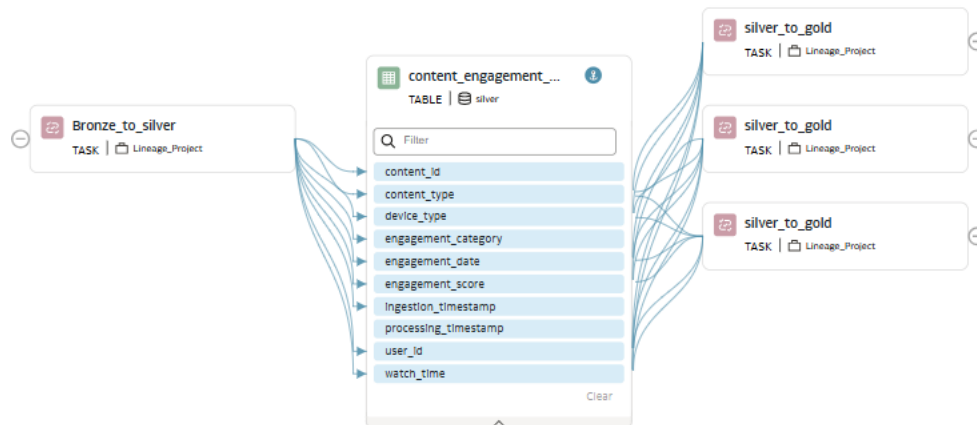


You can view column-level lineage to trace how columns in one data artifact are derived from, transformed by, or propagated to columns in other artifacts.

You can hide the filters at the top of your canvas by clicking the Filter icon in the top-left.



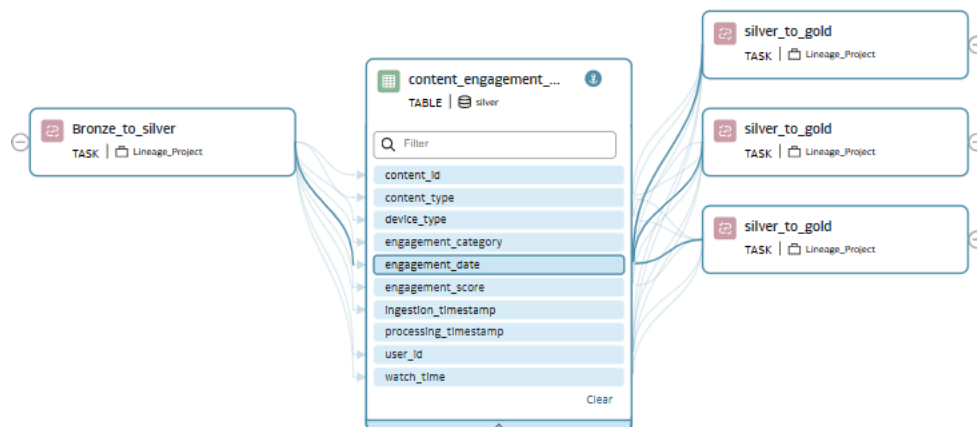
You expand data artifacts in your Lineage flow by clicking the down arrow at the bottom of the artifact card. When the artifact is expanding, you can see upstream and downstream inheritance of specific data columns. This function works for artifacts that contain data columns, like tables and volumes.



For expanded artifact cards, you expand a table or volume to view its columns and the column-level lineage relationships connected to them. You expand data artifacts in your Lineage flow by clicking the down arrow at the bottom of the artifact card. When the artifact is expanding, you can see upstream and downstream data flow for specific columns. This function works only for artifacts that contain data columns, like tables and volumes.

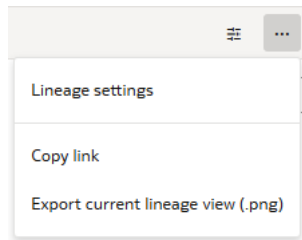
You can expand multiple tables and volumes in your lineage graph to see the data flow from each. When you expand the data artifact, blue arrows show how columns in source artifacts contribute to columns in target artifacts through notebook or workflow executions. You highlight the path of an individual column by double-clicking it.

Blue arrows show column-level lineage relationships between source and target columns. These relationships indicate how data is derived, transformed, or propagated across tables, volumes, notebooks, tasks, and workflows. Double-click a column to highlight its lineage path across the graph.



You can select multiple data columns by Shift- or Ctrl-clicking them to highlight multiple paths.

From the Actions menu in the top-right of the Lineage window, you can control your Lineage settings, which affects the depth of upstream and downstream artifacts displayed, or you can share your lineage diagram, either by copying a link or by exporting a PNG image.

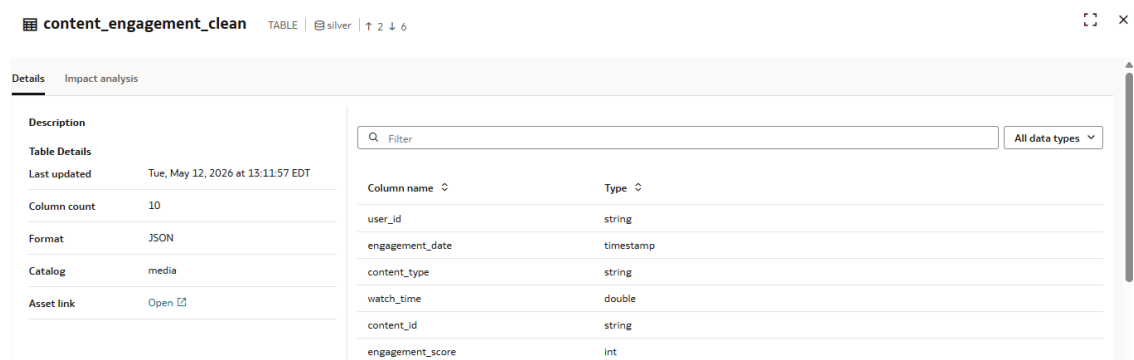


Lineage Details

Double-clicking on an artifact in the lineage diagram shows details for that artifact. For tasks, the details page provides both details for the task and the job it belongs to. For tables and volumes, the details page provides information on the table or volume and their columns.

You can right-click data artifacts to either **View Details** or **Set as Anchor**. Setting the data artifact as anchor changes the currently displayed diagram to center on that node instead.

At the top of the Details window, you can see the artifact type, the schema it belongs to, and the number of upstream and downstream artifacts. In the Description pane, clicking the **Asset link** takes you to the artifact in your workspace.



For Data artifacts, the Details window shows when the artifact was last updated, information on data columns, format, and the catalog to which the data artifact belongs. You can search for specific data columns by name and filter by data type using the drop-down menu.

For Process artifacts, which include tasks and notebooks, the Details window displays information related to the artifact, including the most recent task and job status, duration, task type, job or notebook name and ID, and the attached cluster. In the right pane, you can search for source and target artifacts based on artifact name or using the drop-down menu to filter for transformation type.

Changes the names or description of an artifact are automatically reflected in the lineage diagram.

Knowledge Base Lineage

Knowledge bases function similarly to data artifacts but only show upstream connections, as they have no downstream items. These upstream connections are the data sources selected for the knowledge base, either volumes or folders in a volume.

Double-clicking on a knowledge base artifact shows the details for that knowledge base, including the number of data sources, and the catalog and schema where the knowledge base was created in the workspace.



You can click on the upstream volume and folders to see details on which catalog and schema the volume belongs to, number of contained folders, and if the volume is managed or external.



You can export the lineage for knowledge bases the same way you export other data artifacts. See [Export Impact Analysis \(Preview\)](#).

For more information on the creation and management of knowledge bases, see [Knowledge Bases](#).

Transformation Types

AI Data Platform supports the following transformation types when tracking lineage:

Type	Meaning	Example Scenario	Example Field Mapping
AGGREGATION	The output field is computed by aggregating multiple input records.	Creating summary tables or metrics.	total_sales = SUM(amount)

Type	Meaning	Example Scenario	Example Field Mapping
IDENTITY	The output field is exactly the same as the input field (no change).	Copying a dataset from one table to another.	customer_id → customer_id
TRANSFORMATION	The output is derived from input fields using functions, casts, concatenation, etc.	Standardizing or cleaning data.	full_name = CONCAT(first_name, ' ', last_name)

Impact Analysis

Data artifacts selected as the anchor node have an additional tab in their Details window for Impact Analysis. From the Impact Analysis tab, you can search for specific artifact names or filter by artifact type. You can select Upstream or Downstream to only show artifacts that are upstream or downstream of the currently selected artifact.

The screenshot displays the 'Impact analysis' tab for the 'content_engagement_clean' artifact. It includes a search bar, a filter for 'All artifact types', and buttons for 'Upstream' and 'Downstream'. Below these controls, a table lists the lineage artifacts:

Artifact name	Type	Direction	Depth
Bronze_to_silver	Task	↑ upstream	1
content_engagement_raw	Table	↑ upstream	2
silver_to_gold	Task	↓ downstream	1

Use upstream impact analysis to understand dependencies. Use downstream impact analysis to identify consumers that may be affected by changes to the selected artifact.

Click **Export impact analysis** to export the artifacts related to the selected data artifact. You can export upstream artifacts, downstream artifacts, or all related artifacts.

Enable, Disable, or Re-enable Lineage Capture for a Compute

Lineage capture is managed independently for each compute.

Enable lineage capture

Lineage capture requires no configuration for a newly created compute. If the compute was created before lineage became available, restart it to enable lineage capture.

Disable lineage capture

To disable lineage capture for a compute, open the compute Advanced options and add the following property to the Spark configuration:

```
spark.aidp.lineage.enabled = false
```

Disabling lineage on one compute does not affect notebooks or workflow tasks running on other computes.

Set number of workers

☒ Static amount ☐ Autoscale

Number of Workers

1

Run duration

☒ Forever ☐ Idle timeout

Idle timeout in minutes

▼ Advanced options - Adding Spark configurations and Init script will result in slower startup

Spark

Spark configuration

spark.ai.dp.lineage.enabled = false

Spark environment variables

Init scripts

Add init script

Cancel Save and Restart

Re-enable lineage capture

To re-enable lineage capture after it has been disabled, open the compute Advanced options and change the property to:

```
spark.ai.dp.lineage.enabled = true
```

Limitations

Column lineage is not available for transformation entities that have more than 1000 columns. You can only view entity lineage for these entities.

Entity and Column Lineage (Preview)

In some lineage scenarios where multiple upstream datasets participate in producing a target dataset, only some of those upstream datasets contribute actual column values to the target.

The key distinction between entity lineage and column lineage is the question they answer:

- **Entity lineage answers:** Which datasets participated in creating the target?
- **Column lineage answers:** Which source columns supplied the target column values?

Because these questions are different, entity lineage and column lineage can look different for the same pipeline.

In some transformations, one input provides the rows and column values written to the target, while another input is used only as a reference for filtration. In these cases:

- **Entity lineage** should show all upstream datasets that the target depends on.
- **Column lineage** may show column-level flow only from the value-providing input.
- A reference input can affect the target row set without contributing values to target columns.

This behavior is expected.

Example: Entity and Column Lineage

Assume two source datasets contain the same columns, but not the same rows:

- source_table_1 contains the primary dataset.
- source_table_2 contains a reference set of rows.
- The target table is created by keeping only the rows that exist in both source tables.

For example:

Table 28-1 source_table_1

product_id	sales_date	quantity	total_amount
101	2025-06-01	10	150.0
102	2025-06-02	20	300.0
103	2025-06-03	15	225.0
104	2025-06-04	12	180.0

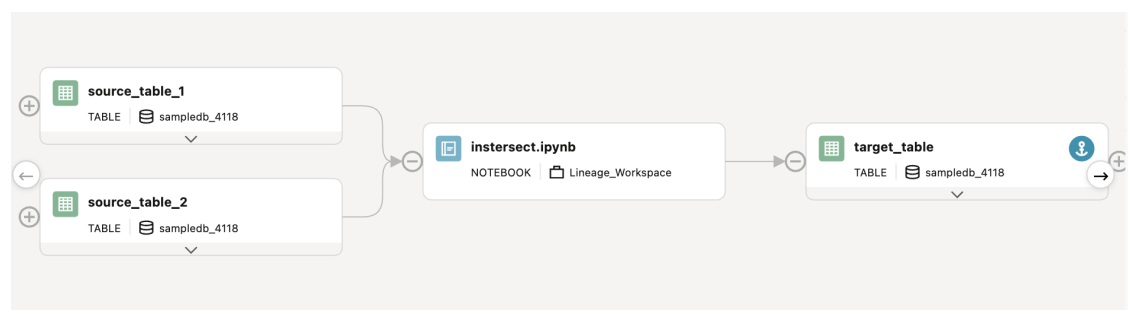
Table 28-2 source_table_2

product_id	sales_date	quantity	total_amount
102	2025-06-02	20	300.0
103	2025-06-03	15	225.0
105	2025-06-05	18	270.0

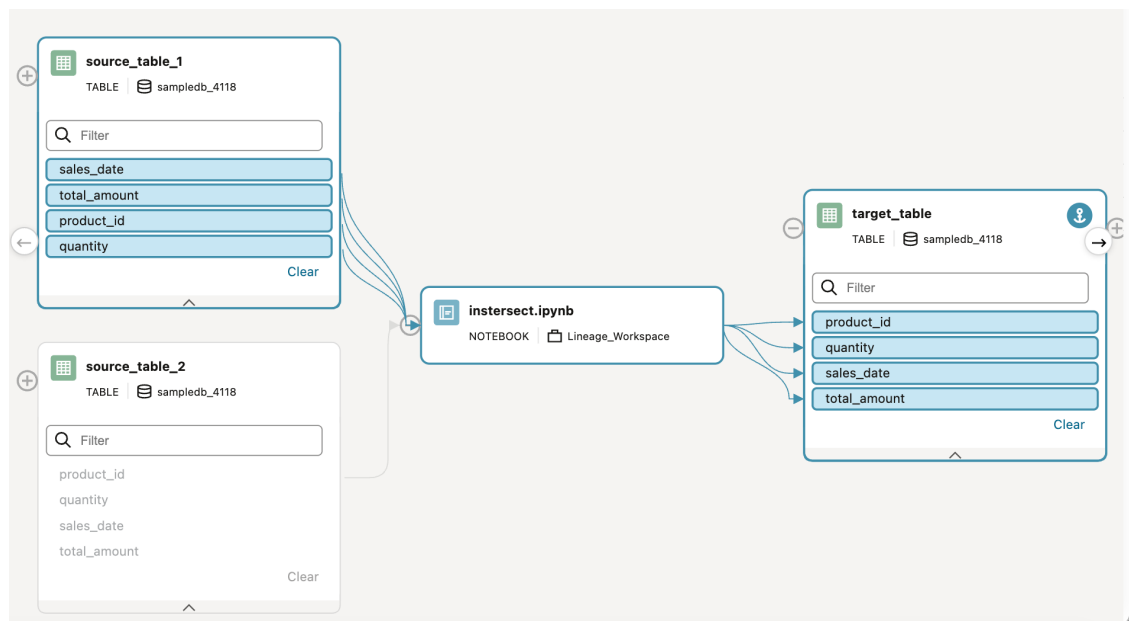
Table 28-3 target_table

product_id	sales_date	quantity	total_amount
102	2025-06-02	20	300.0
103	2025-06-03	15	225.0

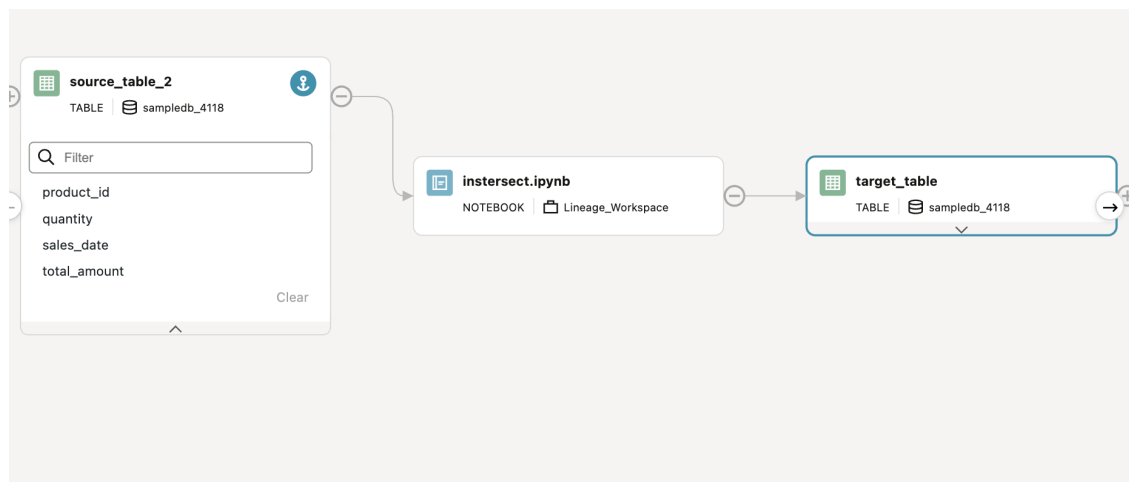
In this example, both source tables participate in creating the target because both are required to determine the final row set.



However, from a column lineage perspective, the target column values may be attributed only to the value-providing input, such as source_table_1. The second input, source_table_2, is used to determine which rows qualify for the target, but its values are not necessarily copied into the target columns.



For these reasons, when the lineage view is anchored on **source_table_2**, no column-level lineage links are displayed, as shown below.



Why Entity Lineage Shows Both Inputs

Entity lineage captures dataset-level dependency. If a processing job reads two datasets and the result depends on both, both datasets are legitimate upstream entities. In this pattern:

- The target cannot be fully explained without **Source Dataset A**.
- The target also cannot be fully explained without **Source Dataset B**, because Source Dataset B determines which records from Source Dataset A are retained.
- Therefore, both Source Dataset A and Source Dataset B should appear as upstream entities for Target Dataset C.

This is **dependency lineage**, not value lineage.

Why Column Lineage Shows Only the Value-Providing Input

Column lineage captures value provenance. It describes where the values in each target column came from.

For example, if the target table is written using rows from Source Dataset A after filtering rows from Source Dataset B, then the target column values still originate from Source Dataset A.

Example column mappings:

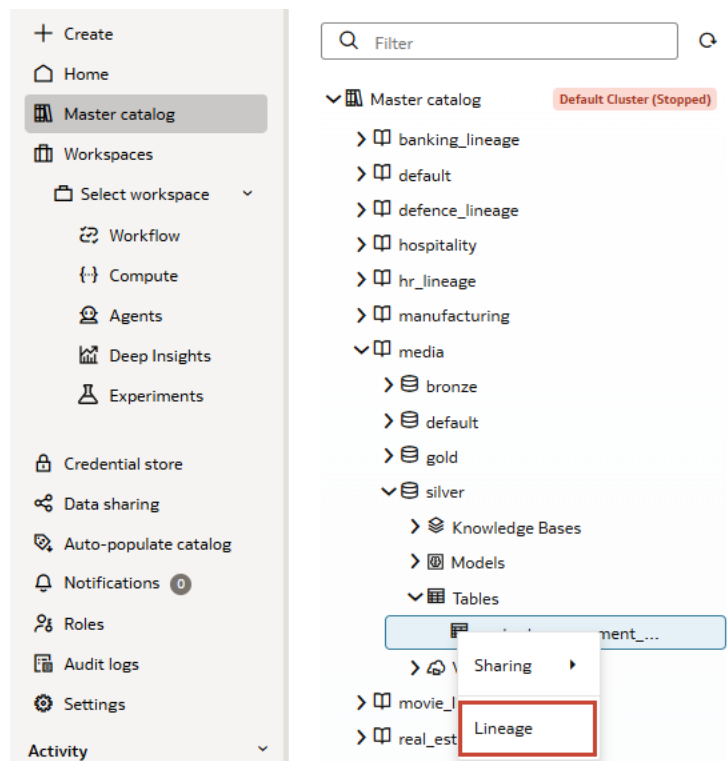
Target Column	Source Column
target.product_id	source_a.product_id
target.sales_date	source_a.sales_date
target.quantity	source_a.quantity
target.total_amount	source_a.total_amount

Source Dataset B influences whether a row is present, but its column values are not copied into the target. As a result, Source Dataset B may appear in entity lineage while not appearing in column lineage.

View Data Lineage (Preview)

You can see the inheritance of data in your workspace as it moves between different Oracle AI Data Platform artifacts.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.



3. The lineage diagram is displayed.

View Lineage for Specific Data Columns (Preview)

You can trace the lineage of a specific data column through your lineage diagram.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. Click the arrow at the bottom of a table or volume artifact to expand it.
4. Double-click the data column you want to highlight the lineage for.

View Details for a Lineage Artifact (Preview)

You can see additional details for an artifact in your lineage diagrams.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. Double-click an artifact on the lineage diagram to view additional details. You can also right-click and click **View Details**.
4. Click the Impact Analysis tab to view the upstream and downstream impact of the artifact. This tab is only available for the anchor node.

Export Impact Analysis (Preview)

You can export the impact analysis for data artifacts while viewing the details of a lineage artifact.

Note

You can only export impact analysis for data artifacts.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. Double-click a data artifact in the lineage diagram. Select the Impact Analysis tab.
4. Click **Export impact analysis**.
5. From the drop-down menu, select if upstream, downstream, or all artifacts should be included.
6. Click **Export**.

Filter Lineage Flow Diagram (Preview)

You can filter your lineage diagram to help focus on more specific data points when examining lineage.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. From the drop-down menus, select specific catalogs, schemas, volumes, or workspaces to filter out results from.

Search for Artifacts in Lineage Flow Diagram (Preview)

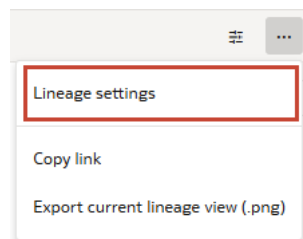
You can search for strings to locate specific artifacts in the lineage diagram when viewing artifact lineage.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. In the **Search** field at the top of your lineage diagram, enter the string to search for.
4. Click a result in the list to center the diagram on that artifact.

Change Lineage Flow Depth (Preview)

You can alter how many levels of upstream or downstream artifacts your lineage diagram displays to help you either expand or narrow the focus of your diagram.

1. Navigate to the artifact in your Master Catalog you want to view the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. Click ... **Actions** in the top-right
4. Click **Lineage Settings**.

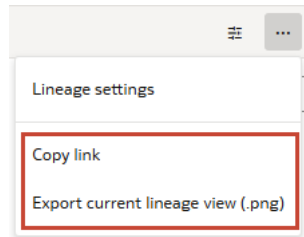


5. Modify **Upstream depth** and **Downstream depth** as needed.
6. Click **Save**.

Share a Lineage Flow Diagram (Preview)

You can share the lineage diagram showing the lineage of a specific object as either a direct link or a PNG image.

1. Navigate to the artifact in your Master Catalog you want to share the lineage for.
2. Right-click the artifact then click **Lineage**. You can also select the artifact and click **Actions** in the top-right, then click **Lineage**.
3. Click ... **Actions** in the top-right.



4. Choose how you want to share your lineage diagram:
 - Click **Copy link** to copy a link directly to your clipboard. Paste the link to share it.
 - Click **Export current lineage view (.png)** to export the current view of your lineage diagram, including any filters you have applied.

Credential Store (Preview)

This chapter covers the using the credential store in Oracle AI Data Platform to create, manage, and provision access credentials.

Topics:

- [About Credential Store \(Preview\)](#)
- [Create Credentials \(Preview\)](#)
- [Use Stored Credentials in a Notebook \(Preview\)](#)
- [Modify Credentials \(Preview\)](#)
- [Share Credentials \(Preview\)](#)
- [Delete Credentials \(Preview\)](#)
- [View Credential Details \(Preview\)](#)
- [View Credential Store History \(Preview\)](#)

About Credential Store (Preview)

The Credential Store in Oracle AI Data Platform allows you to create, manage, and provision access to credentials.

You access the Credential Store page from the left navigation panel in AI Data Platform. From the Credential Store screen, you can view all the existing credentials in your AI Data Platform instance and a history of changes to the credential store, including creation, modification, and deletion of credentials.

AI Data Platform enables you to create and store credentials for use as part of your notebooks and workflows. Credentials include keys, tokens, or passwords used to access sources outside of AI Data Platform, like clouds, databases, or APIs. The credential store provides a safe way to create and store credentials, manage permissions for who can access the credentials and where, and a way to view audit logs of how the credentials are used.

Credentials are managed with strict access controls to ensure secure and authorized use. Operations on the Credential Store are logged in the audit logs of AI Data Platform to ensure adherence to any applicable regulations and compliance requirements.

AI Data Platform Credential Store uses the following types of credentials:

- **Secret tokens:** This category of credentials enables you to access third-party services, like APIs. A secret token type credential is free form in nature and requires you to provide keys and value pairs to store field names and their values.
- **OCI vault references:** You can store references to existing OCI Vault secrets so that they can use utilities APIs to retrieve secrets from the vault. AI Data Platform does not store the secret value, but securely accesses it when necessary.
- **Service accounts:** Service accounts provide a way for you to provide credentials to trigger workflows.

Prerequisites

To create credentials in your AI Data Platform, you need the `CREATE_CREDENTIAL` permission at the Master Catalog level. For more information, see [Master Catalog Permissions](#).

You also need to ensure the following IAM policies are configured in the appropriate compartment:

For Vault References:

```
allow any-user to read secret-bundles in compartment id
<Secret_Compartment_OCID> where all { request.principal.type =
'aidataplatfrom', request.principal.id = target.secret.system-tag.orcl-
aidp.governingAidpId }
```

For Custom Encryption Keys:

```
allow any-user to use keys in compartment id <Key_Compartment_OCID> where
request.principal.type = 'aidataplatfrom'
```

Use Cases for Credential Store

The Credential Store can be used across a variety of scenarios:

- **External System Integration:** Store API tokens or credentials required to connect to third-party systems such as SaaS platforms, databases, or data sources.
- **Pipeline and Workflow Authentication:** Securely reference secrets in data pipelines, jobs, or workflows without hardcoding credentials.
- **Centralized Secret Management:** Maintain a single source of truth for credentials used across teams and environments.
- **Enterprise Vault Integration:** Use Vault References to leverage existing enterprise-grade secret management systems while still integrating with workflows.

Best Practices

To ensure security and maintainability when using Credential Store, follow these best practices:

- **Avoid Hardcoding Secrets:** Always store sensitive values in the Credential Store instead of embedding them in code, configs, or scripts.
- **Use Vault References When Possible:** For highly sensitive or regulated data, prefer vault references.
- **Limit Access with Principle of Least Privilege:** Share secrets only with users that require access, and avoid broad or unnecessary permissions.
- **Use Descriptive Naming Conventions:** Name secrets clearly (e.g., `openai-api-token-prod`) to make them easy to identify and manage.
- **Rotate Secrets Regularly:** Periodically update tokens and credentials to reduce the risk of compromise.
- **Separate Secrets by Environment:** Maintain distinct secrets for development, staging, and production environments to avoid accidental misuse.

- **Audit and Review Usage:** Regularly review who has access to secrets and where they are being used.

By using the Credential Store effectively, teams can improve security posture, simplify credential management, and enable scalable integrations across AI Data Platform.

Create Credentials (Preview)

You can create credentials for accessing other sources by selecting the a credential type providing the required details.

1. Click **Create** in the left navigation pane and select **Credentials**. You can also navigate to the **Credential Store** and click  **Create Credentials**.
2. Provide a name and description.
3. From the **Credential type** drop-down list, select the appropriate credential type.
 - Choose **Secret Token** to store a value directly.
 - Choose **Vault Reference** to reference an external vault secret.
 - Choose **Service Account** to create an account to trigger workflows.
4. Provide the necessary credentials in the fields provided.
 - For **Secret Token**, provide the key name and secret value.
 - For **Vault Reference**, provide the vault OCID.
 - For **Service Account**, provide the user OCID, fingerprint, region, tenancy OCID, and upload the private key file
5. Configure access or visibility settings, if applicable.
6. Click **Create**.

Use Stored Credentials in a Notebook (Preview)

You can call on stored credentials in the code of a notebook using the aidputils utility.

1. Navigate to your notebook.
2. Identify the field that requires a sensitive value. For example, an API key or password.
3. Use aidputils to get the stored credential:
 - For a secret value, use `myKey = aidputils.secrets.get(name=<<cred_name>>, key="key_name")`
 - For a vault reference, use `myKey = aidputils.secrets.get(name=<<cred_name>>, key=VaultSecretReference)`
 - For a service account, use the following to retrieve the service account keys:
 - `userOCID = aidputils.secrets.get(name=<<cred_name>>, key="userId")`
 - `fingerPrint = aidputils.secrets.get(name=<<cred_name>>, key="fingerprint")`
 - `tenancyOCID = aidputils.secrets.get(name=<<cred_name>>, key="tenancyId")`
 - `privateKey = aidputils.secrets.get(name=<<cred_name>>, key="privateKey")`

4. Run the notebook. Oracle AI Data Platform resolves the secret at runtime.

Modify Credentials (Preview)

You can modify the name, description, or configuration of credentials in your Credential Store to keep them up-to-date.

1. On the Home page, navigate to your Credential Store.
2. Next to the Credential you want to modify, click ... **Actions** then click **Edit**. You can also click the credential name, then click **Edit** in the top right.
3. Modify the name, description, or configuration details as needed.
4. Click **Save**.

Share Credentials (Preview)

You can share credentials in your Credential Store and manage who can access them.

1. On the Home page, navigate to your Credential Store.
2. Click the name of the credential you want to share.
3. Click the Permissions tab.
4. Add or modify permissions for the credential as needed.
5. Click **Save**.

Delete Credentials (Preview)

You can delete credentials from your Credential Store to remove unused or obsolete credentials.

1. On the Home page, navigate to your Credential Store.
2. Next to the Credential you want to modify, click ... **Actions** then click **Delete**. You can also click the credential name, then click **Delete** in the top right.
3. Select **Confirm deletion of credential**.
4. Click **Delete**.

View Credential Details (Preview)

You can view configuration, usage, and permission details for credentials in your Credential Store.

1. On the Home page, click **Credential Store** in the left navigation pane.
2. Click the name of the credential you want to view details for.
3. Click the Usage tab to see history of how and when a credential was used and by whom.
4. Click the Permissions tab to see which users or roles can access this credential and their access level.

View Credential Store History (Preview)

You can see a complete history of credentials in your Oracle AI Data Platform instance, including creation, modification, and deletion.

1. On the Home page, click **Credential Store** in the left navigation pane.
2. Click the History tab.
3. Use the **Type** drop-down list or the Search field to filter the displayed credential events.
4. Click an **Event ID** to see more details for that specific event.

Audit Log

Audit logs in Oracle AI Data Platform capture detailed records of user activities. These logs help monitor usage, ensure compliance, and investigate issues such as unauthorized access or configuration changes.



Audit logs are accessed from the **Audit logs** link in the left navigation pane of your AI Data Platform home page. You can track the operations of the following objects in your AI Data Platform:

- Auto-Populate extractors
- Compute clusters
- External Catalogs
- Roles
- Standard Catalogs
 - Schema
 - Tables
 - Views
 - Volumes
 - * Files
 - * Folders
- Workspaces
 - Workspace files
 - Workspace folders

Audit logs tracks the attributes of each object logged, including:

- Object name
- Object type
- Operation
- Operation request details
- Date, time, and timezone
- Details of the user that initiated the object (Initiated by)
- Details of where the object was initiated (UI, Terraform, SDK, API, Notebook, etc.)
- Status of the object (Success, failure)

Filters and search options enable you to narrow the parameters of which objects are displayed to help you locate logs for specific objects or periods of time.

Who can view Audit Logs?

Users assigned the AUDITOR role can view the entire audit trail. When a new AI Data Platform is created, the AUDITOR role is automatically created and the AI_DATA_PLATFORM_ADMIN role is made a member by default. Any users added to the AI_DATA_PLATFORM_ADMIN role are also automatically added as members of the AUDITOR role.

Users with administrator permissions for a specific object, such as Catalog, Schema, Tables, or Volumes, are able to view the audit logs for that object type.

Note

Log groups are created in AI Data Platform, but management of log groups, such as deleting log groups that aren't in use, takes place in Oracle Cloud Infrastructure. See [Logging Overview](#) for more info.

Search and Filter Audit Logs

You can search and filter records in audit logs to narrow the list of displayed logs and identify specific patterns, trends, or issues with your Oracle AI Data Platform objects.

1. On your home page, click **Audit Logs**.
2. Use the provided search fields to locate an object by name or payload.
3. Use the provided filters to help narrow your search by the associated value.

Auto Populate Catalog

This chapter contains information about creating and managing automated extractors to pull data into your catalogs.

Topics:

- [About Auto Populate](#)
- [Create Metadata Extractor](#)
- [Manually Review Extracted Metadata Entities](#)
- [View Reviewed Entities](#)
- [View Metadata Extractor Details](#)
- [Delete Metadata Extractor](#)

About Auto Populate

You can automate the process of extracting metadata from sources directly to your data catalogs.



[Video](#)

Manually creating schema, tables, and partitions from your data sources is time consuming and complicated. Oracle AI Data Platform offers the ability to automatically extract metadata from data sources and create entities in catalogs that you specify in the metadata extractor.

You automatically populate this metadata in your catalog by creating a metadata extractor. As part of creating the extractor, you specify the target catalog to extract metadata to and the source for the metadata. You can choose to have the extractor create tables in a specified schema, or let the system suggest where the tables are created if no schema is specified or detected.

Auto populate can extract metadata from the following file types:

- CSV
- JSON
- Avro
- ORC
- Parquet
- Delta Lake

Note

When you specify the source folder when creating a metadata extractor, all the files in the leaf folder must be the same data format.

Note

Auto-populate only supports underscores (_) as special characters in column names.

You can opt to either manually review entities that are extracted or let the system automatically create the entities from the extracted metadata. When extracting metadata, entities that cause errors are captured in the log. You can view the log to see which entities encountered errors and take action to correct.

Manually reviewing entities allows you to accept or reject entities on an individual basis. You can view entities already approved or rejected in the Reviewed Entities tab.

Extractors display their status to let you know what stage they are currently at and if user intervention is required.

Extractor Status	Description
Not Started	The extractor has not started. Start the extractor to begin.
Running	Extractor is in progress
Ready for review	The extractor has run and you have chosen manual approval. Extracted entities must be reviewed and either accepted or approved.
Reviewing	The extractor has run and you have chosen manual approval. Some entities have been reviewed or approved by a user, but entities remain that require review.
Completed	The extractor has run and entities have either been approved automatically or manually approved by a user

You can view and use metadata extractors created by other users if you have the requisite permissions.

Create Metadata Extractor

You can create metadata extractors to automate extracting entities like schema and tables to your catalogs.

1. On the Home page, click **Auto populate catalog**.
2. Click  **Create Metadata Extractor**.
3. Enter a name for the metadata extractor.
4. Select the target catalog from the **Catalog** dropdown.
5. Select the appropriate source type from **Source Type** dropdown.
6. Next to **Compute**, click **Browse** and choose the cluster the extractor should use. Click **Select**.
7. For **Object Storage URI**, select whether to browse your compartments for the bucket or folder you want to extract metadata to or if you want to specify the URI for the bucket or folder.
 - For **Select bucket or folder**, click **Browse** to select the compartment, bucket, and folder you want metadata extracted to.

- For **Enter URI manually**, enter the URI in the field provided.
8. Select whether entities are created with manual approval or automatically approved by the system.
 9. Optional: Select the schema where external tables are created. If no schema specified, the system creates tables in schema based on folder structure, or in the default schema if no schema is detected.

Manually Review Extracted Metadata Entities

When you choose the manual method of creating entities in a metadata extractor, you need to review the extracted entities and approve or reject adding them to your catalog.

1. On the Home page, click **Auto populate catalog**.
2. Click the name of the metadata extractor.
3. Click the **Entities awaiting review** tab.
4. For each entity, select **Approve** or **Reject**.
5. Optional: Select **Approval All** or **Reject All** to set all entities under review to the selected status.
6. Click **Submit**.

View Reviewed Entities

You can see entities that have been manually or automatically reviewed as part of metadata extraction and see log details, table details, or column schema for that entity.

1. On the Home page, click **Auto populate catalog**.
2. Click on the name of the metadata extractor.
3. Click the **Reviewed entities** tab.
4. Next to an entity, click ... **Actions**.
 - Click **View table details** to see the table details for the selected entity.
 - Click **View column schema** to see the column schema for the selected entity.
 - Click **View logs** to see the metadata extractor logs for the selected entity.

View Metadata Extractor Details

You can view the details of a metadata extractor to see its status, metadata creation method, base location, and creation details.

1. On the Home page, click **Auto populate catalog**.
2. Click the name of the metadata extractor.
3. Click the **Details** tab.

Delete Metadata Extractor

You can delete metadata extractors that are no longer needed.

1. On the Home page, click **Auto populate catalog**.

2. Next to the metadata extractor you want to delete, click ... **Actions** and click **Delete**
3. Click **Delete**.

Share Data

This chapter provides details on managing data shares, assets, and recipients in your Oracle AI Data Platform.

Topics:

- [Data Sharing](#)
- [Create a Share](#)
- [Modify a Share](#)
- [Delete a Share](#)
- [View Share Details](#)
- [Add an Asset to a Share](#)
- [Remove an Asset from a Share](#)
- [Add Recipients to Data Sharing](#)
- [Add an Existing Recipient to a Share](#)
- [Manually Activate a Recipient](#)
- [Resend Activation Link to a Recipient](#)
- [Remove Recipient from a Share](#)
- [Modify a Recipient](#)
- [View Recipient Details](#)
- [Delete a Recipient](#)

Data Sharing

Data Sharing in Oracle AI Data Platform enables secure and efficient real-time data sharing across and within the organization without data duplication.



AI Data Platform uses Delta Sharing protocol to securely share data. To share your data, you must complete these steps:

1. Create a share.
2. Add assets that you want to share. An asset can be a schema or table.
3. Create and add recipients to a share. A share and recipient can be access controlled using permissions model.
4. Activate the recipients.

From the Data Sharing page you can see all Shares in your AI Data Platform, all recipients of Shares in your AI Data Platform, your Shares, and Shares others have made you a recipient for.

From the **Shares** tab, you can see the details of Shares in your AI Data Platform, including the number of assets in a Share, number of recipients, and who owns the share. You can create a new Share by clicking **Create share** next to the search bar.

Clicking on a Share lets you see the shared assets, recipients, details, and RBAC permissions. From the **Assets** tab you can add assets to the share. From the Recipients tab you can manage the recipients for this specific share, create new recipients, or share with existing recipients. The Details tab provides the name, description and creation details of the Share. From the Permissions tab you can manage the RBAC permissions for that Share.

From the **Recipients** tab at the top level of Data Sharing, you can see details of the Share recipients in your AI Data Platform. You can create new recipients by clicking **Create recipient** next to the search bar.

Clicking on a recipients name shows Shares the recipient is added to, details for that recipient, and the RBAC permissions for that recipient.

Consuming a Delta Share

You can consume a Delta Share in AI Data Platform using Python by importing the `delta_sharing` library. Before running the code, the user must activate the Delta Share and save the generated profile file in the workspace. The profile file is used by the Delta Sharing client to authenticate and access the shared tables.

```
import delta_sharing

# Path to your downloaded credentials file
profile_file = "<<path to the share file in workspace>>"

# Initialize the client
client = delta_sharing.SharingClient(profile_file)

#List all schemas and tables available to you
tables = client.list_all_tables()
print("Available Tables:", tables)
```

Handling Token Expiry

AI Data Platform can regenerate bearer tokens for data sharing for active recipients. If the current bearer token has not expired, you can extend the duration of the bearer token by calling the `updateSecretExpiry` API using a curl command:

```
curl
  --header 'Authorization: Bearer <token>'
  --header 'Content-Type: application/json'
  --data '{"existingTokenExpireInSeconds": "<duration>"}'
  --request POST '<delta_share_endpoint>/updateSecretExpiry'
```

Where:

- `<token>` is the bearer token value in the share profile
- `<duration>` is the duration of the new token in seconds
- `<delta_share_endpoint>` is the Delta Share endpoint. You can get this value for a recipient by using `sql("DESCRIBE recipient <<recipient_name>>")` and checking the value of the recipient's `delta_share_endpoint`.

Limitations

- You can send the activation token email using an email application or by copying the activation link and composing an email using their existing email service.
- You cannot share volumes or workspace files (data, notebook, sql or python files) using data sharing.

Create a Share

You must create a Share as a container for managing shared assets as part of Oracle AI Data Platform data sharing.

1. On the Home page, click **Data Sharing**.
2. Click  **Create Share**.
3. Enter a name and description for your Share.
4. Click **Create**.

Modify a Share

You can update the name and description of your Shares after creation.

1. On the Home page, click **Data Sharing**.
2. Next to the share to modify, click  **Actions** then click **Edit**.
3. Enter a new name or description for your Share. Click **Save**.

Delete a Share

You can delete a Share that is no longer needed.

1. On the Home page, click **Data Sharing**.
2. Next to the share to delete, click  **Actions** then click **Delete**.
3. Click **Delete**.

View Share Details

You can view details about a share, such as assets, who has access to the share, and more.

1. On the Home page, click **Data Sharing**.
2. Click to the **Shares** tab.
3. Click the share name and click the **Details** tab.

Add an Asset to a Share

After creating a Share, you can add assets to be shared with recipients.

1. On the Home page, click **Data Sharing**.
2. Click the name of the Share to add assets to.

3. On the Assets tab, click  **Add asset**.
4. Select **Workspace** or **Catalog** as the asset source.
5. Use the search bar or navigate to the asset you want to share. Select it and click **Add**.

Remove an Asset from a Share

You can remove assets you no longer want to be a part of a Share.

1. On the Home page, click **Data Sharing**.
2. Click the name of the Share to remove assets from.
3. On the **Assets** tab, next to asset you want to remove click  **Actions** then click **Delete**.
4. Click **Delete**.

Add Recipients to Data Sharing

You must add recipients before you can include them in Shares.

1. On the Home page, click **Data Sharing**.
2. Click the **Recipients** tab.
3. Click  **Create recipient**.
4. Enter a name and description for the recipient.
5. Optional: Enter the email address to send an activation link to for the recipient.
6. Click **Create**.

Add an Existing Recipient to a Share

You can add recipients that already exist in your Oracle AI Data Platform to a Share.

1. On the Home page, click **Data Sharing**.
2. Click the name of a Share to add an existing recipient to.
3. Click the **Recipients** tab.
4. Click  **Share with existing recipients**.
5. Search for a user and select the box next to the recipient name.
6. Click **Share**.

Manually Activate a Recipient

You can choose to manually activate recipients if there are issues preventing them from using the email activation links.

1. On the Home page, click **Data Sharing**.
2. Click to the **Recipient** tab.
3. Click the recipient name and click the **Details** tab.
4. Next to **Activation status** click **Activate**.
5. Copy the activation link from the prompt and paste into a separate window.

6. Verify the Activation status is now set to **Active**.

Resend Activation Link to a Recipient

If the activation link was lost, deleted, or otherwise not received by a recipient, you can resend the activation link to the provided email address.

The recipient must have an email address as part of their details to send them an activation link. If the email is incorrect or missing, you may need to create a new recipient.

1. On the Home page, click **Data Sharing**.
2. Click to the **Recipient** tab.
3. Click the recipient name and click the **Details** tab.
4. Next to **Activation email link** click **Send activation link**. Your configured email application opens.
5. Review and send the email to the recipient.

Remove Recipient from a Share

You can remove a recipient from a Share, keeping the recipient but removing their access to assets in that Share.

1. On the Home page, click **Data Sharing**.
2. Click the share you want to remove recipients from.
3. Click the **Recipients** tab.
4. Next to the recipient you want to remove, click ... **Actions** then click **Remove**.

Modify a Recipient

You can update the name, description, and email address for a recipient after creation.

1. On the Home page, click **Data Sharing**.
2. Click to the **Recipient** tab.
3. Next to the recipient to update, click ... **Actions** then click **Edit**.
4. Enter a new name or description for your recipient. Click **Save**.

View Recipient Details

You can see information on a recipient, such as name, email address for activation link, active status, and creation details from the Details tab.

1. On the Home page, click **Data Sharing**.
2. Click to the **Recipient** tab.
3. Click the recipient name and click the **Details** tab.

Delete a Recipient

You can delete recipients that are no longer needed in your Oracle AI Data Platform.

1. On the Home page, click **Data Sharing**.
2. Click the **Recipients** tab.
3. Next to the recipient to delete, click ... **Actions** then click **Delete**.
4. Click **Delete**.

Part VII

DevOps

Oracle AI Data Platform provides DevOps tools for integrating your data with Git and CI/CD.

Topics:

- [Git Integration \(Preview\)](#)
- [Oracle AI Data Platform \(AIDP\) Developer Extension](#)

Git Integration (Preview)

You can connect your Oracle AI Data Platform instance to ingest, read, and write data to and from external Git repositories.

Git integration in AI Data Platform works by providing login credentials for Git accounts and creating folders that connect to Git repositories with those credentials to ingest data. You can store multiple sets of credentials to connect to different Git repositories or manage different levels of access to the same repositories. Administrators can select a set of credentials as the default profile for users and delete credentials when they're no longer needed.

You can manage your Git files entirely from within AI Data Platform, with utilities that allow you to push, pull, commit, and merge files as well as manage any merge conflicts that arise without having to leave the platform.

AI Data Platform also enables you to branch code from your Git repository as well as switch, merge, compare, and delete branches.

Note

AI Data Platform currently supports Git repositories hosted on public Git providers such as GitHub, GitLab, Bitbucket, and Oracle Cloud Infrastructure DevOps. Support for Git hosted on customer private networks will be available in a future update.

Git Folders in Oracle AI Data Platform (Preview)

Git folders in Oracle AI Data Platform allow you to interact with data hosted in a Git repository without leaving the platform.

Git folders are created in your Master Catalog using credentials you set up in your AI Data Platform settings. You connect your folder to a Git repository and use the AI Data Platform UI to push, pull, merge, rebase, and reset files in your repository.

You can manage RBAC permissions for your Git folder to control who can see the folder and what Git operations they are able to perform. For more information, see [Git Folder Permissions \(Preview\)](#).

Create a Git Folder (Preview)

You can create a folder in your Oracle AI Data Platform workspace connected to an external Git repository.

1. On the home page, navigate to your workspace.
2. Click  **Create** then click **Git folder**. You can also click **Add** in the top right and click **Git folder**.
3. From the **Select Git credentials** drop-down list, select the Git credentials required to access the Git repository.
4. Provide the Git repository URL. For example: `example.git.path`.

5. Provide a name for the Git folder in your workspace.
6. Provide the branch from the Git repository to connect to your folder.
7. Click **Create**.

Edit Git Folder (Preview)

You can change the name and description for a Git folder in your workspace.

1. On the home page, navigate to your workspace.
2. Next to the Git folder want to modify, click ... **Actions** then click **Edit**.
3. Modify the name and description for the folder.
4. Click **Save**.

Delete Git Folder (Preview)

You can delete Git folders your own from your workspace.

Deleting a Git folder deletes the folder and any changes that have not been committed and pushed to the repository.

1. On the home page, navigate to your workspace.
2. Next to the Git folder want to modify, click ... **Actions** then click **Delete**.
3. Click **Delete**.

Modify Git Folder Settings (Preview)

You can modify the credentials and repository URL used by a Git folder in your workspace.

1. On the home page, navigate to your workspace.
2. Click the Git folder you want to modify settings for.
3. Click the Settings tab.
4. From the drop-down list, select the new credentials for the Git folder.
5. Enter the new Git repository URL.
6. Click **Save**.

View Git Folder History (Preview)

You can view the history of commits to your Git folder from the History tab.

1. Navigate to your Git folder and click the Git tab.
2. Click the History tab.
3. Click the **Git** button to view the history of commits on the Git site.

Create Git Branch (Preview)

You can create new Git branches from existing branches from Git folders in your workspace.

1. Navigate to your Git folder and click the Git tab.

2. From the **Current branch** drop-down list, select the branch you want to base your new branch on.
3. Click **Create Branch**.
4. Provide a name for your new branch,
5. Click **Create**.

Delete Git Branch (Preview)

You can delete Git branches from your Git repository from inside your Oracle AI Data Platform workspace.

1. Navigate to your Git folder and click the Git tab.
2. From the **Current branch** drop-down list, select the branch you want to delete.
3. Click **Delete branch**.
4. Select all branches you want to delete.
5. Optional: Select **Force delete regardless of merge status** to ignore warnings for any branches that have not been merged.
6. Click **Delete**.

Git Pull (Preview)

You can select a branch a perform a Git pull from inside a Git folder in your workspace.

1. Navigate to your Git folder and click the Git tab.
2. Click the Pull tab.
3. From the drop-down list, select the branch you want to pull from.
4. Click **Pull**.
5. Resolve any conflicts.
 - Click **Keep local** to overwrite with the local version.
 - Click **Keep remote** to overwrite with the remote version.
 - You can also make your own changes to resolve, then click **Mark as resolved**.
 - Click **Abort pull** to cancel the pull request.
6. Enter a commit message and, optionally, a description for the pull.
7. Click **Commit pull**.

Git Push (Preview)

You can select a branch a perform a Git push from inside a Git folder in your workspace.

1. Navigate to your Git folder and click the Git tab.
2. Click the Push tab.
3. Select the files you want to push to the Git repository. Click a file to see the changes in the code window.
4. Enter a commit message and, optionally, a description for the push.
5. Click **Commit push**.

Merge Git Branches into Main (Preview)

You can merge branches back in the main trunk from inside a Git folder in your workspace.

1. Navigate to your Git folder and click the Git tab.
2. Click the Merge tab.
3. From the drop-down list, select the branch you want to merge with the main.
4. Click **Merge**.
5. Resolve any conflicts.
 - Click **Keep local** to overwrite with the local version.
 - Click **Keep remote** to overwrite with the remote version.
 - You can also make your own changes to resolve, then click **Mark as resolved**.
 - Click **Abort merge** to cancel the merge request.
6. Enter a commit message and, optionally, a description for the merge.
7. Click **Commit merge**.

Git Rebase (Preview)

You can rebase a branch in your Git folder to apply changes from that Git branch to another branch.

Note

Rebasing a branch requires a force push. Force pushing alters the history on the remote repository and can cause versioning issues for collaborators working on the same repository.

1. Navigate to your Git folder and click the Git tab.
2. Click the Rebase tab.
3. From the **Current branch** drop-down list, select the branch to rebase.
4. From the **Rebase to branch** drop-down list, select the branch to rebase the current branch to.
5. Select **Confirm rebase**.
6. Click **Rebase**.

Git Reset (Preview)

You can reset a branch in your Git folder to revert the content of that branch to a previous state.

① Note

Resetting a branch requires a force push. Force pushing alters the history on the current branch both locally in your workspace and the remote repository. The history of the branch is deleted and reset to the history of the remote reset branch.

1. Navigate to your Git folder and click the Git tab.
2. Click the Reset tab.
3. From the **Current branch** drop-down list, select the branch to reset.
4. From the **Choose the remote branch you want the current branch "main" to be reset to** drop-down list, select the branch in the remote Git repository to reset the current branch to.
5. Select **Confirm hard reset**.
6. Click **Reset**.

Bundles and CI/CD (Preview)

Bundles define how Oracle AI Data Platform resources are packaged, versioned, shared, and deployed across workspaces and environments.

You can create bundles for jobs and agents. A bundle is created in a Git folder so that the bundle files can be committed and pushed to a Git repository. The resources that you include in the bundle can come from any location in the workspace. If the selected resources reference supporting files or dependencies, such as notebooks, scripts, or compute, those references are included in the bundle so the resources can be deployed in another environment without manually recreating the same components.

You create bundles from your workspace and select the AI Data Platform resources to include at the time of creation. Existing bundles can be modified to add or remove resources. Bundling resources allows you to group related jobs and agents together and deploy them to other environments with their dependencies intact.

You can commit and push bundles to your Git repository through your Git folders. Users can then pull the bundles into their own environments and deploy them. Changes to bundle files are available to other environments after the changes are committed and pushed to the Git repository. Whenever users pull the updated bundle from the Git repository, they receive the latest committed resource files.

Bundle Concepts

Concept	Meaning	User Action
Git folder	An AI Data Platform folder connected to a Git repository and branch. Bundles can be created only in Git folders.	Create or open a Git folder before creating a bundle.
Bundle	A deployable package that defines selected resources and their dependencies.	Create the bundle from the workspace, then select resources such as jobs or agent flows.
Bundle content	Generated files that describe resources, variables, target settings, and artifacts required for deployment.	Review generated content after create or sync. Commit and push the bundle through Git when ready.
Target	An environment or workspace configuration used when deploying the bundle outside the source environment.	Provide values that must change by environment, such as workspace key, credential name, token key, or parameter defaults.
Deployed item	A resource created or updated by a bundle deployment in the target workspace.	Use the Deployment tab to deploy, review deployed items, run jobs, and purge deployed resources when needed.
Sync	A bundle action that refreshes bundle resources and dependencies from the latest source resources.	Run sync after changing bundled resources, then commit the refreshed bundle files to Git.

Before You Start

Check this list to make sure you have everything you need to start using Git bundles in your AI Data Platform workspace:

- Create Git credentials and a Git folder that points to the repository and branch used for bundle files.
- Confirm the source resources are in the workspace and are saved before creating or syncing the bundle.
- For cross-environment deployment, confirm the target workspace has required permissions and compute capacity.
- For jobs or notebooks that use secrets, create the required named credentials in the target credential store. Bundle files should reference credentials by name or variable, not contain secret values.
- Commit and push bundle file changes after creating or syncing a bundle so other environments can pull the latest bundle definition.

Bundle Publish Paths

Bundles can publish runtime artifacts to a workspace folder that users can inspect. Runtime artifacts include files such as notebooks, scripts, libraries, and agent flow files that are needed by deployed bundle resources. The folder where artifacts are published is called the publish path. Artifacts are stored under an artifacts folder inside the publish path.

For example, if the publish path is: `/Workspace/Shared/customer_churn_bundle` the bundle artifacts are published under: `/Workspace/Shared/customer_churn_bundle/artifacts`.

The publish path is not the source of truth for the bundle. The bundle source files remain in the bundle folder and should be committed to Git. Files under the publish path are runtime files used by deployed resources. If you manually edit files under the publish path, those edits are not written back to the bundle source. A later deployment can replace the published artifacts with files from the bundle.

Create a Bundle (Preview)

You can bundle Oracle AI Data Platform job and agent resources to create a package that can be deployed to another Git folder.

1. Navigate to your workspace.
2. Click **Actions** then click **Create Bundle**.
3. Provide a name and description for your bundle.
4. Click **Browse** and select the Git folder in your workspace where you want to create your bundle.
5. Select the resources you want to include in the bundle from the listed options.
6. Click **Create**.

Note

Bundles are created in Git folders. If the selected location is not inside a Git folder, you cannot use Git to move the generated bundle files to another environment.

Note

When a bundle is created, AI Data Platform can include a default publish configuration in the bundle manifest. The publish configuration controls where runtime artifacts are placed when the bundle is deployed.

What Gets Generated After Bundle Creation (Preview)

After bundle creation, Oracle AI Data Platform writes your generated files into the bundle folder.

The exact file names vary by resource type, but the created bundle usually contains a top-level bundle manifest, resource definitions, artifact folders, and optional override files.

Generated Item	Description
<code>aidp_workbench.yaml</code>	Bundle manifest. Identifies the bundle, included resources, variables, and target sections used during deployment.
<code>.aidp/resource_origins.yaml</code>	Resource definitions. Describe resources such as jobs, notebook tasks, compute references, or agent flows.
<code>.aidp/overrides.yaml</code>	Override files or target sections. Hold values that differ by environment, such as workspace identifiers, parameter defaults, credential names, or token keys.
<code>.aidp/aidp.state.json</code>	Tracks resources deployed from the bundle. Used by future deploy and purge operations.
<code>jobs/</code>	Job descriptors and job dependencies.
<code>agentflows/</code>	Agent descriptors and agent dependencies.
<code>artifacts/</code>	Artifact folders. Copied code, notebooks, and supporting file artifacts.

Publish Path Configuration (Preview)

You can configure the publish path for a bundle in the bundle manifest file, `aidp_workbench.yaml`.

To configure the publish path for a bundle file, you add a publish block under `defaults` in the `aidp_workbench.yaml` manifest file, as in this example:

```
bundle:
  name: customer_churn_bundle
  description: Customer churn model workflow

resources:
  jobs:
```

```

    - train_churn_model
  agentflows:
    - explain_churn_predictions

defaults:
  publish:
    path: "/Workspace/Shared/customer_churn_bundle"
    overwrite_publish_artifacts: false

```

In this example, when the bundle is deployed, artifacts are published under `/Workspace/Shared/customer_churn_bundle/artifacts`.

Table 34-1 Publish Path Configuration Fields

Field	Description
path	Workspace folder used as the publish root. Artifacts are stored under <code><path>/artifacts</code> .
overwrite_publish_artifacts	Indicates whether deployment can replace existing published artifacts at the configured publish path.

Publish Path Configuration for Different Environments

You can configure different publish paths for different deployment targets, as in this example:

```

bundle:
  name: customer_churn_bundle

resources:
  jobs:
    - train_churn_model

defaults:
  publish:
    path: "/Workspace/Shared/customer_churn_bundle"
    overwrite_publish_artifacts: false

targets:
  prod:
    - aidp_ocid: "ocid1.aidataplatfom.oc1..example"
      workspace_key: "workspace-prod"
      publish:
        path: "/Workspace/Production/customer_churn_bundle"
        overwrite_publish_artifacts: true

```

Alternatively, you can use `.aidp/overrides.yaml` to override the publish path locally:

```

publish:
  path: "/Workspace/Users/me/customer_churn_bundle"
  overwrite_publish_artifacts: true

```

When multiple publish paths are configured, Oracle AI Data Platform resolves the publish path in this order:

1. `defaults.publish.path` in `aidp_workbench.yaml`

2. Matching target publish path
3. `.aidp/overrides.yaml` publish path

How Bundle Artifacts are Published (Preview)

When you deploy a bundle, Oracle AI Data Platform copies the bundle artifacts from the bundle source into the resolved publish path.

The default artifact layout is `<publish.path>/artifacts`. For example:

```
/workspace/Shared/customer_churn_bundle/artifacts
```

Bundle descriptor files can reference published artifacts by using `${publish.path}/artifacts/...`. For example:

```
${publish.path}/artifacts/Notebooks/train_model.ipynb
```

During deployment, `${publish.path}` is resolved to the configured publish path.

Deploy a Bundle (Preview)

You can deploy bundles from a Git folder to share resources and dependencies across workspaces and environments.

1. Navigate to the bundle you want to deploy in your workspace.
2. Click the Deployment tab.
3. Optional: Review target values and overrides, if applicable.
4. Click **Deploy**. You are notified when the deployment is complete.

Note

On first deployment, the Deployment tab may show no deployed items until deployment succeeds. After a successful deployment, review the deployed resources, run jobs or agents as needed, and confirm runtime output.

Note

During deployment, bundle artifacts are published to the configured publish path before resources are deployed. If no custom publish path is configured, Oracle AI Data Platform uses the default publish location for the bundle.

How Bundles are Deployed to Published Paths (Preview)

When a bundle has a publish path configured, deployment performs a fixed set of actions.

The actions Oracle AI Data Platform performs when deploying a bundle with a configured publish path take place in this order:

1. Validate the publish path.

2. Publish bundle artifacts to `<publish.path>/artifacts`.
3. Deploy bundle resources, such as jobs and agent flows, using the published artifact paths.
4. Update deployment status after deployment completes.

Different actions may take place if AI Data Platform encounters any issues with the publish path:

- If the publish path does not exist, AI Data Platform can create it during deployment.
- If the publish path exists and is empty, deployment can use it.
- If the publish path contains artifacts managed by the same bundle, deployment can update them.
- If the publish path contains data that is not owned by the bundle, deployment can fail unless overwrite is enabled.

Sync a Bundle (Preview)

You can sync a bundle to update the resources and dependencies in the bundle with the latest changes.

Sync uses the bundle's recorded origin metadata to rebuild the bundle from the source jobs and agent flows that were captured when the bundle was created. The source metadata is stored in `.aidp/resource_origins.yaml` and must match the requested instance and workspace. The operation refreshes source-controlled bundle content while preserving the bundle identity and runtime metadata.

During sync, the service stages a refreshed bundle snapshot under the bundle `.aidp` directory, compares existing and staged descriptors, preserves existing variable aliases and override references where possible, merges existing manifest default variables, and then promotes the refreshed source-controlled files back into the bundle root.

Sync preserves environment-specific and deployment runtime files such as `.aidp/overrides.yaml` and `.aidp/aidp.state.json`. These files are not replaced by the refreshed source snapshot.

You should sync your bundles when:

- A bundled notebook, job, agent flow, compute reference, parameter, or dependency changed in the source workspace.
- A bundled job now points to a different notebook or has updated task parameters.
- An agent configuration, prompt parameter, model setting, or AI compute dependency changed.
- You need the Git folder to contain the latest generated bundle files before pushing to another environment.

1. Navigate to the bundle you want to sync in your workspace.
2. Click **Actions**.
3. Click **Sync**. You are notified when the sync is complete.

Recommended Promotion Workflow for Bundles (Preview)

We recommend following a workflow of develop, package, version, pull, configure, deploy, and validate for bundles in Oracle AI Data Platform.

Phase	Action	Result
Develop	Create or update notebooks, jobs, compute settings, or agent flows in the source workspace.	The source resources contain the latest intended behavior.
Package	Create the bundle, or run Sync on an existing bundle.	Generated bundle files reflect the current source resources and dependencies.
Version	Review, commit, and push the bundle folder through the Git folder.	The remote repository contains the deployable bundle definition.
Pull	In the target workspace, pull the Git folder changes.	The target workspace has the latest bundle files.
Configure	Set target-specific values such as workspace key, parameter defaults, credential name, token key, or runtime settings.	The bundle is ready for the target environment.
Deploy	Open the bundle Deployment tab and click Deploy.	AI Data Platform instance creates or updates the deployed resources.
Validate	Run deployed jobs or test deployed agent flows and inspect logs or output.	The target environment is verified before promotion continues.

Bundle Variables (Preview)

Use bundle variables when a bundle needs deploy-time values that can change for every target environment without editing the resource definition.

To begin using variables with bundles, you need to define the variable names in the bundle's `aidp_workbench.yaml` file with default values and, when needed, environment-specific overrides. You then reference those variables from generated definition files.

Example Variable Workflow

1. Define variable names under `defaults > variables` in `aidp_workbench.yaml`.
2. Optionally define environment-specific values under `targets > <environment> > variables` in the same `aidp_workbench.yaml` file.
3. Use target overrides to set deploy-time values such as a dev region, QA region, production region, schema name, credential name, token key, or parameter default.
4. When no matching target environment is found, Oracle AI Data Platform uses the default variable value from `aidp_workbench.yaml`.
5. Use the variable in any generated definition file, such as a job definition, compute definition, task definition, or other YAML definition file.
6. Reference the variable with `${var.<<variable_name>>>}`, replacing `<<variable_name>>` with the actual variable name defined in `aidp_workbench.yaml`.
7. Use resource-specific placeholder syntax in places such as agent prompts or tool definitions when the UI expects a named placeholder.
8. For temporary deployment-only changes, create `.aidp/overrides.yaml` and define variable values there.

Bundle Variable Syntax

Variable definitions must be written in `aidp_workbench.yaml`. Other bundle definition files should only reference the variable name.

Definition	Syntax	Example
Variable definition in <code>aidp_workbench.yaml</code>	<pre>defaults: variables: <<variable_name>>: <<default_value>></pre>	<pre>defaults: variables: deployment_region: "us-ashburn-1"</pre>
Target-specific override in <code>aidp_workbench.yaml</code>	<pre>targets: <<target_name>>: variables: <<variable_name>>: <<target_value>></pre>	<pre>targets: qa: "workbenchid": "ocid" "workspaceKey": "key" variables: deployment_region: "us- phoenix-1"</pre>
Variable reference in a definition file	<code>\${var.<<variable_name>>}</code>	<code>value: "\${var.deployment_region}"</code>
Named placeholder in an agent prompt or tool field	<code>{{variable_name}}</code>	Generate a summary about <code>{{topic}}</code> .

Example: Job Parameter Variable

The following example defines one variable and passes it into a job parameter. The resource can remain unchanged while the variable value changes in the bundle configuration.

```
bundle:
  name: namedVariableJobDemo
resources:
  jobs:
    - namedVariableJobDemo
defaults:
  variables:
    job_param_value_default: "cloud data governance"

jobs:
  namedVariableJobDemo:
    tasks:
      - taskKey: Task_956691
        parameters:
          - name: param_key
            value: "${var.job_param_value_default}"
```

A notebook task can read the resolved value through the parameter name used by the job:

```
param_value = odiUtils.parameters.getParameter("param_key", "fallback value")
print(param_value)
```

Example: aidp_workbench.yaml Variable Definition

This example defines the variable `agent_topic_default` in `aidp_workbench.yaml`. The default value applies when no matching target environment is found, or when no selected environment or temporary override supplies the variable. Each target can override the same variable name for that workspace.

```
defaults:
  variables:
    agent_topic_default: "Oracle AI Data Platform"

targets:
  dev:
    - aidp_ocid: "ocid1.aidatapatformdev.oc1.iad.exampledev"
      workspace_key: "6271b138-6f2c-4e8c-a9e5-013e17079955"
      variables:
        agent_topic_default: "AIDP Dev Workspace"
  qa:
    - aidp_ocid: "ocid1.aidatapatformdev.oc1.iad.exampleqa"
      workspace_key: "daa0e9b5-f8d7-4ae8-b728-2984e68a74c0"
      variables:
        agent_topic_default: "AIDP QA Workspace"
  prod:
    - aidp_ocid: "ocid1.aidatapatform.oc1.iad.exampleprod"
      workspace_key: "83be5b8c-bb8d-4ec4-a8fb-16b50b0e4d4c"
      variables:
        agent_topic_default: "AIDP Production Workspace"
```

Any bundle definition file that supports bundle variables can then reference the same variable name. For example:

```
tasks/agentTopic.task.yaml

parameters:
  - name: topic
    value: "${var.agent_topic_default}"
```

Example: File Names and Resolution

In a generated bundle, the variable definition and the resource that uses the variable can be in different files. When you promote or copy the example through Git, include the whole bundle folder so the manifest, job definition, and artifacts stay together.

```
testNamedVarBundleJob/
  aidp_workbench.yaml
  jobs/
    namedVariableJobDemo.job.json
  artifacts/
    namedVarNotebookDemo.ipynb
```

File	What It Contains	Variable or Resolved Value
aidp_workbench.yaml	Defines the variable under defaults > variables and optional target overrides.	Variable: job_param_value_default.
jobs/namedVariableJobDemo.job.json	Uses the variable as the value for job parameter param_key.	Reference: \${var.job_param_value_default}.
artifacts/namedVarNotebookDemo.ipynb	Reads the resolved job parameter at runtime.	Notebook reads param_key after resolution.

For example, the bundle manifest can define the default and the target overrides in aidp_workbench.yaml:

```
defaults:
  variables:
    job_param_value_default: "cloud data governance"
targets:
  dev:
    variables:
      job_param_value_default: "space tourism"
  qa:
    variables:
      job_param_value_default: "data quality validation"
  prod:
    variables:
      job_param_value_default: "customer operations"
```

The job definition can then use that same variable name in jobs/namedVariableJobDemo.job.json. The same syntax can be used in job definitions, compute definitions, task definitions, and other generated YAML definition files that support bundle variables:

```
"parameters": [
  {
    "name": "param_key",
    "value": "${var.job_param_value_default}"
  }
]
```

Deployment Target	Variable Resolution	Notebook Receives
No matching target	Uses default from aidp_workbench.yaml.	param_key = cloud data governance
dev	Uses dev target override.	param_key = space tourism
qa	Uses qa target override.	param_key = data quality validation
prod	Uses prod target override.	param_key = customer operations

Resolution Order

When `.aidp/overrides.yaml` supplies a variable value, AI Data Platform uses that temporary value for the next deployment. Otherwise, AI Data Platform resolves the variable from the selected target first. If no matching target environment is found, or the selected target does not override the variable, AI Data Platform falls back to `defaults > variables` in `aidp_workbench.yaml`. If no value is available for the referenced variable, deployment or job execution can fail before the notebook fallback value is used.

Temporary Overrides for the Next Deployment

The definitions in `aidp_workbench.yaml` are sufficient for normal bundle promotion. If a user needs to override a variable temporarily, create `.aidp/overrides.yaml` in the bundle folder and define the variable values there. AI Data Platform uses those values during the next deployment.

```
.aidp/  
  overrides.yaml
```

```
variables:  
  agent_topic_default: "tuesday"
```

File	Example Value	When Used
<code>aidp_workbench.yaml</code> <code>defaults</code>	Oracle AI Data Platform	Used when no matching target or override file supplies the variable.
<code>aidp_workbench.yaml</code> <code>target dev</code>	AIDP Dev Workspace	Used when deploying to dev and no temporary override is present.
<code>.aidp/overrides.yaml</code>	tuesday	Used as a temporary value during the next deployment.

Note

Use `.aidp/overrides.yaml` for temporary deployment values. Keep durable defaults and environment-specific values in `aidp_workbench.yaml` so the bundle remains reproducible through Git.

Example: Agent Prompt Variable

For agent flow prompt text or tool definitions that support named placeholders, use double braces around the variable name.

Prompt text:

You are a demo assistant. For the given `{{topic}}`, generate 3 short bullet points explaining why it matters.

Tool parameter:

Name: `/{{topic}}`

Default value: sunday

Note

As a best practice, you should use variables for values that are expected to change across deployments. Avoid embedding environment-specific values directly into notebooks, jobs, or agent flow text when a variable can carry the value instead.

Example: Workflow without Compute

You can create a bundle for a workflow that doesn't automatically create a compute cluster as part of deployment. First, create a variable in `aidp_workbench.yaml`:

```
defaults:
  variables:
    job_compute_key: "select_compute"
```

Next, in the `jobs/` folder, locate the job file `<JobName>.job.json`. Edit the job file replace the `"clusterKey"` value with the new created variable:

```
"clusterKey" : "${var.job_compute_key}",
```

Named Variables in Environments (Preview)

Named environment variables enable your Git bundle to keep the same logical variable names while each environment supplies its own deploy-time values.

Use this pattern when a dev region, QA region, and production environment should deploy the same resource definitions with different values.

Example Named Variable Workflow

1. Define the variable once under `defaults > variables` in `aidp_workbench.yaml`.
2. Override the same variable name under each target environment.
3. Deploy with the target whose variable values match the destination workspace.
4. Use defaults for safe fallback values, not for production-only secrets or workspace-specific identifiers.

Example: Dev, QA, and Prod Overrides

The following example uses the same variable names across three targets. The job still references the variables once, while the dev region, QA region, and production environment supply different deploy-time values.

```
bundle:
  name: environmentVariableDemoBundle
resources:
  jobs:
    - namedVariableJobDemo
defaults:
  variables:
    job_param_value_default: "cloud data governance"
    deployment_region: "us-ashburn-1"
targets:
```

```

dev:
  aidp_ocid: "ocid1.aidataplatformdev.oc1.iad.exampledev"
  workspace_key: dev_workspace
  variables:
    job_param_value_default: "space tourism"
    deployment_region: "us-ashburn-1"
qa:
  aidp_ocid: "ocid1.aidataplatformdev.oc1.iad.exampleqa"
  workspace_key: qa_workspace
  variables:
    job_param_value_default: "data quality validation"
    deployment_region: "us-phoenix-1"
prod:
  aidp_ocid: "ocid1.aidataplatformdev.oc1.iad.exampleprod"
  workspace_key: prod_workspace
  variables:
    job_param_value_default: "customer operations"
    deployment_region: "eu-frankfurt-1"

```

Target	deployment_region	job_param_value_default	Runtime Result
dev	us-ashburn-1	space tourism	Job runs with development region and parameter values.
qa	us-phoenix-1	data quality validation	Job runs with QA region and parameter values.
prod	eu-frankfurt-1	customer operations	Job runs with production environment values.

① Note

If no matching target environment is found, or if a target does not define an override, Oracle AI Data Platform uses the configured default variable value when one exists. If neither an override nor a valid default is available, deployment or runtime execution can fail depending on where the missing variable is used.

Bundle Variables and Credential Store (Preview)

You can use variables with the credential store when a bundled job or notebook needs to look up a secret at runtime.

Your variables should carry the credential name and token key. The actual secret value must stay in the target credential store and should not be committed to Git.

Example Workflow: Using Variables with Credential Store

1. Create the credential and token key in credential store.
2. Define variables for the credential name and token key.
3. Pass those variables into the job or notebook as parameters.

4. At runtime, the notebook reads the parameters and calls the credential API to retrieve the secret. If the credential name, token key, or permissions are wrong, the job can fail with a credential lookup error.

Note

Do not document or commit secret values in bundle files. Bundle configuration should reference credential names, token keys, and target variables; the credential store remains the source of secret material.

Example: Credential Name and Token Key Variables

```
bundle:
  name: credentialVariableDemoBundle
resources:
  jobs:
    - credentialLookupJob
defaults:
  variables:
    cred_key: default_credential_name
    token_key: default_token_key
targets:
  dev:
    workspace_key: dev_workspace
    variables:
      cred_key: dev_credential_store_entry
      token_key: dev_token
  qa:
    workspace_key: qa_workspace
    variables:
      cred_key: qa_credential_store_entry
      token_key: qa_token
  prod:
    workspace_key: prod_workspace
    variables:
      cred_key: prod_credential_store_entry
      token_key: prod_token

jobs:
  credentialLookupJob:
    tasks:
      - taskKey: Task_854239
        parameters:
          - name: cred_key
            value: "${var.cred_key}"
          - name: token_key
            value: "${var.token_key}"
```

The notebook reads the parameter values and uses them to look up the secret from the credential store:

```
cred_name = odiUtils.parameters.getParameter("cred_key", "defaultCredValue")
token_name = odiUtils.parameters.getParameter("token_key",
"defaultTokenValue")
```

```
secret = aidpUtils.secrets.get(name=cred_name, key=token_name)
if secret:
    print("Credential lookup succeeded")
else:
    raise Exception("Credential lookup failed")
```

Environment	cred_key Value	token_key Value	Credential Store Requirement
dev	dev_credential_store_entry	dev_token	Create dev_credential_store_entry with key dev_token.
qa	qa_credential_store_entry	qa_token	Create qa_credential_store_entry with key qa_token.
prod	prod_credential_store_entry	prod_token	Create prod_credential_store_entry with key prod_token.

Purge a Bundle (Preview)

You can purge a deployed bundle to remove the bundle resources from your workspace.

1. Navigate to the bundle you want to purge resources for in your workspace.
2. Click the Deployment tab.
3. Click **Purge**.
4. Enter Purge in the prompt. Click **Purge**.

Note

Purging a bundle removes deployed resources first. If the bundle uses a publish path, Oracle AI Data Platform then removes the managed artifacts folder under the publish path. Other files under the publish path are preserved.

Purge Behavior for Published Artifacts (Preview)

When you purge a deployed bundle, Oracle AI Data Platform removes deployed resources first.

After the resources are removed, AI Data Platform removes only the managed artifact folder: <publish.path>/artifacts.

Other files or folders under the publish path are preserved. For example, if the publish path is /Workspace/Shared/customer_churn_bundle, purge removes /Workspace/Shared/customer_churn_bundle/artifacts but preserves other files or folders under /Workspace/Shared/customer_churn_bundle.

If artifact cleanup cannot be completed because of permissions or path access, resource deletion can still complete.

Existing Bundles and Publish Paths (Preview)

Existing bundles that do not have a publish block continue to use the previous artifact behavior until they are synced or manually updated.

If you manually add a publish block to an existing bundle, deployment uses the configured publish path. However, existing descriptor files are not automatically changed during deployment. If existing descriptors still reference the previous artifact location, those resources might continue to use the previous location.

If you want existing resources to use the publish path, update the descriptor references or sync the bundle before manually adding a custom publish block. For example:

Old:

```
${bundle.root}}/.aidp/artifacts/Shared/task1.py
```

New:

```
${publish.path}}/artifacts/Shared/task1.py
```

When sync adds new or refreshed descriptor content for a bundle that already has a publish block, the new descriptor references use the publish path format.

Troubleshooting Bundles (Preview)

If you encounter problems when creating and managing Git bundles, check the following list of issues for potential solutions.

Symptom	Likely Cause	Recommended Action
Create Bundle is unavailable or the bundle cannot be created.	The user is not in a Git folder, or lacks required workspace permissions.	Navigate to a Git folder and confirm the user can create workspace resources.
Sync completes, but another environment does not see the update.	The refreshed bundle files were not committed, pushed, or pulled into the target Git folder.	Commit and push from the source Git folder, then pull in the target workspace before deploying.
Deployment tab shows no deployed items.	The bundle has not been deployed yet, or deployment did not complete successfully.	Click Deploy and wait for the completion notification. Review logs if deployment fails.
Deployed job fails with a credential lookup error.	The target credential name, token key, or access permission does not match the bundle configuration.	Create the credential in the target credential store or update target variables before redeploying.
Agent deployment fails or starts with inactive compute.	The target environment may not have compatible AI compute, region, model, or dependency settings.	Review target overrides, compute status, model availability, and deployment logs.
Changes to a source notebook or agent are not reflected after deployment.	The bundle was deployed without first syncing from the updated source resource.	Run Sync on the bundle, commit and push the refreshed bundle files, pull them in the target, then deploy again.

Oracle AI Data Platform (AIDP) Developer Extension

The Oracle AI Data Platform (AIDP) Developer extension helps you work with agent flows directly from your Integrated Development Environment (IDE).



Oracle AI Data Platform currently supports Microsoft Visual Studio Code and Cursor as IDEs. Installing through the Visual Studio Code Marketplace ensures your developer extension stays up to date. If you manually install the develop extension on an IDE, you need to download and reinstall the VSIX when you want to apply updates.

You can download the developer extension from the [Microsoft Visual Studio Code Marketplace](#).

With the developer extension, you can browse remote workspace files, sync changes, deploy flows, test endpoints, and troubleshoot with built-in logs and traces. The developer extension helps developers build and iterate on agent flows, data engineers manage their workspace assets, and teams validate endpoint behavior before production rollout.

The developer extension provides the following key features:

- Workspace and catalog browsing in your IDE sidebar
- Remote file preview, download, compare, and upload
- Deploy, redeploy, undeploy, and status checks
- Built-in endpoint test chat
- Trace and debug visibility for faster troubleshooting

Best Practices

When using the developer extension, follow these guidelines for best results:

- Start in TEST, then promote to PROD.
- Use compare/diff before uploading in shared projects.
- Keep polling intervals reasonable to reduce noise.
- Treat debug settings/tokens as sensitive.

Prerequisites

Before setting up your developer extension, make sure you have:

- Access to your AI Data Platform instance
- A valid OCI profile/configuration on your machine
- Required workspace and agent permissions
- Region and workspace details

Install Oracle AI Data Platform (AIDP) Developer Extension on Visual Studio Code

You can download and install the Oracle AI Data Platform (AIDP) Developer Extension from the Marketplace in Visual Studio Code.

1. Open Visual Studio Code.
2. Click **Extensions**.
3. Search for **Oracle AI Data Platform (AIDP) Developer Extension**.
4. Click **Install**.
5. Restart Visual Studio Code if prompted.
6. Confirm the AI Data Platform (AIDP) Developer Extension appears on the sidebar.

Install Oracle AI Data Platform (AIDP) Developer Extension Manually

You can download the VSIX file for Oracle AI Data Platform (AIDP) Developer Extension and manually install it to your IDE.

1. In your IDE, open the command palette. In Cursor and Visual Studio Code, the hotkey for this is Ctrl+Shift+P (Cmd+Shift+P in MacOS).
2. Enter `Install from VSIX` and select `Extensions: Install from VSIX...`
3. Browse to the `.vsix` file and click **Install**.
4. Restart your IDE if prompted.
5. Confirm the AI Data Platform (AIDP) Developer Extension appears on the sidebar.

Set Up Oracle AI Data Platform (AIDP) Developer Extension in Visual Studio Code

After you've installed Oracle AI Data Platform (AIDP) Developer Extension, you need set it up to connect to your Oracle AI Data Platform instance.

1. In Visual Studio Code, open a folder, then click the **Oracle AI Data Platform (AIDP) Developer Extension**.
2. Start the setup wizard.
3. Copy and paste your AI Data Platform console URL into the provided field and press Enter.
4. Confirm the detected region or set manually. Press **Enter**.
5. Select workspace, default agent flow, and optional compute.
6. Save your configuration and verify the connection status.

Part VIII

Security

Oracle AI Data Platform offers security options like multi-layered security, role-based access controls, private networking, and more.

Topics:

- [Permissions Model](#)
- [IAM Policies for Oracle AI Data Platform](#)
- [Roles](#)
- [Network Sources](#)

Permissions Model

This chapter describes the permissions model Oracle AI Data Platform uses to manage access.

Topics:

- [About Permissions](#)
- [Workspace Permissions](#)
- [Workspace Folder Permissions](#)
- [Compute Cluster Permissions](#)
- [Job Permissions](#)
- [Notebook Permissions](#)
- [Master Catalog Permissions](#)
- [Standard Catalog Permissions](#)
- [External Catalog Permissions](#)
- [Schema Permissions](#)
- [Table Permissions](#)
- [Volume Permissions](#)
- [Delta Sharing Permissions](#)
- [Agent Permissions](#)
- [AI Compute Cluster Permissions](#)
- [Knowledge Base Permissions](#)
- [Git Folder Permissions \(Preview\)](#)
- [Bundle Permissions \(Preview\)](#)
- [Credential Permissions \(Preview\)](#)

About Permissions

Oracle AI Data Platform permissions follow a similar model for all objects that use them.

You can manage permissions for each object from its **Permissions** tab.

AI Data Platform has two layers of security - access to OCI resources using IAM policies and access to Data Platform objects. Users must have access to OCI resources first before granting them access to AI Data Platform objects. Users of AI Data Platform require access to navigate to resources in OCI console and IAM permissions to list compartments and buckets. To access an AI Data Platform instance, you require at least USE IAM policy permissions. These IAM policies are needed even if you have AI_DATA_PLATFORM_ADMIN role on an AI Data Platform instance.

Permissions in AI Data Platform follow a hierarchy where permissions granted for a parent object or space grant permissions to contained objects and spaces.

Permission to Create Workspaces

Permissions to create workspaces are included in the `AI_DATA_PLATFORM_ADMIN` role by default. If you want users other than the administrator to be able to create workspaces, you need to provide `CREATE_WORKSPACE` permissions to that user. You can assign `CREATE_WORKSPACE` to a user from the Workspace Listing screen.

Permission Inheritance

The permission model in Oracle AI Data Platform is designed to support permission inheritance so that permissions granted at the parent level automatically flow to the child objects.

For example, User A is granted `SELECT` permissions for Catalog A, and Catalog A has the following hierarchy:

```
|CatalogA
|-----schema1
|-----table1
```

In this example, User A has `SELECT` permissions for Catalog A, schema1, and table1.

Permission Expansion

Permissions granted at the child level do not require explicit permissions at the parent object level. For example, you manage a catalog with the following hierarchy:

```
|CatalogA
|-----schema1
|-----table1_1
|-----schema2
|-----table2_1
|-----table2_2
```

In this example, if you grant User A `SELECT` permissions for table2_2, User A only sees the following hierarchy:

```
|CatalogA
|-----schema2
|-----table2_2
```

User A is only granted limited list permissions for Catalog A and schema2. They can see objects that contain table2_2, but do not have access to do anything more to them.

Limitations

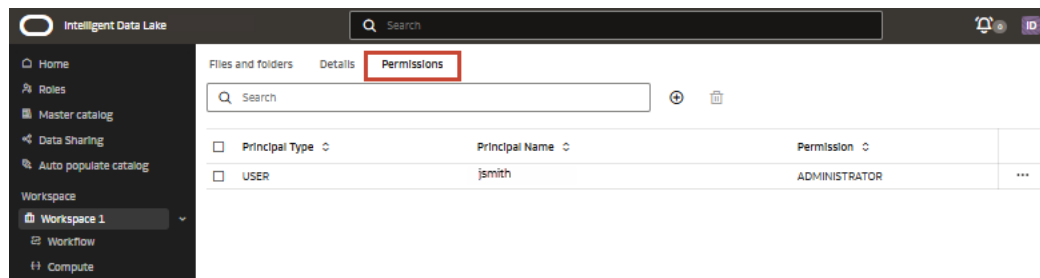
Known limitations of permissions inheritance and expansion are:

- Permissions inheritance does not show up for `ADMIN` permissions for any object type
- Permissions inheritance does not work for volumes
- Permissions expansion does not work for volumes and workspace files and folders

Workspace Permissions

You can set role-based action controls for a workspace you own that apply to all its contents.

Workspace permissions are managed from the Permissions tab, located at the top of your workspace home page.



A user can be granted the following permissions:

- **USER:** You can create folders/files in root, and have **MANAGE** permissions on the Shared Folder.
- **PRIVILEGED_USER:** You have **USER** permissions and can also create compute.
- **ADMINISTRATOR:** You have **ADMIN** permissions on all workspace objects and can update or delete a workspace and manage permissions.

Note

The **USER** permission for workspaces also grants users the **MANAGE** permission on all objects in the Shared Folder. A Shared Folder cannot be deleted, renamed, or moved.

Permissions can be granted to users, groups, or roles. You can either select users from a list of Oracle AI Data Platform users or add a user or role by the OCID.

Create Workspace Permissions

You can grant access to your workspaces to users, roles, or groups.

You must have administrator privileges in the workspace to grant access to others.

1. On the Home page, click **Workspace**.
2. Next to your workspace, click ... **Actions** then click **Permissions**.
3. Click **+ New Permission**.
4. Select the permissions level and principal type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click **Search** and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Workspace Permissions

You can change permission settings for any workspace where you have administrator privileges.

1. On the Home page, click **Workspace**.
2. Next to your workspace, click ... **Actions** then click **Permissions**.
3. Next to the permission, click **Actions** and click **Edit**.
4. Select the new permission level from the **Permissions** dropdown and click **Save**.

Delete Workspace Permissions

You can delete a workspace permission to remove access and actions for all contained users.

1. On the Home page, click **Workspace**.
2. Next to your workspace, click ... **Actions** then click **Permissions**.
3. Next to your permission, click **Actions** and click **Delete**.
4. On the confirmation window, click **Delete**.

Workspace Folder Permissions

You can manage which users, roles, and groups can view and modify folders in your workspaces.

A user can be granted the following permissions for folders in your workspace:

- **READ:** Users can read/list files or folders in the workspace.
- **USE:** Users have READ permissions and can clone folders.
- **MANAGE:** Users have READ and USE permissions and can download contents of a folder.
- **ADMIN:** Users all permissions and can create, rename, move, or delete files and folders, as well as manage other user permissions.

An admin can grant permission to any principal who has at least a workspace USER permissions.

Operation	READ	USE	MANAGE	ADMIN
List	Yes	Yes	Yes	Yes
View object	Yes	Yes	Yes	Yes
Clone folder*	No	Yes	Yes	Yes
Download folder (coming soon)	No	No	Yes	Yes
Create folder in folder	No	No	No	Yes
Create file in folder	No	No	No	Yes
Rename folder	No	No	No	Yes
Move folder	No	No	No	Yes
Delete folder	No	No	No	Yes

Operation	READ	USE	MANAGE	ADMIN
Manage user permissions	No	No	No	Yes
* Users need ADMIN permissions on the Clone destination to successfully clone a folder.				

Permissions granted on a folder in the workspace will cascade to the child objects by default, but you have the option to turn off cascading. If a user has ADMIN permissions on a folder, that user will have ADMIN permissions on all child files and folders. When an ADMIN grants permission on a folder with cascade, all child resources current and future have same permission. When an ADMIN grants permission on a folder without cascade, permission is applicable only on the current object.

For example, given the following folder structure, an ADMIN wants to grant USE access to USER2 for just File1.csv:

- WORKSPACE1
 - F1
 - * F2
 - * File1.csv
 - * F3
 - * File2.csv

Assuming USER2 is already a user in the workspace, the ADMIN grants USE permission to USER2 on File1.csv. For the user to be able to work with File1.csv the ADMIN must also explicitly grant READ permission on folders F1 and F2, without cascading. As a result, the user sees this folder structure:

- WORKSPACE1
 - F1
 - * F2
 - * File1.csv

While granting permission to a user on a folder in the workspace, in the Add Permission to folder drawer, at the last section named Permission inheritance settings, you have the option of selecting one of the following options:

- Inherit to all subfolders and items
- Restrict to this folder only

If you select **Inherit to all subfolders and items**, the selected permission cascades to all the child folders and files.

If the you select **Restrict to this folder only**, the permission is granted at the same folder level only.

Once the permission is granted, you can go to the Permission listed page for the user you granted permissions to. If you selected Inherit to all sub-folders and items, the column **Will be inherited** displays **Yes** for the newly added permission. Otherwise, the column displays **No**.

Workspace File Permissions

You can manage which users, roles, and groups can view and modify files in your workspaces.

A user can be granted the following permissions for a file in workspace:

- **READ:** Users can read/list files in the workspace.
- **USE:** Users have READ permissions and can clone files, clone notebooks, run files, and cancel file runs.
- **MANAGE:** Users have READ and USE permissions and can download and edit files.
- **ADMIN:** Users have all permissions and can move, rename, or delete files, as well as manage other user permissions.

Operation	READ	USE	MANAGE	ADMIN
List	Yes	Yes	Yes	Yes
View details	Yes	Yes	Yes	Yes
Clone file/ notebook*	No	Yes	Yes	Yes
Cancel file run	No	Yes	Yes	Yes
Run file	No	Yes	Yes	Yes
Download file	No	No	Yes	Yes
Edit file	No	No	Yes	Yes
Rename file	No	No	No	Yes
Move file	No	No	No	Yes
Delete file	No	No	No	Yes
Manage user permissions	No	No	No	Yes

* Users need ADMIN permissions on the Clone destination to successfully clone a file or notebook.

Permissions granted on a folder in the workspace will cascade to the child objects by default, but you have the option to turn off cascading. If a user has ADMIN permissions on a folder, that user will have ADMIN permissions on all child files and folders. When an ADMIN grants permission on a folder with cascade, all child resources current and future have same permission. When an ADMIN grants permission on a folder without cascade, permission is applicable only on the current object.

For example, given the following folder structure, an ADMIN wants to grant USE access to USER2 for just File1.csv:

- WORKSPACE1
 - F1
 - * F2
 - * File1.csv
 - * F3
 - * File2.csv

Assuming USER2 is already a user in the workspace, the ADMIN grants USE permission to USER2 on File1.csv. For the user to be able to work with File1.csv the ADMIN must also explicitly grant READ permission on folders F1 and F2, without cascading. As a result, the user sees this folder structure:

- WORKSPACE1
 - F1
 - * F2

* File1.csv

While granting permission to a user on a folder in the workspace, in the Add Permission to folder drawer, at the last section named Permission inheritance settings, you have the option of selecting one of the following options:

- Inherit to all subfolders and items
- Restrict to this folder only

If you select **Inherit to all subfolders and items**, the selected permission cascades to all the child folders and files.

If the you select **Restrict to this folder only**, the permission is granted at the same folder level only.

Once the permission is granted, you can go to the Permission listed page for the user you granted permissions to. If you selected Inherit to all sub-folders and items, the column **Will be inherited** displays **Yes** for the newly added permission. Otherwise, the column displays **No**.

Create File and Folder Permissions

You can set individual permissions for files and folders in your workspaces.

1. Navigate to the file or folder you want to set permissions for.
2. Click ... **Actions** and click **Permissions**.
3. Click + **Create Permission**.
4. Select a permission level, principal type, and the user from the dropdown menus.
5. Click **Save**.

Modify File and Folder Permissions

You can modify existing permissions for files or folders in your workspace.

1. Navigate to the file or folder you want to set permissions for.
2. Click ... **Actions** and click **Permissions**.
3. Next to permission you want to modify, click ... **Actions** and click **Edit**.
4. Change the permissions details as needed and click **Save**.

Delete File and Folder Permissions

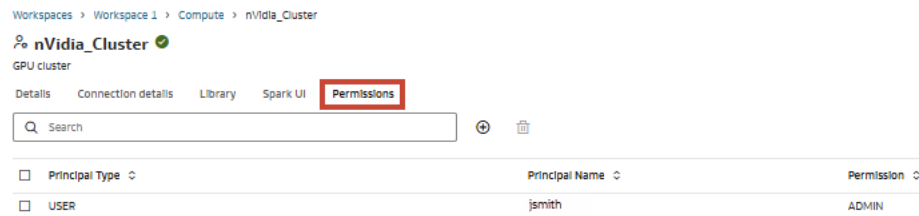
You can delete permissions for files and folders in your workspace.

1. Navigate to the file or folder you want to set permissions for.
2. Click ... **Actions** and click **Permissions**.
3. Next to the permission you want to delete, click ... **Actions** and click **Delete**.
4. Click **Delete**.

Compute Cluster Permissions

You can control which users and roles have view, read, and administrator access to your compute clusters.

You create and manage user permissions from the **Permissions** tab in your cluster.



As an administrator, you can grant permissions to any principal who has at least User workspace permissions.

Operation	Read	Use	Admin
List cluster	Yes	Yes	Yes
Attach cluster to notebook/job	Yes	Yes	Yes
View driver logs, Spark UI	Yes	Yes	Yes
View cluster metrics	Yes	Yes	Yes
Start/Restart cluster	No	Yes	Yes
Terminate cluster	No	Yes	Yes
Edit cluster	No	No	Yes
Attach/Upload library to cluster	No	No	Yes
Grant/Revoke permissions	No	No	Yes

Create Cluster Permissions

You can control which users and roles can see and modify your clusters.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Permissions** tab.
3. Click **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Cluster Permissions

You can modify permissions for users and roles assigned to your cluster.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Permissions** tab.
3. Next to the user or role you want to modify, click **...** **Actions** then click **Edit**.

4. Select a new permission level from the dropdown. Click **Save**.

Delete Cluster Permissions

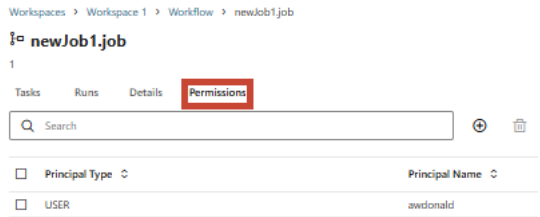
You can remove permissions that are no longer needed for users or roles on your cluster.

1. Navigate to your workspace and click **Compute**.
2. Click your cluster, then click the **Permissions** tab.
3. Next to the user or role you want to delete, click **...** **Actions** then click **Delete**.
4. Click **Delete**.

Job Permissions

Job permissions control which users and roles have access to your jobs.

You manage the users and roles that have access to your job from the Permissions tab in your job.



The following permission levels are available to job users:

- Read
- Use
- Manage
- Admin

Each permission level has access to a different set of operations, outlined below.

Operation	Read	Use	Manage	Admin
List	Y	Y	Y	Y
View details	Y	Y	Y	Y
Execution status	Y	Y	Y	Y
Attach/Detach compute	N	Y	Y	Y
Run	N	Y	Y	Y
View task log	N	Y	Y	Y
Rename job	N	N	Y	Y
Edit job	N	N	Y	Y
Terminate workflow	N	N	Y	Y
Move file	N	N	N	Y
Delete job	N	N	N	Y

Operation	Read	Use	Manage	Admin
Grant/Revoke permissions	N	N	N	Y

Create Job Permissions

You can create permissions to control which users and roles have access to your jobs.

You can only grant access to jobs that you own.

1. Navigate to the job you want to grant access to.
2. Click **Permissions**.
3. Click  **New Permissions**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Job Permissions

You can grant or revoke permissions by changing the permission levels for existing users or roles.

1. Navigate to your workspace and click **Workflow**.
2. Click your job, then click the **Permissions** tab.
3. Next to the user or role you want to modify, click ... **Actions** then click **Edit**.
4. Select a new permission level from the dropdown. Click **Save**.

Delete Job Permissions

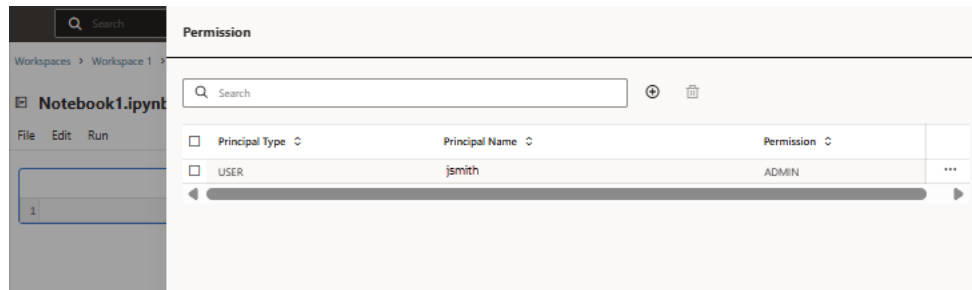
You can remove permissions that are no longer needed for users or roles in your job.

1. Navigate to your workspace and click **Workflow**.
2. Click your job, then click the **Permissions** tab.
3. Next to the user or role you want to delete, click ... **Actions** then click **Delete**.
4. Click **Delete**.

Notebook Permissions

Notebooks permissions determine which users, roles, and groups can view and modify your notebook.

Permissions for a notebook are viewed by clicking **Actions** on the top right of your notebook, and clicking **Permissions**. From the **Permission** page, you can view, create, modify, or delete permissions for your notebook. You can filter the displayed list of users by entering a user in the **Search** bar.



The following permission levels are available to notebook users:

- Read
- Use
- Manage
- Admin

Each permission level has access to a different set of operations, outlined below.

Operation	Read	Use	Manage	Admin
List	Y	Y	Y	Y
View details	Y	Y	Y	Y
Execution status	Y	Y	Y	Y
Attach/Detach compute	N	Y	Y	Y
Run workflow	N	Y	Y	Y
View log	N	Y	Y	Y
Rename notebook	N	N	Y	Y
Edit notebook	N	N	Y	Y
Terminate workflow	N	N	Y	Y
Move file	N	N	N	Y
Delete notebook	N	N	N	Y
Grant/Revoke permissions	N	N	N	Y

Create Notebook Permissions

You can set individual permissions for notebooks you own.

1. Navigate to the notebook you want to set permissions for.
2. Click **Actions** and click **Permissions**.
3. Click **Create Permission**.
4. Select a permission level, principal type, and the user from the dropdown menus.
5. Click **Save**.

Modify Notebook Permissions

You can modify existing permissions for notebooks you own.

1. Navigate to the notebook you want to set permissions for.
2. Click **Actions** and click **Permissions**.
3. Next to permission you want to modify, click **Actions** and click **Edit**.
4. Change the permissions details as needed and click **Save**.

Delete Notebook Permissions

You can delete permissions for notebooks you administer.

1. Navigate to the notebook you want to set permissions for.
2. Click **Actions** and click **Permissions**.
3. Next to the permission you want to delete, click **Actions** and click **Delete**.
4. Click **Delete**.

Master Catalog Permissions

Permissions at the master catalog level determine who can create new standard and external catalogs and grant permissions to others.

You manage permissions for the Master catalog from the **Permissions** tab.



The user that creates your Oracle AI Data Platform is automatically granted ADMIN permissions for the Master catalog. There are two permission levels for Master catalog:

- **CREATE_CATALOG:** User can create standard and external catalogs.
- **CREATE_CREDENTIAL:** User can create credentials in Credential Storage.
- **CREATE_SHARE:** User can create data shares.
- **CREATE_RECIPIENT:** User can create data share recipients.
- **ADMIN:** User can view all catalogs, create, edit, or delete catalogs, credentials, shares, recipients, and their child objects, and grant or revoke permissions.

Master Catalog Permission Inheritance

ADMIN permissions for the Master Catalog confer ADMIN permissions on all child objects in the Master Catalog. When a user with CREATE_CATALOG permissions creates a catalog, they are automatically given ADMIN permission for the newly created catalog and all its child objects.

Create Master Catalog Permissions

You can set permissions to manage who can create, edit, and delete catalogs and grant permissions to others.

1. On the Home page, click **Master Catalog**.
2. Click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Master Catalog Permissions

You can modify the permissions of users or roles for the Master catalog.

1. On the Home page, click **Master Catalog**.
2. Click the **Permissions** tab.
3. Next to the permission, click  **Actions** and click **Edit**.
4. Select the new permission level from the **Permissions** dropdown and click **Save**.

Delete Master Catalog Permissions

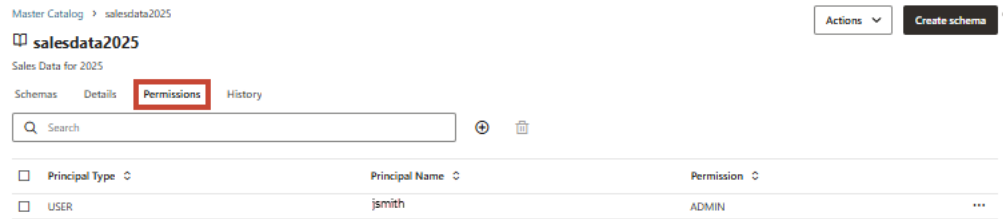
You can delete catalog permissions to remove access and actions for all contained users or roles.

1. On the home page, click **Master Catalog**.
2. Click the **Permissions** tab.
3. Next to your permission, click  **Actions** and click **Delete**.
4. Click **Delete**.

Standard Catalog Permissions

You can manage permissions for standard catalogs to determine which users, roles, and groups can view and modify your catalogs.

You can set permissions for standard catalogs from the **Permissions** tab of your catalog. You can filter the list of users and roles that have access to your catalog by entering a name in the **Search** bar.



Permissions set at the catalog level cascade down to any children of the catalog. Permissions set at the schema level apply to any child objects of the schema.

- **SELECT:** Users can read/list catalogs, schema, and volumes. Users can run select queries on views and tables.
- **MANAGE:** Users have all Select permissions at the Standard catalog level and can alter schema, tables, and views and write to volumes. Users can also insert, update, and delete data in tables.
- **CREATE_SCHEMA:** Users have all Manage permissions at the Standard catalog level and can create new schema in the catalog.
- **ADMIN:** Users have all Create_Schema permissions at the Standard catalog level and can delete schema, as well as manage other user permissions

Operation	SELECT	MANAGE	CREATE_SCHEM A	ADMIN
Read/List	Yes	Yes	Yes	Yes
Run queries	Yes	Yes	Yes	Yes
Edit schema/tables/ volumes/views	No	Yes	Yes	Yes
Create schema	No	No	Yes	Yes
Delete schema	No	No	No	Yes
Manage permissions	No	No	No	Yes

Master Catalog Permission Inheritance

Users with CREATE_CATALOG or ADMIN permissions at the Master catalog level are treated as having the following permissions in standard catalogs:

- SELECT
- MANAGE
- CREATE_SCHEMA
- ADMIN

External Catalog Permissions

You can manage permissions for external catalogs to determine which users, roles, and groups can view and modify your catalogs.

Users with ADMIN permissions for an external catalog can grant permissions to:

- Any IAM user principal or IAM group. Users are loaded in the following order:

1. All users from the selected domain who have opened an Oracle AI Data Platform instance at least once
 2. All remaining users in the selected domain, in alphabetical order
- Roles the ADMIN user can view.

External catalog permissions grant the following actions:

Operation	MANAGE	ADMIN
Read/List & Perform DML operations *	Yes	Yes
DDL (Coming soon)		
Edit catalog name	No	Yes
Edit catalog properties (password, etc.)	No	Yes
Drop catalog	No	Yes
Manage permissions	No	Yes

* External catalog permissions are limited to the permissions of the user used to connect to the external source. If the user of the external source used to create the external catalog has read-only permission, MANAGE permission of the external catalog is also limited to read-only permission.

Master Catalog Permission Inheritance

Users with CREATE_CATALOG or ADMIN permissions at the Master catalog level are treated as having the following permissions in external catalogs:

- MANAGE
- ADMIN

Create Catalog Permissions

You can grant permissions to view and modify catalogs, schema, tables, and volumes.

1. On the Home page, click **Master Catalog**.
2. Navigate to the catalog you want to create a new permission for and click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Catalog Permissions

You can modify the permissions of users or roles for catalogs you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your catalog, then click the **Permissions** tab.
3. Next to the permission, click  **Actions** and click **Edit**.

4. Select the new permission level from the **Permissions** dropdown and click **Save**.

Delete Catalog Permissions

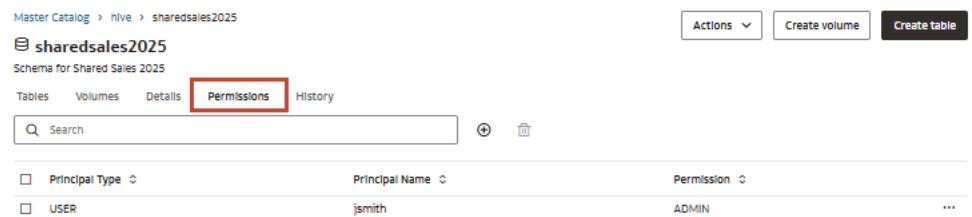
You can delete catalog permissions to remove access and actions for all contained users or roles.

1. On the home page, click **Master Catalog**.
2. Navigate to your catalog, then click the **Permissions** tab.
3. Next to your permission, click ... **Actions** and click **Delete**.
4. Click **Delete**.

Schema Permissions

Schema permissions determine which users, roles, and groups can view and modify your schema and their child objects.

You control the users and roles that can access your schema from the schema **Permissions** tab.



Permissions set at the schema level apply to any child objects of the schema.

Schema permissions grant the following actions:

- **SELECT:** Users can read/list tables, view, and volumes in the schema. Users can run select queries on views and tables.
- **WRITE:** Users have Select permissions and can alter tables or data in tables, write to volumes, and alter views.
- **CREATE_MODEL::** Users can create models in a schema.
- **CREATE_TABLE::** Users can create tables in a schema.
- **CREATE_VIEW:** Users can create views in a schema.
- **CREATE_VOLUME:** Users can create volumes in a schema.
- **ADMIN:** Users have Select, Write, and all Create permissions and can create, modify, or delete other user permissions.

Operation	SELECT	WRITE	CREATE_MODEL	CREATE_TABLE	CREATE_VIEW	CREATE_VOLUME	ADMIN
Read/List	Yes	Yes	Yes	Yes	Yes	Yes	Yes

Operation	SELECT	WRITE	CREATE_ MODEL	CREATE_ TABLE	CREATE_ VIEW	CREATE_ VOLUME	ADMIN
Run queries/ Read volumes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Edit tables/ volumes/ views	No	Yes	Yes	Yes	Yes	Yes	Yes
Create model	No	No	Yes	No	No	No	Yes
Create table	No	No	No	Yes	No	No	Yes
Create view	No	No	No	No	Yes	No	Yes
Create volume	No	No	No	No	No	Yes	Yes
Delete schema	No	No	No	No	No	No	Yes
Manage permissions	No	No	No	No	No	No	Yes

Schema Permission Inheritance

Schema Permission	Catalog Level Permission		
SELECT	SELECT	MANAGE	ADMIN
WRITE	X		
CREATE_VIEW	X	X	
CREATE_VOLUME	X	X	
CREATE_TABLE	X	X	
ADMIN	X	X	

Create Schema Permissions

You can control which users and roles have access to schema you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your schema, then click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Schema Permissions

You can modify the permissions of users or roles for schema you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your schema, then click the **Permissions** tab.
3. Next to the permission, click ... **Actions** and click **Edit**.
4. Select the new permission level from the **Permissions** dropdown and click **Save**.

Delete Schema Permissions

You can delete schema permissions to remove access and actions for all contained users or roles.

1. On the home page, click **Master Catalog**.
2. Navigate to your catalog, then click the **Permissions** tab.
3. Next to your permission, click ... **Actions** and click **Delete**.
4. Click **Delete**.

Table Permissions

Table permissions determine which users, roles, and groups can view and modify your tables.

Table permissions grant the following actions:

- **SELECT:** Users can read/list tables. Users can run select queries on tables.
- **INSERT:** Users can read/list tables and write to tables.
- **UPDATE:** Users can read/list tables and can run updates on table data.
- **DELETE:** Users can read/list tables and can delete data from the table.
- **ALTER:** Users can read/list tables and can modify table names or descriptions.
- **ADMIN:** Users have all permissions and can create, modify, or delete other user permissions.

Operation	SELECT	INSERT	UPDATE	DELETE	ALTER	ADMIN
List table	Yes	Yes	Yes	Yes	Yes	Yes
Read table data	Yes	No	No	No	No	Yes
Write data to table	No	Yes	No	No	No	Yes
Update data in table	No	No	Yes	No	No	Yes
Delete data from table	No	No	No	Yes	No	Yes
Alter table metadata	No	No	No	No	Yes	Yes
Delete table	No	No	No	No	No	Yes
Manage user permissions	No	No	No	No	No	Yes

Table Permission Inheritance

Table Permission	Schema Level Permission		
SELECT	SELECT	MANAGE	ADMIN
INSERT	X		
UPDATE	X		
DELETE	X		
ALTER	X		
ADMIN	X	X	

Create Table Permissions

You can control which users and roles have access to tables you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your table, then click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Table Permissions

You can modify the permissions of users or roles for tables you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your table, then click the **Permissions** tab.
3. Next to the permission, click  **Actions** and click **Edit**.
4. Select the new permission level from the **Permissions** dropdown and click **Save**.

Delete Table Permissions

You can delete table permissions to remove access and actions for all contained users or roles.

1. On the home page, click **Master Catalog**.
2. Navigate to your catalog, then click the **Permissions** tab.
3. Next to your permission, click  **Actions** and click **Delete**.
4. Click **Delete**.

Volume Permissions

Volume permissions determine which users, roles, and groups can view and modify your volumes.

Volume permissions grant the following actions:

- **READ:** Users can list folders/files and read files from volume.
- **WRITE:** Users can list folders/files, read files, create folder and files and write to files in a volume.
- **ADMIN:** User will have READ/WRITE permissions on the volume, delete/create a volume, and will be able to grant/revoke permissions on the volume.

Operation	READ	WRITE	ADMIN
List volume	Yes	Yes	Yes
Read volume data	Yes	Yes	Yes
Write data to volume	No	Yes	Yes
Delete data from volume	No	Yes	Yes
Create folder	No	Yes	Yes
Delete volume	No	No	Yes
Create volume	No	No	Yes
Manage user permissions	No	No	Yes

Create Volume Permissions

You can control which users and roles have access to volumes you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your volume, then click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Volume Permissions

You can modify the permissions of users or roles for volumes you own.

1. On the home page, click **Master Catalog**.
2. Navigate to your volumes, then click the **Permissions** tab.
3. Next to the permission, click  **Actions** and click **Edit**.
4. Select the new permission level from the **Permissions** dropdown and click **Save**.

Delete Volume Permissions

You can delete volume permissions to remove access and actions for all contained users or roles.

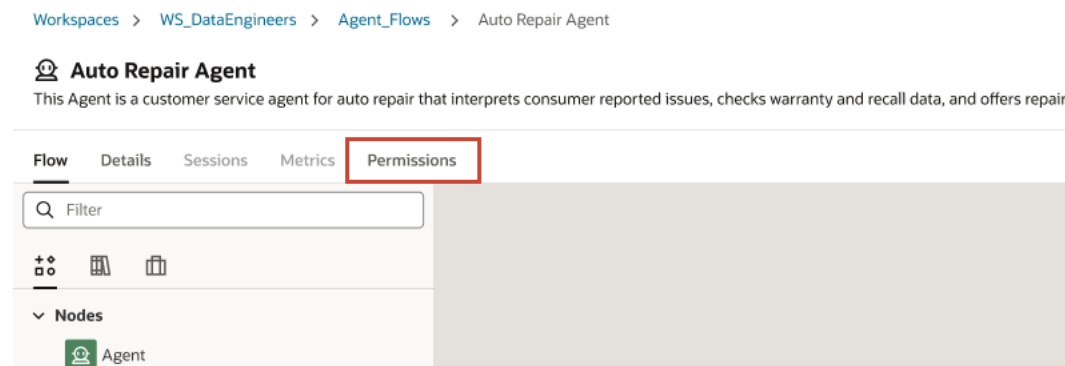
1. On the home page, click **Master Catalog**.

2. Navigate to your volume, then click the **Permissions** tab.
3. Next to your permission, click ... **Actions** and click **Delete**.
4. Click **Delete**.

Agent Permissions

You can set role-based action controls for agents to control who can run, configure, and delete agents.

Agent permissions are managed from the Permissions tab, located at the top of your agent page.



You can grant users the following permissions for agents:

- **READ:** You can view agents and agent details, and invoke deployed agents.
- **USE:** You have **READ** permissions and you can invoke and test agents.
- **MANAGE:** You have **USE** permissions and you can attach or detach agents and AI compute clusters, deploy or remove agents in AI compute instances, activate or deactivate agent URIs, update agent definitions, and test agents attached to AI compute.
- **ADMIN:** You have **MANAGE** permissions and you can delete agents and manage user permissions.

Operation	READ	USE	MANAGE	ADMIN
List	Yes	Yes	Yes	Yes
View object	Yes	Yes	Yes	Yes
Invoke and test agent	No	Yes	Yes	Yes
Attach and detach clusters	No	No	Yes	Yes
Deploy and undeploy	No	No	Yes	Yes
Edit agent	No	No	Yes	Yes
Delete agent	No	No	No	Yes
Manage user permissions	No	No	No	Yes

Create Agent Permissions

You can grant permissions to read, use, manage, or administrate agents to users, roles, or groups.

1. Navigate to the agent you want to grant access to.
2. Click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Agent Permissions

You can modify existing permissions for agents that you own.

1. Navigate to the agent you want to set permissions for.
2. Click the Permissions tab.
3. Next to permission you want to modify, click  **Actions** and click **Edit**.
4. Change the permissions details as needed and click **Save**.

Delete Agent Permissions

You can delete permissions for agents in which you have admin permissions.

1. Navigate to the agent you want to delete permissions for.
2. Click the Permissions tab.
3. Next to the permission you want to delete, click **Actions** and click **Delete**.
4. Click **Delete**.

AI Compute Cluster Permissions

You can control which users and roles have read, use, and administrator access to your AI compute clusters.

You create and manage user permissions from the Permissions tab in your AI compute cluster.



As an administrator, you can grant permissions to any principal who has at least User workspace permissions.

Operation	Read	Use	Admin
List AI compute clusters	Yes	Yes	Yes
View AI compute cluster details	Yes	Yes	Yes
Attach agent	No	Yes	Yes
Deploy agent	No	Yes	Yes
Update AI compute cluster configuration	No	No	Yes
Delete AI cluster	No	No	Yes
Manage permissions	No	No	Yes

Create AI Cluster Permissions

You can control which users and roles can see and modify your AI clusters.

1. Navigate to your workspace and click **Compute**.
2. Click the AI Compute tab.
3. Next to AI compute cluster, click ... **Actions** then click **Permissions**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify AI Cluster Permissions

You can modify permissions for users and roles assigned to your AI compute cluster.

1. Navigate to your workspace and click **Compute**.
2. Click the AI Compute tab.
3. Next to your AI compute cluster, click ... **Actions** then click **Permissions**.
4. Next to the user or role you want to modify, click ... **Actions** then click **Edit**.
5. Select a new permission level from the dropdown. Click **Save**.

Delete AI Cluster Permissions

You can remove permissions that are no longer needed for users or roles on your AI compute cluster.

1. Navigate to your workspace and click **Compute**.
2. Click the AI Compute tab.
3. Next to your AI compute cluster, click ... **Actions** then click **Permissions**.
4. Next to the user or role you want to delete, click ... **Actions** then click **Delete**.
5. Click **Delete**.

Knowledge Base Permissions

You can control which users and roles have read, use, and administrator access to your knowledge bases.

You create and manage user permissions from the Permissions tab in your knowledge base.



Operation	Select	Manage	Admin
List knowledge bases	Yes	Yes	Yes
View knowledge bases (including files and folders)	Yes	Yes	Yes
Search and retrieve	Yes	Yes	Yes
Edit	No	Yes	Yes
Trigger ingestion	No	Yes	Yes
Restart ingestion	No	Yes	Yes
Cancel ingestion	No	Yes	Yes
Delete knowledge base	No	No	Yes
Manage permissions	No	No	Yes

Create Knowledge Base Permissions

You can control which users and roles can see and modify your knowledge bases.

1. Navigate to your knowledge base and click the Permissions tab.
2. Click **New permission**.

3. Select the permissions level and user type from the dropdowns.
4. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
5. Click **Create**.

Modify Knowledge Base Permissions

You can modify existing permissions for knowledge bases that you own.

1. Navigate to the knowledge base you want to modify permissions for.
2. Click the Permissions tab.
3. Next to permission you want to modify, click **...** **Actions** and click **Edit**.
4. Change the permissions details as needed and click **Save**.

Delete Knowledge Base Permissions

You can delete permissions for knowledge bases in which you have admin permissions.

1. Navigate to the knowledge base you want to delete permissions for.
2. Click the Permissions tab.
3. Next to the permission you want to delete, click **Actions** and click **Delete**.
4. Click **Delete**.

Git Folder Permissions (Preview)

You can manage which users, roles, and groups can view and modify Git files and folders in your workspaces.

You can configure two types of Git access controls: Git Folder Operations and Git Bundle Folder Operations. Anyone with USER permissions in an Oracle AI Data Platform workspace can create Git folders or bundle folders. You can grant the following permissions for Git folders and Git bundle folders.

- **READ:** Users can read/list files and folders.
- **USE:** Users can read/write to folders and contained files, and run permitted job types (.ipynb, .py, .sql, .scala, etc).
- **MANAGE:** Users have Read and Use permissions and can rename files/folders and modify files.
- **ADMIN:** Users all permissions and can create, modify, or delete other user permissions.

Table 36-1 Git Folder Permissions

Operation	READ	USE	MANAGE	ADMIN
List Git folder	Yes	Yes	Yes	Yes
View Git folder (metadata)	Yes	Yes	Yes	Yes

Table 36-1 (Cont.) Git Folder Permissions

Operation	READ	USE	MANAGE	ADMIN
Run assets in folder	No	Yes	Yes	Yes
Edit files in a Git folder	No	No	Yes	Yes
Create, delete, rename and move files and folders	No	No	No	Yes
Create a branch in folder	No	No	No	Yes
Switch branches in a folder	No	No	No	Yes
Pull or push a branch into a folder	No	No	No	Yes
Reset, merge, and rebase	No	No	No	Yes
Change Git settings	No	No	No	Yes
Manage user permissions	No	No	No	Yes

Create Git Folder Permissions (Preview)

You can set individual permissions for Git folders in your workspaces.

1. Navigate to your workspace and click the Git folder you want to set permissions for.
2. Click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Git Folder Permissions (Preview)

You can modify permissions for users and roles assigned to your Git folder.

1. Navigate to your workspace and click the Git folder you want to modify permissions for.
2. Click the **Permissions** tab.
3. Next to the user or role you want to modify, click  **Actions** then click **Edit**.
4. Select a new permission level from the dropdown. Click **Save**.

Delete Git Folder Permissions (Preview)

You can remove permissions that are no longer needed for users or roles on your Git folder.

1. Navigate to your workspace and click the Git Folder.
2. Click the **Permissions** tab.
3. Next to the user or role you want to delete, click **...** **Actions** then click **Delete**.
4. Click **Delete**.

Bundle Permissions (Preview)

You can manage which users, roles, and groups can view and modify bundles in your Git folders.

Anyone with USER permissions in a Oracle AI Data Platform workspace can create bundles. You can grant the following permissions for bundles.

- **READ:** Users can read/list bundles.
- **USE:** Users can read/write to bundles, and run sync operations.
- **MANAGE:** Users have Read and Use permissions and can deploy bundles.
- **ADMIN:** Users all permissions and can create, modify, or delete other user permissions.

Note

Users deploying a bundle must have permission to deploy the bundle and write to the configured publish path. If the publish path is not accessible, deployment can fail during validation.

Table 36-2 Bundle Permissions

Operation	READ	USE	MANAGE	ADMIN
List bundle folder	Yes	Yes	Yes	Yes
View bundle	Yes	Yes	Yes	Yes
Run files in bundle folder	No	Yes	Yes	Yes
Edit bundle files	No	No	Yes	Yes
Sync bundle	No	No	Yes	Yes
Deploy bundle	No	No	Yes	Yes
Create, delete, and move assets	No	No	No	Yes
Prune/delete bundle	No	No	No	Yes
Manage user permissions	No	No	No	Yes

Create Bundle Permissions (Preview)

You can set individual permissions for bundles in your Git folders.

1. Navigate to your Git folder and click the bundle you want to set permissions for.
2. Click the **Permissions** tab.
3. Click **+ New Permission**.

4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.
 - For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Git Bundle Permissions (Preview)

You can modify permissions for users and roles assigned to your Git bundle.

1. Navigate to your Git folder and click the Git bundle you want to modify permissions for.
2. Click the **Permissions** tab.
3. Next to the user or role you want to modify, click ... **Actions** then click **Edit**.
4. Select a new permission level from the dropdown. Click **Save**.

Delete Bundle Permissions (Preview)

You can remove permissions that are no longer needed for users or roles on your bundle.

1. Navigate to your Git folder and click the bundle you want to delete permissions for.
2. Click the **Permissions** tab.
3. Next to the user or role you want to delete, click ... **Actions** then click **Delete**.
4. Click **Delete**.

Credential Permissions (Preview)

You can control which users and roles have use, manage, and administration access to your stored credentials.

You create and manage user permissions from the Permissions tab in your stored credentials.

Operation	Use	Manage	Admin
List credentials	Yes	Yes	Yes
View credential details	Yes	Yes	Yes
Edit credentials	No	Yes	Yes
Delete credentials	No	No	Yes
Manage permissions	No	No	Yes

Create Credential Permissions (Preview)

You can control which users and roles have access to credentials you own.

1. Navigate to your Credential Store and click the credential you want to set permissions for.
2. Click the **Permissions** tab.
3. Click  **New Permission**.
4. Select the permissions level and user type from the dropdowns.
5. Select whether to add the user by user name or OCID.

- For **User name**, click Search and enter a user name. Select the user from the list.
 - For **Enter OCID**, enter the OCID of the user.
6. Click **Create**.

Modify Credential Permissions (Preview)

You can modify permissions for users and roles assigned to your credentials.

1. Navigate to your Credential Store and click the credential you want to modify permissions for.
2. Click the **Permissions** tab.
3. Next to the user or role you want to modify, click ... **Actions** then click **Edit**.
4. Select a new permission level from the drop-down list. Click **Save**.

Delete Credential Permissions (Preview)

You can remove permissions that are no longer needed for users or roles on your credentials.

1. Navigate to your Credentials Store and click the credential you want to delete permissions for.
2. Click the **Permissions** tab.
3. Next to the user or role you want to delete, click ... **Actions** then click **Delete**.
4. Click **Delete**.

IAM Policies for Oracle AI Data Platform

Oracle AI Data Platform is managed in OCI and requires the provided IAM policies.



[Video](#)



[LiveLabs Sprint](#)

To create new AI Data Platform instances, a user needs at least **MANAGE** enabled in IAM policies:

```
allow group <aidpAdminIdentityDomain>/<aidpAdminGroup> to manage ai-data-
platforms in compartment id <aidpCompartmentId>
```

Once your AI Data Platform instance is created, you can grant use permissions to any users you want to access your instance. **manage** is only required for the user that creates the instance.

```
allow group <aidpAdminIdentityDomain>/<aidpAdminGroup> to use ai-data-
platforms in compartment id <aidpCompartmentId>
```

AI Data Platform allows users two different combination of policies they can choose from to set up their instance.

Option 1: Tenancy-level Policies (Broad Scope)

With this option, your policies are defined at the **tenancy** (root) level, giving your AI Data Platform instance broad access across compartments.

- Minimizes the need to write new IAM policies every time you add new workloads, data sources, or compartments.
- Easiest onboarding experience; requires the least changes after initial setup.
- Users have a broader scope of permissions.
- May not meet strict least-privilege requirements in regulated environments.

1. Allow AI Data Platform service to view OCI IAM resources to configure role-based access control of AI Data Platform managed resources:

```
allow any-user TO {AUTHENTICATION_INSPECT, DOMAIN_INSPECT, DOMAIN_READ,
DYNAMIC_GROUP_INSPECT, GROUP_INSPECT, GROUP_MEMBERSHIP_INSPECT,
USER_INSPECT, USER_READ} IN TENANCY where all
{request.principal.type='aidataplatform'}
```

2. Allow AI Data Platform service to create OCI logging log group and provide logs to users:

```
allow any-user to manage log-groups in compartment id <aidpCompartmentId>
where ALL { request.principal.type='aidataplatform' }
```

```
allow any-user to read log-content in compartment id <aidpCompartmentId>
where ALL { request.principal.type='aidataplatfom' }
```

3. Allow AI Data Platform service to provide metrics to users:

```
allow any-user to use metrics in compartment id <aidpCompartmentId> where
ALL {request.principal.type='aidataplatfom',
target.metrics.namespace='oracle_aidataplatfom'}
```

4. Allow AI Data Platform service on create and manage OCI Object Store Bucket for workspace and managed data in Master Catalog:

```
allow any-user to manage buckets in tenancy where all
{ request.principal.type='aidataplatfom', any {request.permission =
'BUCKET_CREATE', request.permission = 'BUCKET_INSPECT', request.permission
= 'BUCKET_READ', request.permission = 'BUCKET_UPDATE'}}
```

5. Allow AI Data Platform service to govern/manage data in Workspace and Master Catalog with restricted access to per AI Data Platform instance level:

```
allow any-user to {TAG_NAMESPACE_USE} in tenancy where all
{request.principal.type = 'aidataplatfom'}
allow any-user to manage buckets in tenancy where all
{ request.principal.id=target.resource.tag.orcl-aidp.governingAidpId, any
{request.permission = 'BUCKET_DELETE', request.permission = 'PAR_MANAGE',
request.permission = 'RETENTION_RULE_LOCK', request.permission =
'RETENTION_RULE_MANAGE'}}
allow any-user to read objectstorage-namespaces in tenancy where all
{ request.principal.type='aidataplatfom', any {request.permission =
'OBJECTSTORAGE_NAMESPACE_READ'}}
allow any-user to manage objects in tenancy where all
{ request.principal.id=target.bucket.system-tag.orcl-
aidp.governingAidpId }
```

6. Allow AI Data Platform service to call LLMs hosted on the OCI GenAI service:

Note

If you are upgrading your instance to access AI features, you can either include both listed policies or select the one relevant to your instance. You can determine which principal type (datalake or aidataplatfom) is relevant to your instance by inspecting the OCID of your instance.

9. Allow AI Data Platform service to access and use existing Oracle Autonomous AI Lakehouse instances in a given compartment:

```
allow group <aidpAdminGroup> to use autonomous-databases in <aidp
compartment>
```

Note

Oracle Autonomous AI Lakehouse instances must be 26ai or above to enable AI features in AI Data Platform.

10. Allow AI Data Platform service to configure Compute Cluster to access data in a private network (Optional):

```
allow any-user to manage vnics in compartment id <aidpCompartmentId> where
all { request.principal.type='aidataplatfom'}
allow any-user to use subnets in compartment id <aidpCompartmentId> where
all { request.principal.type='aidataplatfom'}
allow any-user to use network-security-groups in compartment id
<aidpCompartmentId> where all { request.principal.type='aidataplatfom'}
```

11. Allows the Object Storage service to automatically apply lifecycle actions (such as permanent deletion or archival) to your AI Data Platform workspace data, reducing manual maintenance effort and supporting compliance with data retention best practices (Optional):

```
allow service objectstorage-<<region_identifier>> to manage object-family
in compartment id <<aidp-compartment-ocid>>
```

Option 2: Compartment-Level Policies (Fine-grained Scope)

With this option, your policies are defined at the **compartment** level, meaning the compartment where your AI Data Platform instance is created.

- Provides you a tighter security boundary; limits the access of your AI Data Platform to a single compartment by default.
- You can add new compartment policies incrementally when workflows need to span additional compartments.
- Requires you to make manual IAM updates whenever you need your AI Data Platform to access a different compartment.
- Requires more operational overhead during expansion.

1. Allow AI Data Platform service to view OCI IAM resources to configure role-based access control of AI Data Platform-managed resources:

```
allow any-user TO {AUTHENTICATION_INSPECT, DOMAIN_INSPECT, DOMAIN_READ,
DYNAMIC_GROUP_INSPECT, GROUP_INSPECT, GROUP_MEMBERSHIP_INSPECT,
USER_INSPECT, USER_READ} IN TENANCY where all
{request.principal.type='aidataplatfom'}
```

2. Allow AI Data Platform service to create OCI logging log group and provide logs to users:

```
allow any-user to manage log-groups in compartment id <aidpCompartmentId>
where ALL { request.principal.type='aidatapatform' }
allow any-user to read log-content in compartment id <aidpCompartmentId>
where ALL { request.principal.type='aidatapatform' }
```

3. Allow AI Data Platform service to provide metrics to users:

```
allow any-user to use metrics in compartment id <aidpCompartmentId> where
ALL {request.principal.type='aidatapatform',
target.metrics.namespace='oracle_aidatapatform'}
```

4. Allow AI Data Platform service on create and manage OCI Object Store Bucket for workspace and managed data in Master Catalog:

```
allow any-user to manage buckets in compartment id <aidpCompartmentId>
where all { request.principal.type='aidatapatform', any
{request.permission = 'BUCKET_CREATE', request.permission =
'BUCKET_INSPECT', request.permission = 'BUCKET_READ', request.permission =
'BUCKET_UPDATE'}}}
```

5. Allow AI Data Platform service to govern/manage data in Workspace and Master Catalog with restricted access to per AI Data Platform instance level:

```
allow any-user to {TAG_NAMESPACE_USE} in tenancy where all
{request.principal.type = 'aidatapatform'}
allow any-user to manage buckets in compartment id <aidpCompartmentId>
where all { request.principal.id=target.resource.tag.orcl-
aidp.governingAidpId, any {request.permission = 'BUCKET_DELETE',
request.permission = 'PAR_MANAGE', request.permission =
'RETENTION_RULE_LOCK', request.permission = 'RETENTION_RULE_MANAGE'} }
allow any-user to read objectstorage-namespaces in compartment id
<aidpCompartmentId> where all { request.principal.type='aidatapatform',
any {request.permission = 'OBJECTSTORAGE_NAMESPACE_READ'}}
allow any-user to manage objects in compartment id <aidpCompartmentId>
where all { request.principal.id=target.bucket.system-tag.orcl-
aidp.governingAidpId }
```

6. Allow AI Data Platform service to configure Compute Cluster to access data in a private network (Optional):

```
allow any-user to manage vnics in compartment id <aidpCompartmentId> where
all { request.principal.type='aidatapatform'}
allow any-user to use subnets in compartment id <aidpCompartmentId> where
all { request.principal.type='aidatapatform'}
allow any-user to use network-security-groups in compartment id
<aidpCompartmentId> where all { request.principal.type='aidatapatform'}
```

7. Allows the Object Storage service to automatically apply lifecycle actions (such as permanent deletion or archival) to your AI Data Platform workspace data, reducing

manual maintenance effort and supporting compliance with data retention best practices (Optional):

```
allow service objectstorage-<<region_identifier>> to manage object-family
in compartment id <<aidp-compartment-ocid>>
```

Additional Policies for External Tables

If your AI Data Platform instance needs to access data stored in a different compartment, you must grant additional policies for that external compartment. These policies allow AI Data Platform to inspect, read, and manage buckets and objects in the external compartment to use it inside AI Data Platform workspace.

```
allow any-user to manage buckets in compartment id <external-data-
CompartmentId> where all { request.principal.type='aidatapatform', any
{request.permission = 'BUCKET_INSPECT', request.permission = 'BUCKET_READ',
request.permission = 'BUCKET_UPDATE'}}
allow any-user to manage buckets in compartment id <external-data-
CompartmentId> where all { request.principal.id=target.resource.tag.orcl-
aidp.governingAidpId, any {request.permission = 'PAR_MANAGE',
request.permission = 'RETENTION_RULE_LOCK', request.permission =
'RETENTION_RULE_MANAGE'} }
allow any-user to manage objects in compartment id <external-data-
CompartmentId> where all { request.principal.id=target.bucket.system-tag.orcl-
aidp.governingAidpId }
allow service objectstorage-<<region_identifier>> to manage object-family in
compartment id <external-data-CompartmentId>
```

Note

If you are using a custom identity domain (non-default), you must prefix the group name with the domain name in your IAM policy. For example:

```
allow group <aidpAdminIdentityDomain>/<aidpAdminGroup> to manage ai-data-
platforms in compartment id <aidpCompartmentId>
```

For more information on IAM policies, see [IAM Policies Overview](#).

To see and login to an AI Data Platform, you need to be granted access by the administrator of that AI Data Platform.

38

Roles

This chapter describes how to use Oracle AI Data Platform role-based access controls (RBAC) to manage user roles and access.

Topics:

- [About Roles](#)
- [Map Active Directory Groups to IAM Groups](#)
- [Create a Role](#)
- [Modify a Role](#)
- [Delete a Role](#)
- [Assign Members to a Role](#)
- [Remove Member from Role](#)

About Roles

Oracle AI Data Platform lets you manage your users and permissions using role based access controls (RBAC).



[LiveLabs Sprint](#)



[Video](#)

You manage RBAC through the Roles interface, where you can create new roles, modify existing, or delete unused roles. After you've provisioned a role, you can assign members by individual user, group, or other roles. You can review and modify the assigned members for any role you have created. You can check the permissions assigned to the role from the Permissions tab.

By default, AI Data Platform has two system roles, `AI_DATA_PLATFORM_ADMIN` and `AUDITOR`:

- `AI_DATA_PLATFORM_ADMIN` is automatically assigned to the user that created the data platform. This user has administrator permissions to all data platform objects and can grant or revoke permissions to other users, groups, or AI Data Platform roles. To create an AI Data Platform instance you need `MANAGE AI Data Platform IAM` permission.
- `AUDITOR` users are able to view the entire audit trail of objects in your AI Data Platform. The `AI_DATA_PLATFORM_ADMIN` is automatically made a member of the `AUDITOR` role when you create your AI Data Platform instance. Any users added to the `AI_DATA_PLATFORM_ADMIN` role are added to the `AUDITOR` role as well.

Note

Your AI Data Platform instance can only have one `AI_DATA_PLATFORM_ADMIN` system role. If the `AI_DATA_PLATFORM_ADMIN` role needs to pass to another user, a user with `MANAGE AI Data Platform IAM` permissions can reassign it to another user by logging in to OCI and viewing the details of the AI Data Platform instance.

RBAC permissions are passed down to contained objects. Permissions granted at the Master Catalog level cascade down to all contained objects.

Map Active Directory Groups to IAM Groups

To map Active Directory (AD) groups to Oracle Cloud Infrastructure (OCI) Identity and Access Management (IAM) groups, you need to establish a federation between your AD and your OCI tenancy.

To map AD groups to IAM groups, see [Federating with Microsoft Active Directory](#).

This process involves creating mappings between AD groups and corresponding IAM groups in OCI, allowing users in your AD groups to access OCI resources with appropriate permissions. Once federated, your AD groups are visible in OCI and you can add group mappings by following the steps in **To add group mappings for an identity provider** under [Managing Identity Providers in the Console](#).

Once you have added group mappings, you can assign permissions to IAM groups in Oracle AI Data Platform.

Create a Role

You can create new role as part of RBAC management.

1. On the Home page, click **Roles**.
2. Click  **New Role**.
3. Provide a name and description for the role.
4. Click **Create**.

Modify a Role

You can modify settings of a role you own.

1. Navigate to **Roles**.
2. Next to the role you want to modify, click ... **Actions** then **Edit**.
3. Make your changes to the role, then click **Save**.

Delete a Role

You can delete Oracle AI Data Platform roles that you own.

1. Navigate to **Roles**.
2. Next to the role you want to delete, click ... **Actions** then **Delete**.

3. Click **Delete**.

Assign Members to a Role

You can assign users, groups, or other roles to a role you created.

1. Navigate to **Roles** and click the role you want to add members to.
2. Click **Members** then click  **Add Members**.
3. From the Principle Type, select **User**, **Group**, or **Role**.
 - For **User**, you search for an individual user by name or you provide the OCID of the user.
 - To assign a user by name, you select the compartment and domain then select the user from the list. Enter the user name in the search bar to narrow results.
 - To assign a user by OCID, enter their OCID in the provided field.
 - For **Group**, you search for a group name or you provide the OCID of the group.
 - To assign a group by name, you select the compartment and domain then select the group from the list. Enter the group name in the search bar to narrow results.
 - To assign a group by OCID, enter their OCID in the provided field.
 - For **Role**, you select a role from the list provided.
4. Click **Create**.

Remove Member from Role

You can remove assigned members from roles you own.

1. Navigate to **Roles** and click the role you want to remove members from.
2. Click **Members**.
3. Next to the member you want to remove, click ... **Actions** then **Remove assignee**.
4. Click **Delete**.

Network Sources

You can enable network source-based access control for your Oracle AI Data Platform instances.

A network source is a set of trusted IP addresses. The IP addresses can be public IP addresses or IP addresses from Virtual Cloud Networks (VCNs) within your tenancy. OCI Network Source support in AI Data Platform enables defining a set of IP ranges (public or VCN-based) as trusted IP addresses from which traffic should be allowed. Your tenancy admins can then define IAM policies to control traffic to your private networks only from the originating IP ranges that are related to your AI Data Platform use-cases.

To restrict access to requests made from a set of IP addresses or network sources, you need to follow these steps:

1. Create a network source that defines the allowed IP addresses for controlling access to your AI Data Platform instances. For setup instructions for network source, refer to [OCI Network Source documentation](#).
2. Write a policy that uses the network source variable in a condition. The following is an example of the IAM policy including a network source variable you can add to the your existing AI Data Platform-related IAM policies:

```
allow group <aidpAdminIdentityDomain>/<aidpAdminGroup> to manage ai-data-  
platform-family in tenancy where all {request.networkSource.name=<AIDP-  
related-network-source-name>}
```

Part IX

Administration

This section covers the administration of your Oracle AI Data Platform and its contents.

In order to create and manage AI Data Platform instances, you will need to have an existing cloud account (tenancy) with Oracle Cloud Infrastructure (OCI) and have familiarity with the essential OCI concepts. For more information, see [Service Essentials](#).

Workspaces

A workspace in AI Data Platform acts as an isolated logical container where users can manage and organize their data resources, including workflows, notebooks, and libraries. Workspaces enable efficient collaboration and governance by keeping resources grouped logically.

Networking

The data assets that are read from and written to by your AI Data Platform can be in an OCI virtual cloud network. For more information, see [Virtual Cloud Network](#).

AI Data Platform workspaces can be enabled for a private network where one or more target data assets belong to. Enabling for private networks allows workspaces to access target data assets located in the private network. In order to enable private network for a workspace, you will need three networking related details from your administrator: the existing virtual cloud network of the data asset, the subnet and the network security group.

In order to create an external catalog for a data asset that is in a private network, you will first need to create a workspace enabled for the same private network.

Security and Identity

In order to create, manage or use AI Data Platform instances, you need to have either a local user or federated user credential for console. The existing security features (e.g. Multi factor authentication or MFA) and access control features (e.g. pre-requisite policies, sign-on policies) are applicable for AI Data Platform in OCI Console. For more information, review the following OCI documentation:

- [User Credentials](#)
- [Federating with Identity Providers](#)
- [Securing IAM](#)
- [Managing Access to Resources](#)

On top of OCI Security and Identity and Access Management, AI Data Platform implements additional Role Based Access Control (RBAC) to enforce fine-grained access control across different resources. In your AI Data Platform, you can define roles and permissions for workspaces, catalogs, and datasets to ensure secure collaboration.

Limitations

AI Data Platform implements and manages limits for the resources created and managed by AI Data Platform instances.

Topics:

- [Workspaces](#)

- [Administration Settings \(Preview\)](#)
- [Notifications](#)
- [Network Sources](#)
- [Limits](#)
- [Known Issues](#)

40

Pricing

AI Data Platform Unit (AIDP Unit) is the unit of measurement used to meter consumption in your Oracle AI Data Platform.



AIDP Unit is a measure of the AI, data management, and data processing work done in your AI Data Platform instance. The rate at which AIDP Units are charged depends on the use of OCPU and memory resources per hour. 100 GB block volume is included per OCPU.

You can see how different cluster configurations are priced by using our [Pricing Widget](#).

Table 40-1 AI Data Platform Units Mapping

AI Data Platform Definition	AIDP Units Charged
AMD OCPU per Hour (with included 100 GB Block storage per OCPU)	67 AIDP Units
ARM OCPU per Hour (with included 100 GB Block storage per OCPU)	27 AIDP Units
Intel OCPU per Hour (with included 100 GB Block storage per OCPU)	87 AIDP Units
NVIDIA GPU per Hour (with included Block Volume, CPU, Memory)	4110 AIDP Units
AMD Memory GB per hour	3 AIDP Units
ARM Memory GB per hour	
Intel Memory GB per hour	

AMD shapes are used for AI computes. Use the unit mapping for AMD in the table above to calculate the cost of AI computes.

You pay AIDP Units to AI Data Platform for all your compute clusters, AI compute, and features based on the definitions in the table above. For object storage, logging, metrics, Generative AI models, or any databases on OCI, you pay directly to OCI.

For all-purpose compute clusters, you select Compute Shape options between AMD, Intel, ARM, and NVIDIA GPUs. More specific or granular shapes are not currently available. Your selection of compute shapes is subject to the availability of compute shapes in the OCI region where your AI Data Platform instance is located. Oracle reserves the right to change the specific compute shapes available in an OCI region.

For AI compute hosting agent flows, you pick the number of OCPUs and memory. Compute shape is set as AMD for all AI compute.

All-purpose compute clusters are a combination of compute OCPU and memory. When you create an all-purpose compute cluster in AI Data Platform, you see an estimated cost of that compute cluster per hour in AIDP Units. For example:

- The cost Per Hour for an AMD 2 OCPU 32 GB Memory cluster will be 2 OCPU x 67 AIDP Units + 32 GB Memory x 3 AIDP Units = 230 AIDP Units.

- The cost Per Hour for an Intel 4 OCPU 32 GB Memory cluster will be $4 \text{ OCPU} \times 87 \text{ AIDP Units} + 32 \text{ GB Memory} \times 3 \text{ AIDP Units} = 444 \text{ AIDP Units}$.
- The cost Per Hour for an ARM 4 OCPU 32 GB Memory cluster will be $4 \text{ OCPU} \times 27 \text{ AIDP Units} + 32 \text{ GB Memory} \times 3 \text{ AIDP Units} = 204 \text{ AIDP Units}$.
- The cost Per Hour for a NVIDIA GPU cluster with 2 GPU will be $2 \text{ GPU} \times 4110 \text{ AIDP Units} = 8220 \text{ AIDP Units}$.

AI Computes are a combination of OCPU and Memory. Examples:

- The cost Per Hour for a 1 OCPU and 16GB Memory AI Compute will be $1 \text{ OCPU} \times 67 \text{ AIDP Units} + 16 \text{ GB Memory} \times 3 \text{ AIDP Units} = 115 \text{ AIDP Units}$.
- The cost Per Hour for a 2 OCPU and 32GB Memory AI Compute will be $2 \text{ OCPU} \times 67 \text{ AIDP Units} + 32 \text{ GB Memory} \times 3 \text{ AIDP Units} = 230 \text{ AIDP Units}$.

By default, one Master Catalog Default Cluster is present in each AI Data Platform instance. This cluster is responsible for the essential AI Data Platform functions, like search crawls, refreshing catalog objects, creating, editing, and deleting objects, testing connections etc. Default Master Catalog Compute Clusters have AMD 2 OCPU 32 GB Memory. For the Default Master Catalog Compute Cluster, you are charged hourly for a minimum of 230 AIDP Units. Deleting an AI Data Platform instance halts the ongoing cost of that Default Master Catalog Compute Cluster. If you add more capacity to the Default Master Catalog Compute cluster to meet your performance and scalability needs, the AIDP Units cost of the Default Master Catalog scales up relative to the increase.

Sample Pricing Scenario

In this scenario, you have 2 compute clusters running for a month (all calculations are in USD):

- 1 Default Master Catalog Compute Cluster: AMD 2 OCPU with 32 GB Memory
- 1 custom all-purpose compute cluster for workloads: Intel 4 OCPU with 32 GB Memory
- 1 AI compute for hosting agent flows: 1 OCPU with 16GB Memory

In this case, the cost calculation for your AMD Cluster (for Default Master Catalog Compute Cluster) is:

- $2 \text{ OCPUs } 32 \text{ GB Memory Per Hour} = (2 \times 67 + 32 \times 3) \text{ AIDP Units} = 230 \text{ AIDP Units}$
- $230 \text{ AIDP Units/hour} = 171,120 \text{ AIDP Units/month}$
- $\$0.230/\text{hour} = \$171.12/\text{month}$

The cost calculation for your all-purpose compute cluster (for workloads) is:

- $4 \text{ OCPUs } 32 \text{ GB Memory Per Hour} = (4 \times 87 + 32 \times 3) \text{ AIDP Units} = 444 \text{ AIDP Units}$
- $\text{AIDP Units } 444 \text{ AIDP Units/hour} = 330,370 \text{ AIDP Units/month}$
- $\$0.444/\text{hour} = \$330.37/\text{month}$

Calculation for AI compute (for hosting agent flows)

- $1 \text{ OCPU } 16\text{GB Memory Per Hour} = (1 \times 67 + 16 \times 3) \text{ AIDP Units} = 115 \text{ AIDP Units}$
- $115 \text{ AIDP Units / hour} = 85,560 \text{ AIDP Units / month}$
- $\$0.115 / \text{hour} = \$85.56 / \text{month}$

Adding your monthly costs together, your total monthly costs would amount to: $\$171.12/\text{month} + \$330.37/\text{month} + \$85.56/\text{month} = \$587.05/\text{month}$.

Administration Settings (Preview)

You can manage administration level settings for your Oracle AI Data Platform from your Settings menu.

The Administration section of the Settings menu in AI Data Platform enables you to monitor, enable, and disable important aspects of your instance.

Git Linked Accounts (Preview)

You can use Git linked accounts to add your own Git credentials in Oracle AI Data Platform.

After you add a linked account, you can use those credentials when working with remote Git repositories from AI Data Platform folders. Git settings are modified from the **Settings** page on the left navigation pane.

Best Practices for Git Credentials

- Use a clear credential name so you can easily identify the account later.
- Use a valid personal access token created in your Git provider.
- If you manage multiple accounts, set the most commonly used account as the default.
- Update or delete credentials when tokens expire or are rotated.

Add Git Credentials by Using Personal Access Tokens (Preview)

You can add your Git account credentials to your user settings by using a personal access token.

1. On the home page, click **Settings**.
2. Under **User**, click **Git linked accounts**.
3. Click **Add credential**.
4. Provide a descriptive name for this credential. For example, Github Personal.
5. Select **Use personal access token**.
6. Provide the email address for your Git provider.
7. Provide your personal access token.
8. Click **Add**.

Set Default Git Credentials (Preview)

You can specify a set of Git credentials added to Oracle AI Data Platform settings as the default credentials used when connecting to Git.

1. On the home page, click **Settings**.
2. Under **User**, click **Git linked accounts**.
3. Click the **Actions** next to the Git credentials you want to set as default.

4. Click **Set as default**.

Edit Git Credentials (Preview)

You modify Git credentials you have previously added to Oracle AI Data Platform settings.

1. On the home page, click **Settings**.
2. Under **User**, click **Git linked accounts**.
3. Click the **Actions** next to the Git credentials you want to modify.
4. Click **Edit**.
5. Modify your Git credentials, then click **Save**.

Delete Git Credentials (Preview)

You delete Git credentials to remove redundant or unwanted credentials from your Oracle AI Data Platform settings.

1. On the home page, click **Settings**.
2. Under **User**, click **Git linked accounts**.
3. Click the **Actions** next to the Git credentials you want to delete.
4. Click **Delete**.
5. Click **Delete**.

AI Feature Management

You can choose to either enable or disable AI support in your existing Oracle AI Data Platform instance.

AI Data Platform requires an Oracle Autonomous AI Lakehouse instance of version 26ai or above to enable AI features like knowledge bases, agent flows, vector search and embedding storage, and retrieval-augmented generation (RAG) workflows.

If you didn't create or link an Autonomous AI Lakehouse instance to your AI Data Platform instance at creation, you can create and configure one or select an existing instance from your home page.

To create or choose an Autonomous AI Lakehouse as part of creating your AI Data Platform, see [Create Your First Oracle AI Data Platform](#).

When disabling AI features, if the Autonomous AI Lakehouse was provisioned as part of creating your AI Data Platform instance or created as part of the upgrade flow in [Create an Oracle Autonomous AI Lakehouse Instance to Enable AI](#), you have two options:

- You can disable and detach the Autonomous AI Lakehouse, leaving the instance available to reattach later to enable AI features, or
- You can disable, detach, and permanently delete the Autonomous AI Lakehouse instance.

If you selected an existing Autonomous AI Lakehouse as part of creating your AI Data Platform instance or as part of the upgrade flow in [Choose an Existing Oracle Autonomous AI Lakehouse Instance to Enable AI](#), you can only choose to disable and detach your Autonomous AI Lakehouse.

Note

AI Data Platform deletes any existing AI resources when disabling AI features, including agent flows, AI compute, and knowledge bases.

If your Autonomous AI Lakehouse instance is in use for data or gold-layer workloads, choosing to disable and detach the instance allows the instance to continue supporting those workloads while still disabling support for AI features.

Note

If the option to disable AI features is not shown on your home page, contact your AI Data Platform support to have the option enabled in your environment.

You can re-enable AI features by clicking **Upgrade** from the AI Integration card on your home page.

Create an Oracle Autonomous AI Lakehouse Instance to Enable AI

You can create and attach an Oracle Autonomous AI Lakehouse instance to your Oracle AI Data Platform to enable AI features.

New Oracle Autonomous AI Lakehouse instances created by this method are automatically 26ai or above.

1. On the Home page, in the AI Integration card, click **Upgrade**
2. From the **ALH Instance** drop-down list, select **Create new**.
3. Enter a password for the ADMIN user of your new Autonomous AI Lakehouse instance.
4. Enter the group name for any listed IAM policies
5. Click **Add** for any required IAM policies.
6. Click **Upgrade**.

Choose an Existing Oracle Autonomous AI Lakehouse Instance to Enable AI

You can choose an existing Oracle Autonomous AI Lakehouse instance to enable AI features in your Oracle AI Data Platform instance.

Your Autonomous AI Lakehouse instance must be 26ai or above to enable AI in your AI Data Platform instance.

1. On the Home page, in the AI Integration card, click **Upgrade**.
2. From the **ALH Instance** drop-down list, select **Choose existing**.
3. Select the compartment and instance for your Autonomous AI Lakehouse.
4. Provide the password for the ADMIN user of your existing Autonomous AI Lakehouse instance.
5. Enter the group name for any listed IAM policies
6. Click **Add** for any required IAM policies.
7. Click **Upgrade**.

Disable AI Features

You can choose to disable support for AI features in your Oracle AI Data Platform and detach the associated Oracle Autonomous AI Lakehouse.

Note

Any AI resources, such as agent flows, AI compute, and knowledge bases are deleted when disabling AI features in your AI Data Platform instance.

1. On the homepage, in the Disable AI features card, click **Disable**.

2. Select an option for **ALH instance detach options**:
 - **Detach and keep the instance:** The Autonomous AI Lakehouse instance is detached and AI features are disabled. You can choose to enable AI features with the same instance later.
 - **Detach and delete the instance:** The Autonomous AI Lakehouse instance is detached, AI features are disabled, and the Autonomous AI Lakehouse instance is permanently deleted. You will have to create a new instance or select another existing instance if you want to re-enable AI features.
3. Click **Disable**.

Notifications

Oracle AI Data Platform keeps you updated about the status of your long running operations through notifications to the users or administrators of workspaces and catalogs in your AI Data Platform instance.

Notifications keep you in the loop on the events that matter most to you related to operations you initiated. You get an instant heads-up when something finishes, breaks, or needs your attention.

Notifications in AI Data Platform are one of two types, based on the audience:

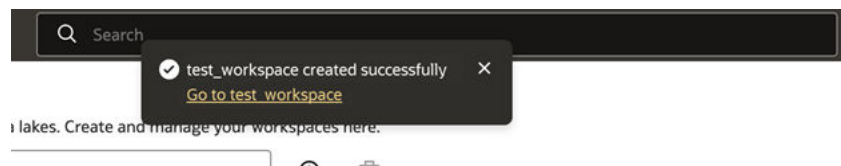
- Notifications to users of workspace or catalogs, for the success, failure, or in-progress status of user initiated operations
- Notifications to administrators of workspaces or catalogs for unexpected events (warnings, errors) detected by the AI Data Platform service

Notifications tell you what's happening with your long-running or asynchronous tasks. Notifications let you know if something you started, scheduled, or set up in a workflow has succeeded, failed, or is still in progress. Notifications can be related to different user-triggered actions, such as:

- Creating catalogs, schema, tables, volumes, workspaces, and compute clusters
- Starting or stopping compute clusters
- Starting, progress, and completion of file uploads

AI Data Platform also sends notifications to workspace or catalog admins about unexpected errors or warnings affecting resources across all the workspaces or catalogs they manage. These are typically error or warning notifications regarding health, capacity, and performance related issues that users don't directly trigger. For example, your AI Data Platform sends notifications when an external catalog can't authenticate, or certain cluster shapes become unavailable, or if a Default Master Catalog cluster hits any issue. These notifications can provide you early warning signals when workflow jobs are running with no data or clusters are slowing down under heavy load, along with clear next-step guidance, like advising an increase OCPUs or contacting Oracle support. These proactive messages let you fix problems before they disrupt your workloads.

Notifications appear as brief toast messages and disappear automatically after a short period. When you see a notification, you can click on it to view more details, you can ignore it and let it disappear, or you can click the X to close and dismiss it.



All the notifications from the last 14 days are stored in your Notifications page, so you don't need to worry about missing a notification while you are away or attend to every notification immediately while you're occupied with other tasks.

Intelligent Data Lake - IDL-SCT

Create

Home

Master catalog

Workspace

Data-Playground

Workflow

Compute

Data Sharing

Auto populate catalog

Notifications

Roles

Audit logs

Activity

Unlisted_1746855049...

CalendarPrinter_1748...

CalendarPrinter_1748...

Dev_cluster_1

CalendarPrinter_1748...

Unlisted_1746855955...

Dev_1747054498540...

etl_notebook.ipynb

Execute Oracle ADW S...

OCJ_Gen_AI_Sentiment...

OCJ_Gen_AI_Sentiment...

Notifications

Notifications for workspace and catalogs for your requests are shown below for the last 14 days

Resource type

All

Operation type

All

Status

All

From

All

To

All

Filter Resource name

Create Table

Edit Table

Resource name

Delete Table

Create Workspace

Edit Workspace

Delete Workspace

Create Cluster

Edit Cluster

Start Cluster

Run Workflow Job

File upload

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Data-Playground > Workflow > Scheduled calendar job job > Schedu...

Operation type

Status

Message

Time

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Mon, Jul 28, 2025 at 12:34:12 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Sun, Jul 27, 2025 at 12:34:13 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Fri, Jul 26, 2025 at 12:34:13 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Fri, Jul 25, 2025 at 12:34:14 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Thu, Jul 24, 2025 at 12:34:14 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Wed, Jul 23, 2025 at 12:34:13 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Tue, Jul 22, 2025 at 12:34:13 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Sat, Jul 21, 2025 at 12:34:13 UTC

job > Sche...

RUN_WORKFLOW_JOB

Failed

Scheduled calendar job job ran failed with the error - Please resolve ...

Sat, Jul 19, 2025 at 12:33:43 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Fri, Jul 18, 2025 at 12:34:12 UTC

job > Sche...

RUN_WORKFLOW_JOB

Failed

Scheduled calendar job job ran failed with the error - Please resolve ...

Thu, Jul 17, 2025 at 12:33:44 UTC

job > Sche...

RUN_WORKFLOW_JOB

Failed

Scheduled calendar job job ran failed with the error - Please resolve ...

Wed, Jul 16, 2025 at 12:33:45 UTC

job > Sche...

RUN_WORKFLOW_JOB

Failed

Scheduled calendar job job ran failed with the error - Please resolve ...

Tue, Jul 15, 2025 at 12:33:44 UTC

job > Sche...

RUN_WORKFLOW_JOB

Success

Scheduled calendar job job ran successfully

Mon, Jul 14, 2025 at 12:34:11 UTC

The count in the notification page indicates the number of notifications that are new for you. After you visit the notification page, the notification count will reset. In the notification page, you can filter by resource types, operation types, status and the date-time ranges. Filtering helps you narrow down to the displayed notifications so you can locate specific notifications you want to review. If you are not able to view the full content of a notification message, you can hover over it to view the complete text. Clicking on the notification takes you to the respective resource details page.

Limits

Oracle AI Data Platform imposes limits on its resources to ensure efficient and error-free use.

All limits are independently applied. The following examples illustrate how limits are applied:

- **Example 1:**
 - You have 1 AI Data Platform instance with of 10 workspaces and 5 workspaces enabled for private networks, or
 - You have 2 AI Data Platform instances with 5 workspaces each (totaling 10 workspaces). The first AI Data Platform has 5 workspaces enabled for private networks. Then the second one can have only 3 workspaces enabled for private networks, to remain within the limit for up to 8 workspaces enabled for private networks.
- **Example 2:**
 - You have 1 all-purpose compute clusters with 100 driver and worker nodes, or
 - You have 25 all-purpose compute clusters each with 4 nodes (totalling 100 nodes across all clusters).

The following table lists various numerical limits for AI Data Platform resources in a single tenancy.

Resource	Metric	Scope	Limit
AI Data Platform	Maximum number of AI Data Platform	Region	100
Workspace	Maximum number of workspaces Includes default workspace and workspaces enabled for private network	Region	100
Workspaces enabled for private network	Number of workspaces enabled for private network	Region	8
Catalog	Maximum number of catalogs that can be created	AI Data Platform	1000
Schema	Maximum number of schema that can be created	Catalog	10000
Table	Maximum number of tables that can be created	Schema	100000
Volume	Maximum number of volumes that can be created	Schema	10000

Resource	Metric	Scope	Limit
Jobs	Maximum number of jobs that can be created per hour	Workspace	10000
Jobs: Concurrent tasks	Maximum number of tasks that are running simultaneously	Workspace	1000
Jobs: Saved jobs	Maximum number of jobs that can be saved	Workspace	12000
Jobs: Tasks	Maximum number of tasks Excludes tasks inside a pipeline task.	Job	100
Jobs: Nesting level	Maximum levels of nesting	Job	3
Jobs: Schedule	Minimum frequency	Job	30 minutes
Jobs: Log retention period	Maximum number of days JobRun logs will be retained	AI Data Platform	30
Workspace file	Maximum number of files Includes notebooks	Workspace	10000
Workspace folder	Maximum number of folders	Workspace	10000
Workspaces enabled for private network	Number of SCAN proxies per workspace	Workspace	15
Workspace file: File size	Maximum size allowed for a single file	Workspace	500 MB
Workspace folder: folder depth	Maximum depth of folders	Workspace	25
Workspace: Object name length	Maximum length of an object name	Workspace	500 characters
Compute: Clusters	Maximum number of clusters	Region	200
Compute: Driver & Worker node	Maximum total number of driver and worker nodes	Region	500
Compute: OCPU per node: (driver or worker)	Maximum number of OCPU	Node (Driver or worker)	64
Compute: Memory per OCPU	Memory per OCPU	OCPU	Up to 16 times OCPU (in GB)
Compute: Block storage per OCPU	Block storage per OCPU	OCPU	100 times OCPU (in GB)
Compute: NVIDIA GPU	Maximum number of NVIDIA GPU	Region	3
Compute: Logs retention	Maximum number of days compute logs will be retained	AI Data Platform	30
AI Features: Knowledge Bases	Number of knowledge bases per AI Data Platform instance	Service	1000
AI Features: Agents	Number of agents per AI Data Platform instance	Service	1000

Resource	Metric	Scope	Limit
AI Features: AI Compute	Maximum AI Computes per tenancy per region	Service	500
AI Compute: Maximum OCPUs per AI Compute	Maximum OCPUs per AI Compute	AI Compute	64
AI Compute: Maximum memory per AI Compute OCPU	Memory per AI Compute OCPU	OCPU	Up to 16 times OCPU (in GB)

Part X

Resources

This section details additional documentation resources you can review for Oracle AI Data Platform, like API reference, known issues, and FAQs.

Topics:

- [Use APIs, SDK, and CLI](#)
- [Known Issues](#)

Use APIs, SDK, and CLI

Oracle AI Data Platform offers REST APIs, SDK, and CLI to integrate with your applications.

Topics:

- [REST APIs](#)
- [SDK and CLI](#)

REST APIs

Use the REST APIs to automate processes and programmatically access features and functionality in Oracle AI Data Platform.

API Version

The base path of the endpoint includes the API version (for example, 20260430). Here's an example GET request to get workbook details for a given workbook:

GET

```
https://aidp.<oci-region>.oci.oraclecloud.com/api/20260430/aiDataPlatforms/  
<aiDataPlatformId>/workspaces/<workspaceKey>
```

REST API Reference

See [REST API for Oracle AI Data Platform](#).

SDK and CLI

Oracle AI Data Platform SDK and CLI provide generated clients and command line tools for working with Oracle AI Data Platform public APIs.

SDK and CLI Reference

See [Oracle AI Data Platform SDK and CLI](#). For language-specific documentation, refer to the following:

Language	Documentation	Installation
CLI	AIDP CLI Command Reference	AIDP CLI Installation
Java	AIDP Java SDK Operations Reference	AIDP Java Installation
Node / TypeScript	AIDP TypeScript SDK Operations Reference	AIDP TypeScript Installation
Python	AIDP Python SDK Operations Reference	AIDP Python Installation

Reference

This section provides additional reference material on Oracle AI Data Platform.

Topics:

- [SQL Grammar](#)
- [Aidutils for Agents and Tools Sample Code](#)
- [Agent LangGraph Code Samples](#)

SQL Grammar

Oracle AI Data Platform users can use SQL to automate their DDL workloads.

Topics:

- [Catalog SQL Grammar](#)
- [Schema SQL Grammar](#)
- [Volume SQL Grammar](#)
- [Table SQL Grammar](#)
- [DML Queries](#)

Catalog SQL Grammar

Catalog objects support the listed SQL grammar for DDL workloads.

Oracle AI Data Platform supports all standard Spark SQL data types. For more information, see [Apache Spark Documentation - Supported Data Types](#).

Table 46-1 Standard and External Catalog SQL Grammar

Operation	Grammar
Create Catalog	Catalog <pre>CREATE CATALOG [IF NOT EXISTS] <<catalog_name>> [PROPERTIES (DESCRIPTION = description)] OPTIONS ({ option_name = option_value } [, ...])</pre> External Catalog <pre>CREATE EXTERNAL CATALOG [IF NOT EXISTS] <<catalog_name>> [PROPERTIES (DESCRIPTION description)]OPTIONS ({ option_name = option_value } [, ...])</pre> OPTIONS will have connection details External Catalog - Oracle Autonomous AI Lakehouse Example <pre>wt = base64 encoded wallet contents create_sql="create external catalog if not exists catalog_adw options ('wallet.content' = '{wt}', 'type' = 'ORACLE_ADW', 'user.name' = 'ADMIN', 'tns' = 'adw23ai_high', 'password' = 'xxxxx', 'wallet.password' = 'xxxxx')"</pre> Response <pre>Catalog <<catalog_name>> created successfully</pre> Error <pre><<SQL Command>> failed due to <<reason>></pre>

Table 46-1 (Cont.) Standard and External Catalog SQL Grammar

Operation	Grammar
Alter Catalog	Alter catalog name
	ALTER CATALOG old_catalog_name RENAME new_catalog_name;
	Alter catalog description
	ALTER CATALOG <catalog-name> set properties (DESCRIPTION=<property- value>)
	Alter catalog options (conn)
	ALTER CATALOG <catalog-name> set options (option_name = option_value)
Delete Catalog	Response
	Catalog <<catalog_name>> updated successfully
	Error
	<<SQL Command>> failed due to <<reason>>
Delete Catalog	Drop catalog
	DROP CATALOG [IF EXISTS] catalog_name
	By default during DROP catalog, all child objects will also get deleted
	Response
	Catalog <<catalog_name>> dropped successfully
Delete Catalog	Error
	<<SQL Command>> failed due to <<reason>>

Table 46-1 (Cont.) Standard and External Catalog SQL Grammar

Operation	Grammar								
List Catalogs	SHOW CATALOGS [[LIKE] [regex_pattern] [TYPE = EXTERNAL CATALOG CATALOG] regex_pattern: A regular expression pattern that is used to filter the results of the statement. Response: <table><tr><th>Catalog</th><th>Type</th></tr><tr><td><<catalog_name>></td><td>Catalog External Catalog</td></tr><tr><td><<catalog_name>></td><td>Catalog External Catalog</td></tr><tr><td><<catalog_name>></td><td>Catalog External Catalog</td></tr></table> Error <<SQL Command>> failed due to <<reason>>	Catalog	Type	<<catalog_name>>	Catalog External Catalog	<<catalog_name>>	Catalog External Catalog	<<catalog_name>>	Catalog External Catalog
Catalog	Type								
<<catalog_name>>	Catalog External Catalog								
<<catalog_name>>	Catalog External Catalog								
<<catalog_name>>	Catalog External Catalog								

Table 46-1 (Cont.) Standard and External Catalog SQL Grammar

Operation	Grammar																																				
Describe Catalog	DESC CATALOG <<catalog_name>> DESCRIBE CATALOG <<catalog_name>> Response (Standard Catalog):<table><tr><th>Attribute</th><th>Value</th></tr><tr><td>Name</td><td>Standard catalog name</td></tr><tr><td>Type</td><td>Standard Catalog</td></tr><tr><td>Description</td><td>Standard catalog description</td></tr><tr><td>Created by</td><td>Principal that created the standard catalog</td></tr><tr><td>Created on</td><td>Date and time created</td></tr><tr><td>Updated by</td><td>Principal that last updated the standard catalog</td></tr><tr><td>Updated on</td><td>Date and time last updated</td></tr></table> Response (External catalog):<table><tr><th>Attribute</th><th>Value</th></tr><tr><td>Name</td><td>External catalog name</td></tr><tr><td>Type</td><td>External Catalog</td></tr><tr><td>Source type</td><td>Source of external catalog (e.g. Oracle Autonomous AI Lakehouse)</td></tr><tr><td>Description</td><td>External catalog description</td></tr><tr><td>Created by</td><td>Principal that created the external catalog</td></tr><tr><td>Created on</td><td>Date and time created</td></tr><tr><td>Updated by</td><td>Principal that last updated the external catalog</td></tr><tr><td>Updated on</td><td>Date and time last updated</td></tr><tr><td>Connection details</td><td>Connection .json file</td></tr></table> Error: <<SQL Command>> failed due to<<reason>>	Attribute	Value	Name	Standard catalog name	Type	Standard Catalog	Description	Standard catalog description	Created by	Principal that created the standard catalog	Created on	Date and time created	Updated by	Principal that last updated the standard catalog	Updated on	Date and time last updated	Attribute	Value	Name	External catalog name	Type	External Catalog	Source type	Source of external catalog (e.g. Oracle Autonomous AI Lakehouse)	Description	External catalog description	Created by	Principal that created the external catalog	Created on	Date and time created	Updated by	Principal that last updated the external catalog	Updated on	Date and time last updated	Connection details	Connection .json file
Attribute	Value																																				
Name	Standard catalog name																																				
Type	Standard Catalog																																				
Description	Standard catalog description																																				
Created by	Principal that created the standard catalog																																				
Created on	Date and time created																																				
Updated by	Principal that last updated the standard catalog																																				
Updated on	Date and time last updated																																				
Attribute	Value																																				
Name	External catalog name																																				
Type	External Catalog																																				
Source type	Source of external catalog (e.g. Oracle Autonomous AI Lakehouse)																																				
Description	External catalog description																																				
Created by	Principal that created the external catalog																																				
Created on	Date and time created																																				
Updated by	Principal that last updated the external catalog																																				
Updated on	Date and time last updated																																				
Connection details	Connection .json file																																				

Table 46-1 Standard and External Catalog SQL Grammar

Operation	Grammar
Refresh Catalog	<pre>REFRESH [EXTERNAL] CATALOG <<catalog_name>></pre> Response: <pre>Refresh Initiated</pre>

Schema SQL Grammar

Schema support the listed SQL grammar for DDL workloads.

Oracle AI Data Platform supports all standard Spark SQL data types. For more information, see [Apache Spark Documentation - Supported Data Types](#).

Table 46-2 List of Schema SQL Grammar

Operation	Grammar
Create Schema	<pre>CREATE SCHEMA [IF NOT EXISTS] catalog_name.schema_name</pre> Response <pre><<SQL Command>> was successfully executed</pre> Error <pre>Error: <<SQL Command>> failed due to <<reason>></pre>

Table 46-2 (Cont.) List of Schema SQL Grammar

Operation	Grammar								
Alter Schema	Alter Schema Description ALTER SCHEMA <schema-name> set dbproperties (DESCRIPTION=<property- value>) Response <<SQL Command>> was successfully executed Error Error: <<SQL Command>> failed due to <<reason>>								
Delete Schema	DROP SCHEMA [IF EXISTS] <<schema_name>> By default during DROP schema, all child objects will also get deleted								
List Schemas	SHOW SCHEMAS [{ FROM IN } catalog_name] [[LIKE] regex_pattern] Examples: - SHOW SCHEMAS FROM defaultcatalog1 LIKE 'd*' - SHOW SCHEMAS IN defaultcatalog1 LIKE 'd*' Response: <table><tr><th></th><th>Schema</th></tr><tr><td>1</td><td><<schema_1>></td></tr><tr><td>2</td><td><<schema_2>></td></tr><tr><td>2</td><td><<schema_3>></td></tr></table> Error Error: <<SQL Command>> failed due to <<reason>>		Schema	1	<<schema_1>>	2	<<schema_2>>	2	<<schema_3>>
	Schema								
1	<<schema_1>>								
2	<<schema_2>>								
2	<<schema_3>>								

Table 46-2 List of Schema SQL Grammar

Operation	Grammar																		
Describe Schema (get details)	DESCRIBE SCHEMA <<catalog_name>>.<<schema_name>> DESCRIBE SCHEMA <<schema_name>> DESCRIBE SCHEMA <<schema_name>> in Catalog <<catalog_name>> <table><tr><th>Attribute</th><th>Value</th></tr><tr><td>Catalog name</td><td>Catalog name</td></tr><tr><td>Schema</td><td>Schema name</td></tr><tr><td>Description</td><td>Schema description</td></tr><tr><td>Created by</td><td>User that created the catalog</td></tr><tr><td>Created on</td><td>Date and time created</td></tr><tr><td>Updated by</td><td>User that last updated the catalog</td></tr><tr><td>Updated on</td><td>Date and time last updated</td></tr><tr><td>Location</td><td>Location in the catalog</td></tr></table>	Attribute	Value	Catalog name	Catalog name	Schema	Schema name	Description	Schema description	Created by	User that created the catalog	Created on	Date and time created	Updated by	User that last updated the catalog	Updated on	Date and time last updated	Location	Location in the catalog
Attribute	Value																		
Catalog name	Catalog name																		
Schema	Schema name																		
Description	Schema description																		
Created by	User that created the catalog																		
Created on	Date and time created																		
Updated by	User that last updated the catalog																		
Updated on	Date and time last updated																		
Location	Location in the catalog																		
Drop Schema	drop schema [IF EXISTS] <<schema_name>> cascade																		
Refresh Schema	REFRESH SCHEMA IN [EXTERNAL] CATALOG <<catalog_name.schema_name>> Response: Refresh Initiated																		

Volume SQL Grammar

Volume objects support the listed SQL grammar for DDL workloads.

Oracle AI Data Platform supports all standard Spark SQL data types. For more information, see [Apache Spark Documentation - Supported Data Types](#).

Table 46-3 Volume SQL Grammar

Operation	Grammar
Create volume	<pre>CREATE [EXTERNAL] VOLUME [IF NOT EXISTS] <<catalog_name.schema_name.volume_name>> [LOCATION location_path] [PROPERTIES (DESCRIPTION = description)]</pre>
Alter volume properties	<pre>ALTER VOLUME <<volume_name>> { RENAME TO <<new_volume_name>> [set properties (DESCRIPTION = description)] }</pre>
Drop volume	<pre>DROP VOLUME [IF EXISTS] <<volume_name>> OR DROP VOLUME <<catalog_name>>.<<schema_name>>.<<volume_name>> By default during DROP volume, all child objects will also get deleted</pre>
List Volumes	<pre>SHOW VOLUMES [{ FROM IN } catalog_name.schema_name] [[LIKE] regex_pattern }]</pre>

Table 46-3 Volume SQL Grammar

Operation	Grammar														
Describe Volume	DESCRIBE VOLUME volume_name <table><tr><th>Attribute</th><th>Value</th></tr><tr><td>Catalog name</td><td>Catalog name</td></tr><tr><td>Schema name</td><td>Schema name</td></tr><tr><td>Volume name</td><td>Volume name</td></tr><tr><td>Description</td><td>User defined description of volume</td></tr><tr><td>Location</td><td>Location in catalog</td></tr><tr><td>Volume type</td><td>Type of volume</td></tr></table> Error: Error: <<SQL Command>> failed due to <<reason>>	Attribute	Value	Catalog name	Catalog name	Schema name	Schema name	Volume name	Volume name	Description	User defined description of volume	Location	Location in catalog	Volume type	Type of volume
Attribute	Value														
Catalog name	Catalog name														
Schema name	Schema name														
Volume name	Volume name														
Description	User defined description of volume														
Location	Location in catalog														
Volume type	Type of volume														

Table SQL Grammar

Table objects support the listed SQL grammar for DDL workloads.

Oracle AI Data Platform supports all standard Spark SQL data types. For more information, see [Apache Spark Documentation - Supported Data Types](#).

Operation	Grammar
Create Table	<pre>CREATE [EXTERNAL] TABLE [IF NOT EXISTS] <catalog_name>.<schema-name>.<table-name> [(<column1-name><column1-type> [comment <column1-comment>], ...)] USING [HIVE DELTA, CSV, TXT, ORC, JDBC, PARQUET, etc.] [options (<key1>=<val1>[, ...])] [PARTITIONED BY (<par-column-name>[, ...])] [CLUSTERED BY (<clus-column-name>[, ...]) [SORTED BY (<sort-column-name> [asc desc][, ...])] INTO <num_buckets> buckets] [LOCATION '<path>'] [TBLPROPERTIES (DESCRIPTION = 'some-description', '<property-name>'='<property-value>'[, ...])]</pre> Response: <<SQL Command>> was successfully executed Error: Error: <<SQL Command>> failed due to <<reason>>
Create Managed Table	<pre>Create Managed Table CREATE TABLE <catalog>.<schema>.<table-name> [(<column1-name><column1-type> [comment <column1-comment>], ...)] USING <format>;</pre> Response: <<SQL Command>> was successfully executed Error: Error: <<SQL Command>> failed due to <<reason>>

Operation	Grammar
Create Managed Table with Data	<pre>create datatable <<catalog_name>>.<<schema_name>>.<<table_name>> [(<column1-name><column1-type> [comment <column1-comment>], ...)] tblproperties ('lakehouse_storage_format'='PARQUET') using parquet with select (<column1-name>], ...) from parquet.'oci://bucket@namespace/ folder/'</pre> Response: <<SQL Command>> was successfully executed Error: Error: <<SQL Command>> failed due to <<reason>>
Create Table with Uniform Support	<pre>CREATE [EXTERNAL] TABLE [IF NOT EXISTS] <catalog_name>.<schema-name>.<table-name> [(<column1-name> <column1-type> [comment <column1-comment>], ...)] [TBLPROPERTIES ('delta.universalFormat.enabledFormats' = 'iceberg')]</pre>

Operation	Grammar
Alter Table	ALTER TABLE table_old_name RENAME TO table_new_name ALTER TABLE table_name ADD COLUMNS (col_spec [, ...]) ALTER TABLE table_name DROP { COLUMN COLUMNS } [(] col_name [, ...] [)] ALTER TABLE table_name RENAME COLUMN col_name TO col_name ALTER TABLE table_name ADD [IF NOT EXISTS] (partition_spec [partition_spec ...]) ALTER TABLE table_name DROP [IF EXISTS] partition_spec [PURGE] ALTER TABLE table_name set tblproperties (description ='some-description')
Drop Table	DROP TABLE [IF EXISTS] table_name [PURGE] Response: <<SQL Command>> was successfully executed Error: Error: <<SQL Command>> failed due to <<reason>>
List Tables in a schema	SHOW TABLES in catalog_name.schema_name [LIKE <regex_pattern>] regex_pattern: A regular expression pattern that is used to filter the results of the statement. Response: <<namesake>>, tableName, isTemporary Error: <<SQL Command>> failed due to <<reason>>

Operation	Grammar
Describe Table	DESCRIBE TABLE [FORMAT] catalog_name.schema_name.table_name [PARTITION (<partition_col_name> = <partition_col_val>, ...)] [catalog_name.schema_name.table_name .column_name] Format: If EXTENDED is specified as the format, additional metadata information (such as parent database, owner, and access time) is returned. DESCRIBE TABLE catalog.schema.table Response: col_name,data_type,comment DESCRIBE TABLE catalog.schema.table column Response: info_name,info_value
Refresh Table	REFRESH TABLE IN EXTERNAL CATALOG <<catalog_name.schema_name.table_name >> Response: Refresh Initiated

DML Queries

You can run select, insert and delete queries on data using Oracle AI Data Platform notebooks, SL, Python, and Spark scripts.

AI Data Platform currently supports Spark 3.5 with Delta Lake 3.2.0. For more information on DML queries, see:

- [Apache Spark SQL Syntax - DML Statements](#)
- [Delta Lake - Table deletes, updates, and merges](#)
- [Delta Lake - Table utility commands](#)
- [Delta Lake - Use liquid clustering for Delta tables](#)

Aidputils for Agents and Tools Sample Code

The sample code provided is used demonstrate how you can use the aidputils library for building agents and tools.

For aidputils API reference, see [Aidputils API for Oracle AI Data Platform](#).

Topics:

- [Agent without Tools](#)
- [SQL Tool Test](#)
- [Prompt \(LLM\) Tool Test](#)
- [Custom Code Tool - Hello World](#)
- [Custom Code Tool - Developer Toolkit](#)
- [Agent with Tool Registration in Oracle AI Data Platform](#)
- [Observability: Logging, Tracing, and Metrics](#)
- [Agent Instantiation and Usage with Aidputil Packages](#)
- [Guardrails Configuration](#)

Agent without Tools

You can use the provided sample code to test an Oracle AI Data Platform AI agent that does not include tools, like prompt, SQL, or RAG.

```
# Generated code for SIMPLE_AGENT operator muse_agent_node
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.agent_helper import init_oci_llm,
pre_tool_setup, post_tool_setup, pre_invoke_setup
from aidputils.agents.toolkit.configs import AIDPToolConf, OCIAIConf,
ModelArgs
from langgraph.prebuilt import create_react_agent
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
import logging

logger = logging.getLogger('SingleAgentNoTool')
class_name = 'SingleAgentNoTool'
checkpointner = globals().get("checkpointner", None)

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",
    "policies" : [ ]
}
##### End Guardrails Configuration #####

##### Start Generated code for Agent Flow #####
##### Generated code for OCI Gen AI LLM
model_args = {
    "temperature" : 0.8,
    "max_tokens" : 500,
    "frequency_penalty" : 0,
    "presence_penalty" : 0,
    "top_p" : 1.0,
    "top_k" : 0
}
```

```

llm_conf = OCIAIConf(model_provider='cohere',
                     compartment_id='<your-compartment-ocid>',
                     model_args=model_args,
                     endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
                     model_id='<your-model-id>')

## Agent class definition
class SingleAgentNoTool:
    def __init__(self) -> None:
        self.agent = None
        """
        Setup for LangGraph agent. This includes returns react_agent or compiled
        langgraph object.
        """
    def setup(self) -> None:
        logger.info(llm_conf)
        # TODO: Handle other kinds of llms, for example openAI or gemini
        oci_llm = init_oci_llm(llm_conf)
        system_prompt = """
        You are an AI Agent
        """
        try:
            if checkpointer:
                self.agent = create_react_agent(model=oci_llm, tools=[],
                prompt=system_prompt, debug=True, checkpointer= checkpointer)
            else:
                self.agent = create_react_agent(model=oci_llm, tools=[],
                prompt=system_prompt, debug=True)
        except Exception as e:
            # Fallback compile without checkpointer if wiring fails
            self.agent = create_react_agent(model=oci_llm, tools=[],
            prompt=system_prompt, debug=True)
            logger.warning(f"Checkpointer could not be initialized {e}")
            logger.info(f"Setup for agent completed {self.agent}")

    async def invoke(self, user_query: str, **kwargs):
        token = pre_tool_setup(**kwargs)
        config = pre_invoke_setup(**kwargs)
        user_message = HumanMessage(content=user_query)
        message = {"messages": [dict(user_message)]}
        try:
            return await self.agent.ainvoke(input=message, config = config)
        except Exception as e:
            logger.error(f"Exception while calling invoke {e}")
        finally:
            post_tool_setup(token)

#####End Generated code for Agent Flow#####

```

SQL Tool Test

This sample code demonstrates how you can use aidputils to test your SQL tool.

```

from aidputils.agents.tools import utils
from aidputils.agents.auth.util import auth_utils

tool_conf = {'catalogKey': 'aidp_tools_dev',
              'schemaKey': 'aidpuser',
              'query': 'select * from employees where
SALARY>={{SALARY_RANGE}}'}
runtime_params = {"SALARY_RANGE": 60000}
context_vars = {'datalake_id': 'YOUR_DATALAKE_ID'}

try:
    tool_result = utils.call_tool_by_class('SQLTool', tool_conf,
runtime_params, **context_vars)
    print(tool_result)
except Exception as e:
    print(f"SQLTool execution failed: {e}")

```

Prompt (LLM) Tool Test

This sample code demonstrates how you can use aidputils to test your prompt tool.

```

from aidputils.agents.tools import utils
from aidputils.agents.auth.util import auth_utils

tool_conf = {
    'prompt_template': 'What is the capital of {{country}}',
    'llm': {
        'model_id': 'cohere.command-r-08-2024',
        'model_provider': 'cohere',
        'model_args': {
            'temperature': 1,
            'max_tokens': 600,
            'frequency_penalty': 0,
            'presence_penalty': 0,
            'top_k': 0,
            'top_p': 0.75
        },
        'compartment_id': '<your-compartment-ocid>',
        'auth_type': 'REMOTE',
        'endpoint': 'https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
        'auth_profile': 'DEFAULT'
    }
}
runtime_params = {'country': 'India'}
context_vars = {'datalake_id': 'YOUR_DATALAKE_ID'}

try:
    tool_result = utils.call_tool_by_class('PromptTool', tool_conf,

```

```
runtime_params, **context_vars)
    print(tool_result)
except Exception as e:
    print(f"PromptTool execution failed: {e}")
```

Custom Code Tool - Hello World

This sample code demonstrates how you can use aidputils to test your Custom Code tool.

The Hello World example is the simplest possible Custom Code tool. It defines a single tool class that accepts a name parameter and returns a greeting. Use it as a starting point for your own tool.

tool_implementation.py

```
from aidputils.agents.tools.custom_tools.base import CustomToolBase
```

```
@BaseTool.register
class HelloTool(CustomToolBase):
    """A simple greeting tool."""

    @classmethod
    def _execute_tool(cls, conf, runtime_params, **context_vars):
        name = runtime_params.get("name", "World")
        return {"greeting": f"Hello, {name}!"}
```

tool_config.json

```
{
  "displayName": "Hello Tool",
  "description": "A simple hello world tool",
  "tools": [
    {
      "toolClassName": "HelloTool",
      "displayName": "Hello Tool",
      "description": "Returns a hello world greeting",
      "version": "1.0.0",
      "schema": [
        {
          "name": "name",
          "type": "string",
          "description": "Name to greet"
        }
      ],
      "conf": {}
    }
  ]
}
```

requirements.txt

```
# no deps
```

Package the three files at the root of a ZIP archive and upload the ZIP through the Package tab. Once uploaded, switch to the Parameters tab, fill in the Description if you want to override the default, and switch to the Test tab to invoke the tool. With name="Alice", the tool returns:

```
{"greeting": "Hello, Alice!"}
```

Custom Code Tool - Developer Toolkit

This sample code demonstrates how you can use aidputils to test your Custom Code tool.

The Developer Toolkit example demonstrates a multi-tool package and the use of helper modules in a utils/ directory. The package registers three tools — a bash command runner, a file operations tool, and a Python code runner — and uses shared helper functions for output truncation and path sanitization.

Note

The Developer Toolkit is an illustrative example. Bash command execution and Python code execution have significant security implications. In production, restrict the AI compute, sandbox the operations, and apply strict allow-lists for the commands and code patterns the tool will execute.

Package Layout

```
advanced_tool.zip
├── tool_implementation.py
├── tool_config.json
├── requirements.txt          # stdlib only
├── utils/
│   ├── __init__.py
│   └── text_utils.py        # truncate_output, sanitize_path
```

tool_implementation.py

```
import subprocess
import os

from aidputils.agents.tools.custom_tools.base import CustomToolBase
from .utils.text_utils import truncate_output, sanitize_path

def _get_cfg(conf, key, default):
    """Read a config value from either the outer dict or the
    nested user conf. Coerces numeric settings to int to avoid
    type mismatches when values are rendered as strings by the
    template substitution layer."""
    inner = conf.get("conf") if isinstance(conf, dict) else None
    if isinstance(inner, dict) and key in inner:
        value = inner[key]
    elif isinstance(conf, dict) and key in conf:
        value = conf[key]
    else:
        value = default
```

```

        if isinstance(default, int) and not isinstance(value, bool):
            try:
                return int(value)
            except (TypeError, ValueError):
                return default
        return value

@BaseTool.register
class BashTool(CustomToolBase):
    """Execute bash commands and return output."""

    @classmethod
    def _execute_tool(cls, conf, runtime_params, **context_vars):
        command = runtime_params.get("command", "")
        timeout = _get_cfg(conf, "timeout", 30)
        max_lines = _get_cfg(conf, "max_output_lines", 200)
        try:
            result = subprocess.run(
                ["bash", "-c", command],
                capture_output=True, text=True, timeout=timeout
            )
        except subprocess.TimeoutExpired:
            # Surface the timeout as a tool failure rather than
            # returning {"error": ...}, which would be treated as
            # a successful response.
            raise RuntimeError(f"Command timed out after {timeout}s")
        output = result.stdout or ""
        if result.stderr:
            output += "\n[stderr]\n" + result.stderr
        return {"output": truncate_output(output, max_lines)}

@BaseTool.register
class FileTool(CustomToolBase):
    """Read, write, or list files in the workspace."""

    @classmethod
    def _execute_tool(cls, conf, runtime_params, **context_vars):
        operation = runtime_params.get("operation", "")
        path = runtime_params.get("path", "")
        content = runtime_params.get("content", "")
        base_dir = _get_cfg(conf, "base_dir", "/workspace")
        max_size = _get_cfg(conf, "max_file_size_kb", 1024) * 1024

        safe_path = sanitize_path(base_dir, path)
        if safe_path is None:
            raise ValueError("Invalid path: path traversal detected")

        if operation == "read":
            with open(safe_path, "r") as f:
                return {"output": f.read()}
        if operation == "write":
            parent = os.path.dirname(safe_path)
            if parent:
                os.makedirs(parent, exist_ok=True)

```

```

        with open(safe_path, "w") as f:
            f.write(content)
        return {"output": f"Written {len(content)} chars to {path}"}
    if operation == "list":
        target = safe_path if os.path.isdir(safe_path) else
os.path.dirname(safe_path)
        return {"output": "\n".join(sorted(os.listdir(target)))}
        raise ValueError(f"Unknown operation: {operation}. Use read/write/
list")

```

```

@BaseTool.register
class PythonTool(CustomToolBase):
    """Execute Python code in an isolated subprocess."""

    @classmethod
    def _execute_tool(cls, conf, runtime_params, **context_vars):
        code = runtime_params.get("code", "")
        timeout = _get_cfg(conf, "timeout", 60)
        max_lines = _get_cfg(conf, "max_output_lines", 500)
        try:
            result = subprocess.run(
                ["python3", "-c", code],
                capture_output=True, text=True, timeout=timeout
            )
        except subprocess.TimeoutExpired:
            raise RuntimeError(f"Execution timed out after {timeout}s")
        output = result.stdout or ""
        if result.stderr:
            output += "\n[stderr]\n" + result.stderr
        return {"output": truncate_output(output, max_lines)}

```

tool_config.json

```

{
  "displayName": "Developer Toolkit",
  "description": "A collection of tools for bash commands, file operations,
and Python execution",
  "tools": [
    {
      "toolClassName": "BashTool",
      "displayName": "Bash Tool",
      "description": "Executes a bash command and returns stdout/stderr
output",
      "version": "1.0.0",
      "schema": [
        {
          "name": "command",
          "type": "string",
          "description": "The bash command to execute"
        }
      ],
      "conf": {
        "timeout": 30,
        "max_output_lines": 200
      }
    }
  ]
}

```

```

    }
  },
  {
    "toolClassName": "FileTool",
    "displayName": "File Tool",
    "description": "Read, write, or list files in the workspace",
    "version": "1.0.0",
    "schema": [
      {"name": "operation", "type": "string",
        "description": "Operation to perform: read, write, or list"},
      {"name": "path", "type": "string",
        "description": "File or directory path"},
      {"name": "content", "type": "string",
        "description": "Content to write (for write operation)"}
    ],
    "conf": {
      "base_dir": "/workspace",
      "max_file_size_kb": 1024
    }
  },
  {
    "toolClassName": "PythonTool",
    "displayName": "Python Tool",
    "description": "Executes Python code in an isolated subprocess and
returns the output",
    "version": "1.0.0",
    "schema": [
      {"name": "code", "type": "string",
        "description": "The Python code to execute"}
    ],
    "conf": {
      "timeout": 60,
      "max_output_lines": 500
    }
  }
]
}

```

utils/text_utils.py

```

def truncate_output(text, max_lines=200):
    if not text:
        return ""
    try:
        max_lines = int(max_lines)
    except (TypeError, ValueError):
        max_lines = 200
    lines = text.strip().split("\n")
    if len(lines) > max_lines:
        lines = lines[:max_lines] + [f"... ({len(lines) - max_lines} lines
truncated)"]
    return "\n".join(lines)

```

```

def sanitize_path(base_dir, relative_path):

```

```

import os
if not relative_path:
    return base_dir
full = os.path.normpath(os.path.join(base_dir, relative_path))
if not full.startswith(os.path.normpath(base_dir)):
    return None
return full

```

utils/__init__.py

Empty file. Required for Python to treat utils/ as a package.

requirements.txt

stdlib only

After uploading the ZIP, the Package tab shows the three discovered tools and lets you enable or disable each one. The Parameters tab shows a **Tool Class** drop-down that switches between BashTool, FileTool, and PythonTool, and exposes the per-tool configuration (timeout, max_output_lines, base_dir, max_file_size_kb) on the right.

Agent with Tool Registration in Oracle AI Data Platform

Oracle AI Data Platform supports flexible agent construction and internal tool orchestration. This topic provides a sample recommended approach for defining, registering, and using tools within an agent.

1. Describe tools via configuration

Each tool is a Python dictionary:

```

my_tool = {
    "name": "blog_idea_tool",
    "description": "Generate blog ideas for a topic.",
    "class": "PromptTool",
    "conf": {...}, # tool-specific settings
    "params": [
        {"name": "topic", "type": "string", "description": "Blog topic"}
    ]
}

```

2. Register tools in a registry/config

All user tools are collected in a registry for agent lookup:

```

tool_conf = {
    "blog_idea_tool": my_tool,
    "social_post_tool": another_tool,
    # ... more tools
}

```

3. Framework wrapping: Create agent-consumable tool objects

Agent construction requires converting these dicts to executable tool objects (StructuredTool or similar):

```

from langchain_core.tools import StructuredTool

def create_langgraph_tool(tool):
    def tool_fn(**kwargs):
        # Example implementation: you would use utils.call_tool_by_name/tool
        runner, etc.
        return f"Executed {tool['name']} with inputs: {kwargs}"
    return StructuredTool.from_function(
        func=tool_fn,
        name=tool['name'],
        description=tool['description'],
        args_schema=None, # Build a pydantic schema if detailed validation
        required
        infer_schema=False
    )

```

4. Memory and Using a Checkpointer

Agents in AI Data Platform often need memory to persist intermediate state, enable resumability, and allow for recovery after failures or across long-running workflows. The typical mechanism is a checkpointer object, which saves and restores agent state.

```

# Suppose you have a 'checkpointer' object available:
# It might be provided to your agent context directly, or created via aidp-
agent-runtime utilities

```

```

# During agent run:
state = {"step": "tool_invoked", "result": tool_result}

```

```

if checkpointer:
    checkpointer.save(state)
    # To restore later:
    loaded_state = checkpointer.load()
    print(f"Restored state: {loaded_state}")

```

```

# You can persist any serializable agent context, params, or partial results

```

Usage pattern:

- Pass the 'checkpointer' to agent code/class at construction or as a global/context variable.
- Save state after every critical agent event, like tool output, prompt step, or LLM generation.
- Restore state on agent restart, if available.

Typical sources of the checkpointer:

- In AI Data Platform demo code, a `checkpointer` may be injected via workflow configuration or globals, e.g. `checkpointer = globals().get("checkpointer", None)`

- For complex use-cases, the checkpointer may wrap external storage, databases, or cloud state to allow robust failure recovery.

```
# Inside agent code
checkerpoint = globals().get("checkerpoint", None)
if checkerpoint:
    checkerpoint.save({"step": "after_tool", "context": context_vars})
    # ...
    restored_state = checkerpoint.load()
```

Observability: Logging, Tracing, and Metrics

Observability is seamlessly integrated into Oracle AI Data Platform applications via the `aidp_observability` package, enabling automatic telemetry (logs, traces, metrics) collection with minimal setup.

Initialization

Import and initialize as shown:

```
from observability.aidp_observability import AIDPObservability
from observability.config import CollectorConfig

config = CollectorConfig()
config.service_name = "dummy_name"
observability = AIDPObservability(config)
observability.initialize()
```

Upon initialization:

- OpenTelemetry exporters for traces, metrics, and logs are created.
- The collector endpoint is configured for all telemetry data (port 4317, GRpc protocol).
- Application loggers are set up.
- Playground mode enables in-memory exporter for instant trace display.
- The collector is pre-configured for log rotation, buffering, and includes a sink for telemetry exporting.
- Default metrics, logs, and AI Data Platform metadata are included in all telemetry signals.
- Default span/session attributes (e.g., `sessionId`, `traceId`) are set for correlation.

Usage pattern:

No changes are needed in application logic to emit telemetry. As a user:

- Use the OpenTelemetry Meter for metrics.
- Use Python's standard ``logging`` for logs.
- Use the OpenTelemetry Tracer for traces.

Example

```
import logging
import time
from opentelemetry import trace, metrics
```

```

tracer = trace.get_tracer(__name__)
meter = metrics.get_meter(__name__)

request_counter = meter.create_counter(
    name="requests_total",
    description="Number of requests processed",
    unit="1",
)

logging.basicConfig(level=logging.INFO)
logger = logging.getLogger("sample-app")

def process_request(user_id: str):
    logger.info("Processing request for user %s", user_id)
    request_counter.add(1, {"user.id": user_id})
    with tracer.start_as_current_span("process_request") as span:
        span.set_attribute("user.id", user_id)
        time.sleep(0.1)
        span.add_event("request_completed", {"status": "ok"})

if __name__ == "__main__":
    for i in range(3):
        process_request(f"user-{i}")
        time.sleep(1)

```

Note

Application telemetry is automatically exported; no instrumentation changes are required by the user. The observability package auto-instruments LLM frameworks and LangGraph applications for trace reporting.

Agent Instantiation and Usage with Aidputil Packages

The following samples demonstrate how you can create and use agents with aidputil packages.

```

from aidputils.agents.toolkit.agent_helper import invoke, get_client
from aidputils.agents.toolkit.configs import OCIAIConf
from langchain_core.tools import StructuredTool
from langgraph.prebuilt import create_react_agent
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
from langchain_community.chat_models.oci_generative_ai import ChatOCIGenAI
import logging
import json

logger = logging.getLogger('muse_agent_flow')
checkpointer = globals().get("checkpointer", None)

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",

```

```

    "policies" : [ ]
}
##### End Guardrails Configuration #####

##### Start Generated code for Agent Flow #####
##### Start Tool configuration for blog_idea_tool
##### Start PROMPT Tool configuration
blog_idea_tool_def = {
    "llm": {
        "model_id" : "<your-model-id>",
        "model_provider" : "cohere",
        "compartment_id" : "<your-compartment-ocid>",
        "endpoint" : "https://inference.generativeai.<oci-
region>.oci.oraclecloud.com",
        "auth_type" : "SECURITY_TOKEN",
        "auth_profile" : "DEFAULT",
        "model_args" : {
            "temperature" : 1,
            "max_tokens" : 600,
            "frequency_penalty" : 0,
            "presence_penalty" : 0,
            "top_k" : 0,
            "top_p" : 0.75
        }
    }, "prompt_template": ""
You are a master blog strategist.
Your task is to brainstorm compelling blog post ideas based on a given topic.
For the given topic '{{topic}}', generate 5 unique blog post titles.
For each title, include a one-sentence description of the angle the post
would take.
Present the output as a numbered list.
""
}

blog_idea_tool_params = [ {
    "name" : "topic",
    "type" : "string",
    "description" : "The central theme or subject for which to generate blog
ideas."
} ]

blog_idea_tool_dict = {
    "name": "blog_idea_tool",
    "description": "Use this tool to generate several distinct and engaging
blog post titles and concepts based on a topic ",
    "tool_class": "PromptTool",
    "conf": blog_idea_tool_def,
    "params": blog_idea_tool_params
}
blog_idea_tool = create_langgraph_tool(blog_idea_tool_dict)
##### End PROMPT Tool configuration
# Set tool_var_name = blog_idea_tool
# set ns.tool_var_list = [blog_idea_tool]
##### End Tool configuration for Blog idea tool
##### End Tool configuration

```

```

##### Start tool List#####
tools_agent1 = [blog_idea_tool]
##### End tool List#####

##### Generated code for OCI Gen AI LLM
model_args = {
    "temperature" : 0.8,
    "max_tokens" : 500,
    "frequency_penalty" : 0,
    "presence_penalty" : 0,
    "top_p" : 1.0,
    "top_k" : 0
}

llm_conf = OCIAIConf(model_provider='cohere',
                    compartment_id='<your-compartment-ocid>',
                    auth_type='SECURITY_TOKEN',
                    auth_profile='DEFAULT',
                    model_args=model_args,
                    endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
                    model_id='<your-model-id>')

## Agent class definition
class MuseAgentFlow:
    def __init__(self) -> None:
        self.agent = None

    def setup(self) -> None:
        # TODO: Handle other kinds of llms, for example openAI or gemini
        oci_llm = init_oci_llm(llm_conf)
        system_prompt = ""

    **Task:**
    For the given topic '{{topic}}', generate 5 unique blog post titles. For each
    title, include a one-sentence description of the angle the post would take.
    Present the output as a numbered list.

    **Example Input:**
    topic: "AI in marketing"

    **Example Output:**
    1.  **Title:** "Beyond the Hype: 3 Practical Ways to Use AI in Your Marketing
    Today"
        * **Angle:** This post will focus on simple, actionable AI tools that
        small businesses can implement immediately.
    2.  **Title:** "Is AI Coming for Your Marketing Job? A Realistic Look at the
    Future"
        * **Angle:** This post will explore how AI will change marketing roles,
        not just replace them, focusing on new skills.
    3.  **Title:** "We Let an AI Write Our Marketing Emails for a Week. Here's
    What Happened."
        * **Angle:** A case-study style post detailing the results of an
        interesting experiment.
    4.  **Title:** "The Ethics of AI Marketing: Are You Crossing a Line with

```

```

Personalization?"
    * **Angle:** A thought-leadership piece that discusses the important
ethical considerations of using AI.
5. * **Title:** "How to Personalize at Scale: A Guide to AI-Powered Customer
Journeys"
    * **Angle:** A tactical guide on using AI to create highly personalized
marketing campaigns.
"""

    try:
        if checkpointer:
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
prompt=system_prompt, debug=True, checkpointer= checkpointer)
        else:
            self.agent = self.agent = create_react_agent(model=oci_llm,
tools=tools_agent1, prompt=system_prompt, debug=True)
        except Exception as e:
            # Fallback compile without checkpointer if wiring fails
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
prompt=system_prompt, debug=True)
            logger.warning(f"Checkpointer could not be initialized {e}")
            logger.info(f"Setup for agent completed {self.agent}")

    async def invoke(self, user_query: str, **kwargs):
        try:
            return await self.agent.invoke(input=user_query, **kwargs)
        except Exception as e:
            logger.error(f"Exception while calling invoke {e}")

    def init_oci_llm(llm_conf: OCIAIConf):

        chat = ChatOCIGenAI(
            model_id='<your-model-id>',
            provider='cohere',
            service_endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
            compartment_id='<your-compartment-ocid>',
            client=get_client(llm_conf=llm_conf),
            model_kwargs=model_args
        )

        return chat

    def create_langgraph_tool(tool):
        def tool_fn(**kwargs):
            # Example implementation: you would use utils.call_tool_by_name/tool
runner, etc.
            return f"Executed {tool['name']} with inputs: {kwargs}"
        return StructuredTool.from_function(
            func=tool_fn,
            name=tool['name'],
            description=tool['description'],
            args_schema=None, # Build a pydantic schema if detailed validation
required
            infer_schema=False
        )

```

Guardrails Configuration

You can configure guardrails using aidputils as part of selecting a foundational model using `OCIAIConf()`.

Guardrails configuration is provided when selecting a foundational model from the OCI Generative AI service. In this example, we select the `xai.grok-4` model:

```
from aidputils.agents.toolkit.configs import OCIAIConf
guardrails_config = {
    "name" : "<guardrailsName>",
    "description" : "<guardrailsDescription>",
    "policies" : [ ]
}
model_args = {}
llm_conf = OCIAIConf(model_provider='generic',
                      compartment_id='<compartment_ocid>',
                      model_args=model_args,
                      endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
                      model_id='xai.grok-4',
                      guardrails_config=guardrails_config)
```

The guardrails config is a JSON-like string consisting of an array of policies. In the above example, it is defined in this block of code where `<guardrailsName>` and `<guardrailsDescription>` are a user-defined name and description:

```
guardrails_config = {
    "name" : "<guardrailsName>",
    "description" : "<guardrailsDescription>",
    "policies" : [ ]
}
```

Each policy has the following keys:

Key	Required	Description	Data Type	Default value
policyName	No	Custom name for the policy	String	N/A
policyType	Yes	Type of guardrail policy to apply. Allowed values include: - CONTENT_MODERATION - PROMPT_ATTACKS_PREVENTION - PII_DETECTION	ENUM	
policyDescription	No	A description for the policy	String	

Key	Required	Description	Data Type	Default value
scope	No	The scope defines where the guardrails are applied. Allowed values include: • USER_REQUEST • AGENT_RESPONSE • BOTH	ENUM	
action	No	The action to take when the policy is violated Allowed values include: • INFORM • BLOCK • ALLOW_MASK (only for PII_DETECTION)	ENUM	
threshold	No	Threshold for detection. Range is a probability between 0 and 1.	float	
piiCategories	Yes	Category of PII data to be detected along with their action and enablement.	Array	

piiCategories is also an array of JSON-like objects that uses the following keys:

Key	Required	Description	Data Type	Default value
category	Yes	The PII category to detect. Allowed values include: • PERSON • ADDRESS • TELEPHONE_NUMBER • EMAIL	String	N/A
isEnabled	No	Enable the detection of the PII category. Allowed values include: • True • False	ENUM	

Key	Required	Description	Data Type	Default value
action	No	Action to take if PII category is detected. Override the action above. Allowed values include: • INFORM • BLOCK • ALLOW • MASK	String	

Example: Complete Guardrails Configuration

In this case, we apply all three policies:

- content moderation is only applied on the agent response,
- prompt injection will block user requests if detected,
- PII is detected on both agent response and user request. Each PII category is treated differently.

```
guardrails_config = {
  "policies" : [ {
    "policyType" : "CONTENT_MODERATION",
    "policyName" : "Content Moderation prevention",
    "policyDescription" : "Choose an action to take when hate, sexual,
violence, toxic, derogatory, or harassment content is detected in either the
user input query or the agent response.",
    "scope" : "AGENT_RESPONSE",
    "action" : "INFORM",
    "threshold" : 0.5,
    "categories" : [ ]
  }, {
    "policyType" : "PROMPT_ATTACKS_PREVENTION",
    "policyName" : "Prompt Injection prevention",
    "policyDescription" : "Choose action when prompt injection is detected
on the user query.",
    "scope" : "USER_REQUEST",
    "action" : "BLOCK",
    "threshold" : 0.5
  }, {
    "policyType" : "PII_DETECTION",
    "policyName" : "Personally Identifiable Information (PII) detection",
    "policyDescription" : "Choose an action to take when PII entities are
detected in either the user input query or the agent response.",
    "scope" : "AGENT_RESPONSE",
    "action" : "INFORM",
    "threshold" : 0.5,
    "piiCategories" : [ {
      "category" : "PERSON",
      "isEnabled" : False,
      "action" : "INFORM"
    }, {
      "category" : "ADDRESS",
```

```

        "isEnabled" : False,
        "action" : "INFORM"
    }, {
        "category" : "TELEPHONE_NUMBER",
        "isEnabled" : True,
        "action" : "MASK"
    }, {
        "category" : "EMAIL",
        "isEnabled" : True,
        "action" : "MASK"
    } ]
    }, {
        "policyType" : "PII_DETECTION",
        "policyName" : "Personally Identifiable Information (PII) detection",
        "policyDescription" : "Choose an action to take when PII entities are
detected in either the user input query or the agent response.",
        "scope" : "USER_REQUEST",
        "action" : "INFORM",
        "threshold" : 0.5,
        "piiCategories" : [ {
            "category" : "PERSON",
            "isEnabled" : True,
            "action" : "INFORM"
        }, {
            "category" : "ADDRESS",
            "isEnabled" : True,
            "action" : "INFORM"
        }, {
            "category" : "TELEPHONE_NUMBER",
            "isEnabled" : True,
            "action" : "BLOCK"
        }, {
            "category" : "EMAIL",
            "isEnabled" : False,
            "action" : "INFORM"
        } ]
    } ]
}

```

Agent LangGraph Code Samples

You can use the provided LangGraph code samples as reference or to get started with agents in Oracle AI Data Platform.

Topics:

- [Hello World with LangGraph](#)
- [ReAct Agent with Prompt Tool LangGraph Sample Code](#)
- [ReAct Agent with RAG Tool LangGraph Sample Code](#)
- [ReAct Agent with SQL Tool LangGraph Sample Code](#)
- [ReAct Agent with Custom Tools LangGraph Sample Code](#)
- [Multi-agent System – Supervisor Pattern Sample Code](#)

- [Remote Oracle Analytics Cloud MCP Server Connection Sample Code](#)

Hello World with LangGraph

You can use this LangGraph code sample in an agent flow to test and debug output.

```
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
from langgraph.graph import StateGraph, MessagesState, START, END

def mock_llm(state: MessagesState):
    return {"messages": [{"role": "ai", "content": "hello world"}]}

class AgentBasic:
    def __init__(self) -> None:
        self.graph = None

    def setup(self) -> None:
        self.graph = StateGraph(MessagesState)
        self.graph.add_node(mock_llm)
        self.graph.add_edge(START, "mock_llm")
        self.graph.add_edge("mock_llm", END)
        self.graph = self.graph.compile()
        system_prompt = "Be a helpful assistant."

    async def invoke(self, user_query: str, **kwargs):
        user_message = HumanMessage(content=user_query)
        messages = {"messages": [dict(user_message)]}
        try:
            return self.graph.invoke(messages)
        except Exception as e:
            import traceback
            logger.error(f"Exception while calling invoke {e}", exc_info=True)
            print("Stack trace:\n", traceback.format_exc())

import asyncio
async def main():
    test_agent = AgentBasic()
    test_agent.setup()
    result = await test_agent.invoke("Hi there")
    print("Agent response:", result)
if __name__ == "__main__":
    asyncio.run(main())
```

ReAct Agent with Prompt Tool LangGraph Sample Code

You can use this LangGraph sample code to test and debug prompt tools in react agents.

```
from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.agent_helper import init_oci_llm,
pre_invoke_setup
from aidputils.agents.toolkit.configs import AIDPToolConf, OCIAIConf,
ModelArgs
from langgraph.prebuilt import create_react_agent
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
```

```

import logging
import json

logger = logging.getLogger('agent_with_prompt_tool')
checkpointer = globals().get("checkpointer", None)

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",
    "policies" : [ ]
}
##### End Guardrails Configuration #####

##### Start PROMPT Tool configuration
blogger_def = {
    "llm": {
        "model_id" : "xai.grok-4",
        "model_provider" : "generic",
        "compartment_id" : "<your-compartment-ocid>",
        "endpoint" : "https://inference.generativeai.<oci-
region>.oci.oraclecloud.com"
    }, "prompt_template": ""
    You are a master blog strategist. Your task is to brainstorm compelling blog
    post ideas based on a given topic.
    For the given {{topic}}, generate 5 unique blog post titles. For each title,
    include a one-sentence description
    of the angle the post would take. Present the output as a numbered list"
    ""
}

blogger_params = [ {
    "name" : "topic",
    "type" : "string",
    "description" : "Blog topic",
    "defaultValue" : "golf"
} ]

blogger_conf= AIDPToolConf(name="blogger",
                           description= "PROMPT_description_794368 ",
                           tool_class = "PromptTool",
                           conf=blogger_def, params=blogger_params)
blogger = create_langgraph_tool(blogger_conf.model_dump())
##### End PROMPT Tool configuration

##### Start tool List#####
tools_agent1 = [blogger]
##### End tool List#####

model_args = {}
llm_conf = OCIAIConf(model_provider='generic',
                     compartment_id='<your-oci-compartment-ocid>',
                     model_args=model_args,
                     endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',

```

```

        model_id='xai.grok-4',
        guardrails_config=guardrails_config)

## Agent class definition
class AgentWithPromptTool:
    def __init__(self) -> None:
        self.agent = None

    def setup(self) -> None:
        logger.info(llm_conf)
        # TODO: Handle other kinds of llms, for example openAI or gemini
        oci_llm = init_oci_llm(llm_conf)
        system_prompt = """
Be a helpful assistant.
"""

        try:
            if checkpointer:
                self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
prompt=system_prompt, debug=True, checkpointer= checkpointer)
            else:
                self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
prompt=system_prompt, debug=True)
        except Exception as e:
            # Fallback compile without checkpointer if wiring fails
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
prompt=system_prompt, debug=True)
            logger.warning(f"Checkpointer could not be initialized {e}")
            logger.info(f"Setup for agent completed {self.agent}")

    async def invoke(self, user_query: str, **kwargs):
        config = pre_invoke_setup(**kwargs)
        user_message = HumanMessage(content=user_query)
        message = {"messages": [dict(user_message)]}
        try:
            return await self.agent.ainvoke(input=message, config = config)
        except Exception as e:
            import traceback
            logger.error(f"Exception while calling invoke {e}", exc_info=True)
            print("Stack trace:\n", traceback.format_exc())

# The following is used when executing the python from within the agent code
editor
import asyncio
async def main():
    # Instantiate and initialize the agent
    test_agent = AgentWithPromptTool()
    test_agent.setup()

    # You can customize this user query or prompt for input
    user_query = "Give me 3 ideas for a robotics blog"

    # Run the asynchronous invoke method and print the result
    result = await test_agent.invoke(user_query)

```

```

    print("Agent response:", result)

if __name__ == "__main__":
    asyncio.run(main())

```

ReAct Agent with RAG Tool LangGraph Sample Code

You can use this LangGraph sample code to test and debug RAG tools in react agents.

```

from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.agent_helper import init_oci_llm,
pre_invoke_setup
from aidputils.agents.toolkit.configs import AIDPToolConf, OCIAIConf,
ModelArgs
from langgraph.prebuilt import create_react_agent
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
import logging
import json

logger = logging.getLogger('agent_with_rag_tool')
checkpointer = globals().get("checkpointer", None)

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",
    "policies" : [ ]
}
##### End Guardrails Configuration #####

##### Start RAG Tool configuration

rag_params = [ {
    "name" : "query",
    "type" : "string",
    "description" : "RAG query",
    "defaultValue" : "find matching doc artifacts for ..."
} ]

conf = {
    "catalog": "default",
    "schema": "default",
    "knowledgeBase": "acme_kb",
    "top_k": 5,
    "llm": {
        "model_id" : "xai.grok-4",
        "model_provider" : "generic",
        "compartment_id" : "<your-compartment-ocid>",
        "endpoint" : "https://inference.generativeai.<oci-
region>.oci.oraclecloud.com"
    }
}
rag_conf= AIDPToolConf(name="rag",
                        description= "RAG Tool ",
                        tool_class = "RAGTool", conf=conf,

```

```

params=rag_params)
rag_tool = create_langgraph_tool(rag_conf.model_dump())

##### End RAG Tool configuration

##### Start tool List#####
tools_agent1 = [rag_tool]
##### End tool List#####

model_args = {}

llm_conf = OCIAIConf(model_provider='generic',
                    compartment_id='<your-compartment-ocid>',
                    model_args=model_args,
                    endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
                    model_id='xai.grok-4',
                    guardrails_config=guardrails_config)

## Agent class definition
class AgentWithRAGTool:
    def __init__(self) -> None:
        self.agent = None

    def setup(self) -> None:
        logger.info(llm_conf)
        oci_llm = init_oci_llm(llm_conf)
        system_prompt = """
Be a helpful assistant.
"""

        try:
            if checkpointer:
                self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
                prompt=system_prompt, debug=True, checkpointer= checkpointer)
            else:
                self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
                prompt=system_prompt, debug=True)
        except Exception as e:
            # Fallback compile without checkpointer if wiring fails
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
            prompt=system_prompt, debug=True)
            logger.warning(f"Checkpointer could not be initialized {e}")
            logger.info(f"Setup for agent completed {self.agent}")

    async def invoke(self, user_query: str, **kwargs):
        config = pre_invoke_setup(**kwargs)
        user_message = HumanMessage(content=user_query)
        message = {"messages": [dict(user_message)]}
        try:
            return await self.agent.ainvoke(input=message, config = config)
        except Exception as e:
            import traceback
            logger.error(f"Exception while calling invoke {e}", exc_info=True)

```

```

        print("Stack trace:\n", traceback.format_exc())

import asyncio
async def main():
    # Instantiate and initialize the agent
    test_agent = AgentWithRAGTool()
    test_agent.setup()

    # You can customize this user query or prompt for input
    user_query = "Summarize SOX Compliance and Reconciliation at Acme Corp"

    # Run the asynchronous invoke method and print the result
    result = await test_agent.invoke(user_query)
    print("Agent response:", result)

if __name__ == "__main__":
    asyncio.run(main())

```

ReAct Agent with SQL Tool LangGraph Sample Code

You can use this LangGraph sample code to test and debug SQL tools in ReAct agents.

```

from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.agent_helper import init_oci_llm,
pre_invoke_setup
from aidputils.agents.toolkit.configs import AIDPToolConf, OCIAIConf,
ModelArgs
from langgraph.prebuilt import create_react_agent
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
import logging
import json
import os

compartment_id = os.getenv('AIDP_USER_COMPARTMENT_ID')
compartment_id="<your-compartment-ocid>"
logger = logging.getLogger('agent_with_sql_tool')
checkpointer = globals().get("checkpoint", None)

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",
    "policies" : [ ]
}
##### End Guardrails Configuration #####

#####
##### Start SQL Tool configuration

sql_params = [ {
    "name" : "MAX_SALARY",
    "type" : "string",

```

```

        "description" : "Maximum salary",
        "defaultValue" : "50000"
    } ]

conf = {
    "catalogKey": "adw23ai_phx",
    "schemaKey": "gold",
    "query": """ Select * from (
        Select 101 employee_id, 'John' first_name, 'Doe' last_name,
        'john.doe@acme.com' email_address, 75000 salary from DUAL
        UNION
        Select 102 employee_id, 'Jane' first_name, 'Smith' last_name,
        'jane.smith@acme.com' email_address, 100000 salary from DUAL
        UNION
        Select 103 employee_id, 'Peter' first_name, 'Jones' last_name,
        'peter.jones@acme.com' email_address, 45000 salary from DUAL
    ) employees where salary >= {{MAX_SALARY}}
    """
}

sql_conf= AIDPToolConf(name="query_employees",
                        description= "Query employees using SQL
Tool ",
                        tool_class = "SQLTool", conf=conf,
                        params=sql_params)
sql_tool = create_langgraph_tool(sql_conf.model_dump())

##### End SQL Tool configuration

##### Start tool List#####
tools_agent1 = [sql_tool]
##### End tool List#####

model_args = {}

llm_conf = OCIAIConf(model_provider='generic',
                     compartment_id='<your-compartment-ocid>',
                     model_args=model_args,
                     endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
                     model_id='xai.grok-4',
                     guardrails_config=guardrails_config)

## Agent class definition
class AgentWithSQLTool:
    def __init__(self) -> None:
        self.agent = None

    def setup(self) -> None:
        logger.info(llm_conf)
        oci_llm = init_oci_llm(llm_conf)
        system_prompt = """
        Be a helpful assistant.

```

```

"""

    try:
        if checkpointer:
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
            prompt=system_prompt, debug=True, checkpointer= checkpointer)
        else:
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
            prompt=system_prompt, debug=True)
    except Exception as e:
        # Fallback compile without checkpointer if wiring fails
        self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
        prompt=system_prompt, debug=True)
        logger.warning(f"Checkpointer could not be initialized {e}")
        logger.info(f"Setup for agent completed {self.agent}")

    async def invoke(self, user_query: str, **kwargs):
        config = pre_invoke_setup(**kwargs)
        user_message = HumanMessage(content=user_query)
        message = {"messages": [dict(user_message)]}
        try:
            return await self.agent.ainvoke(input=message, config = config)
        except Exception as e:
            import traceback
            logger.error(f"Exception while calling invoke {e}", exc_info=True)
            print("Stack trace:\n", traceback.format_exc())

import asyncio
async def main():
    # Instantiate and initialize the agent
    test_agent = AgentWithSQLTool()
    test_agent.setup()

    # You can customize this user query or prompt for input
    user_query = "Which employees have a salary greater than 50000"

    # Run the asynchronous invoke method and print the result
    result = await test_agent.invoke(user_query)
    print("Agent response:", result)

if __name__ == "__main__":
    asyncio.run(main())

```

ReAct Agent with Custom Tools LangGraph Sample Code

You can use this LangGraph code sample in an agent flow to test and debug output for custom tools.

```

from aidputils.agents.toolkit.tool_helper import create_langgraph_tool
from aidputils.agents.toolkit.agent_helper import init_oci_llm,
pre_invoke_setup
from aidputils.agents.toolkit.configs import AIDPToolConf, OCIAIConf,

```

```

ModelArgs
from langgraph.prebuilt import create_react_agent
from langchain_core.tools import tool
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
import logging
import json

logger = logging.getLogger('agent_with_prompt_tool')
checkpointeer = globals().get("checkpointer", None)

SYSTEM_PROMPT = (
    "You are a customer operations assistant.\n"
    "Use tools to gather accurate information.\n"
    "Think step by step.\n"
    "Respond clearly and concisely.\n"
)

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",
    "policies" : [ ]
}
##### End Guardrails Configuration #####

@tool
def lookup_customer(customer_id: str) -> str:
    """Lookup customer profile (stub)"""
    return f"Customer {customer_id}: Enterprise, ARR $250k, healthy"

@tool
def fetch_usage(customer_id: str) -> str:
    """Fetch customer usage metrics (stub)"""
    return f"Customer {customer_id}: 42 jobs/day, 3.1TB processed"

@tool
def open_ticket(reason: str) -> str:
    """Create support ticket (stub)"""
    return f"Ticket created for issue: {reason}"

TOOLS = [
    lookup_customer,
    fetch_usage,
    open_ticket
]

model_args = {}
llm_conf = OCIAIConf(model_provider='generic',
                     compartment_id='<your-compartment-ocid>',
                     model_args=model_args,

```

```

        endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
        model_id='xai.grok-4',
        guardrails_config=guardrails_config)

## Agent class definition
class AgentWithTools:
    def __init__(self) -> None:
        self.agent = None

    def setup(self) -> None:
        logger.info(llm_conf)
        # TODO: Handle other kinds of llms, for example openAI or gemini
        oci_llm = init_oci_llm(llm_conf)

        try:
            if checkpointer:
                self.agent = create_react_agent(model=oci_llm, tools=TOOLS,
prompt=SYSTEM_PROMPT, debug=True, checkpointer= checkpointer)
            else:
                self.agent = create_react_agent(model=oci_llm, tools=TOOLS,
prompt=SYSTEM_PROMPT, debug=True)
        except Exception as e:
            # Fallback compile without checkpointer if wiring fails
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent1,
prompt=system_prompt, debug=True)
            logger.warning(f"Checkpointer could not be initialized {e}")
            logger.info(f"Setup for agent completed {self.agent}")

    async def invoke(self, user_query: str, **kwargs):
        config = pre_invoke_setup(**kwargs)
        user_message = HumanMessage(content=user_query)
        message = {"messages": [dict(user_message)]}
        try:
            return await self.agent.ainvoke(input=message, config = config)
        except Exception as e:
            import traceback
            logger.error(f"Exception while calling invoke {e}", exc_info=True)
            print("Stack trace:\n", traceback.format_exc())

# The following is used when executing the python from within the agent code
editor
import asyncio
async def main():
    # Instantiate and initialize the agent
    test_agent = AgentWithTools()
    test_agent.setup()

    # You can customize this user query or prompt for input
    user_query = "Get customer 123 profile and usage"

    # Run the asynchronous invoke method and print the result
    result = await test_agent.invoke(user_query)
    print("Agent response:", result)

```

```
if __name__ == "__main__":
    asyncio.run(main())
```

Multi-agent System – Supervisor Pattern Sample Code

This example follows a supervisor agent pattern and is a data agent which uses SQLite. The SQL functions can be replaced with SQL tool and leverage Oracle database. For example, this uses SQLite for demonstration as it works out of the box.

Try:

- What was the total revenue in EMEA?
- How did the revenue in EMEA compare APAC in 2024?

The design of this agent flow looks like this:

- Router
 - → SQL agent
 - * → Comparison agent
 - * → Insight agent
 - * → Final
 - * → Final
 - → Other agent → Final

```
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
from langgraph.graph import StateGraph, MessagesState, START, END
import sqlite3
from typing import TypedDict, Literal
from aidutils.agents.toolkit.agent_helper import init_oci_llm,
pre_invoke_setup
from aidutils.agents.toolkit.configs import AIDPToolConf, OCIAIConf,
ModelArgs
from typing import TypedDict, Literal
from typing import List
from langchain_core.messages import BaseMessage

## Replace compartment id and endpoint
##
compartment_id = '<your-compartment-ocid>'
endpoint = 'https://inference.generativeai.<oci-region>.oci.oraclecloud.com'
####

checkpointer = globals().get("checkpointer", None)
conn = sqlite3.connect("file::memory:?cache=shared", uri=True)

def setup_db():
    cur = conn.cursor()

    cur.execute("""
        CREATE TABLE IF NOT EXISTS orders (
            order_id INTEGER,
            region TEXT,
```

```

        revenue REAL,
        order_date TEXT
    )
    """

cur.executemany(
    "INSERT INTO orders VALUES (?, ?, ?, ?)",
    [
        (1, "EMEA", 1200, "2024-10-01"),
        (2, "EMEA", 900, "2024-10-02"),
        (3, "AMER", 1500, "2024-10-01"),
        (4, "EMEA", 400, "2024-10-03"),
        (5, "APAC", 800, "2024-10-02"),
        (6, "EMEA", 1400, "2025-10-01"),
        (7, "AMER", 1200, "2025-10-01"),
        (8, "EMEA", 900, "2025-10-03"),
        (9, "APAC", 300, "2025-10-02"),
    ],
)

conn.commit()
#conn.close()

class State(TypedDict):
    question: str
    route: Literal["sql", "other"]
    sql: str
    rows: list
    content: str
    comparison: str
    insight: str
    messages: List[BaseMessage]

model_args = {}
guardrails_config = {
    "name": "Default Guardrails",
    "description": "Default empty guardrails configuration",
    "policies": [ ]
}

llm_conf = OCIAIConf(model_provider='generic',
                     compartment_id='<your-compartment-ocid>',
                     model_args=model_args,
                     endpoint=endpoint,
                     model_id='xai.grok-4',
                     guardrails_config=guardrails_config)
llm = init_oci_llm(llm_conf)

def supervisor(state: State) -> State:
    messages = [
        HumanMessage(
            content=(
                "Decide whether the following question requires SQL
analysis.\n\n"
                "Respond with ONLY one word:\n"
                "- sql\n"
            )
        )
    ]

```

```

        "- other\n\n"
        f"Question:\n{state['question']}"
    )
]

response: AIMessage = llm.invoke(messages)
route = response.content.strip().lower()

return {"route": route}

def other_agent(state: State) -> State:
    response: AIMessage = llm.invoke(
        [HumanMessage(content=state["question"])]
    )
    return {"content": response.content.strip()}

def final(state: State) -> State:
    messages = state.get("messages", [])

    combined_answer = state["content"]
    print(state.get("insight"))
    if state.get("insight") and state["insight"] != "No additional insight.":
        combined_answer += f"\n\nInsight: {state['insight']}"

    messages.append(HumanMessage(content=state["question"]))
    messages.append(AIMessage(content=combined_answer))

    return {
        **state,
        "messages": messages,
    }

def execute_sql(query: str):
    conn = sqlite3.connect("file::memory:?cache=shared", uri=True)
    cur = conn.cursor()
    cur.execute(query)

    columns = [desc[0] for desc in cur.description]
    rows = cur.fetchall()

    conn.close()

    return columns, rows

def sql_agent(state: State) -> State:
    # 1. Generate SQL
    sql_messages = [
        HumanMessage(
            content=(
                "You are a senior data analyst.\n\n"

```

```

        "Database schema:\n"
        "orders(order_id, region, revenue, order_date)\n\n"
        "Write a SQLite-compatible SQL query that contents the
question below.\n"
        "Return ONLY the SQL starting with the SELECT statement.\n\n"
        f"Question:\n{state['question']}"
    )
    )
]

sql_response: AIMessage = llm.invoke(sql_messages)
sql = sql_response.content.strip()

# 2. Execute SQL
columns, rows = execute_sql(sql)
results = [dict(zip(columns, row)) for row in rows]

# 3. Format content (LLM owns presentation)
format_messages = [
    HumanMessage(
        content=(
            "You are a professional analytics assistant.\n\n"
            "The following data is the FINAL, correct result.\n\n"
            f"Data:\n{results}\n\n"
            f"User Question:\n{state['question']}\n\n"
            "Formatting Rules:\n"
            "- Format entire response using Markdown syntax. Include
headers, bold text, and a bulleted list\n"
            "- Start with a short headline (max 12 words)\n"
            "- On the next line, give a complete sentence contenting the
question\n"
            "- Clearly state the numeric value\n"
            "- Use commas in numbers\n"
            "- Do NOT mention SQL, databases, tables, or queries\n"
            "- Do NOT explain how the data was obtained\n"
            "- Do NOT add disclaimers\n\n"
            "Output Format (EXACT):\n"
            "<Headline>\n"
            "<Sentence>"
        )
    )
]

formatted_response: AIMessage = llm.invoke(format_messages)
content = formatted_response.content.strip()

return {
    "sql": sql,
    "rows": results,
    "content": content,
}

def comparison_agent(state: State) -> State:
    rows = state.get("rows", [])

```

```

        messages = [
            HumanMessage(
                content=(
                    "You are a business analyst.\n\n"
                    "Analyze the result data and produce a comparative
insight.\n\n"
                    f"Result Data:\n{rows}\n\n"
                    "Rules:\n"
                    "- Format entire response using Markdown syntax. Include
headers, bold text, and a bulleted list\n"
                    "- If multiple rows exist, identify the highest, lowest, or
notable difference\n"
                    "- If only one value exists, explain what it represents and
how it could be compared\n"
                    "- Do NOT mention SQL or databases\n"
                    "- One concise sentence\n"
                    "- Never say 'no additional insight'\n"
                )
            )
        ]

        response: AIMessage = llm.invoke(messages)

        return {
            "comparison": response.content.strip()
        }

def insight_agent(state: State) -> State:
    messages = [
        HumanMessage(
            content=(
                "You are a senior analytics advisor.\n\n"
                f"User Question:\n{state['question']}\n\n"
                f"Comparison Insight:\n{state.get('comparison', '')}\n\n"
                "Turn this into a clear business insight.\n"
                "- One sentence\n"
                "- No speculation\n"
                "- No technical language\n"
            )
        )
    ]

    response: AIMessage = llm.invoke(messages)

    return {
        "insight": response.content.strip()
    }

class AgentBasic:
    def __init__(self) -> None:
        self.graph = None

    def setup(self) -> None:

```

```

setup_db()
builder = StateGraph(State)

builder.add_node("supervisor", supervisor)
builder.add_node("sql_agent", sql_agent)
builder.add_node("comparison_agent", comparison_agent)
builder.add_node("insight_agent", insight_agent)
builder.add_node("other_agent", other_agent)
builder.add_node("final", final)

builder.set_entry_point("supervisor")

builder.add_conditional_edges(
    "supervisor",
    lambda s: s["route"],
    {
        "sql": "sql_agent",
        "other": "other_agent",
    },
)

builder.add_edge("sql_agent", "comparison_agent")
builder.add_edge("comparison_agent", "insight_agent")
builder.add_edge("insight_agent", "final")
builder.add_edge("other_agent", "final")
builder.add_edge("final", END)

if checkpointer:
    self.graph = builder.compile(checkpointer= checkpointer)
else:
    self.graph = builder.compile()

async def invoke(self, user_query: str, **kwargs):
    config = pre_invoke_setup(**kwargs)
    initial_state = {
        "question": user_query
    }
    try:
        return await self.graph.ainvoke(initial_state, config=config)
    except Exception as e:
        import traceback
        #logger.error(f"Exception while calling invoke {e}", exc_info=True)
        print("Stack trace:\n", traceback.format_exc())

import asyncio

async def main():
    test_agent = AgentBasic()
    test_agent.setup()
    result = await test_agent.invoke("What was the total revenue in EMEA?")
    print("Agent response:", result)
if __name__ == "__main__":
    asyncio.run(main())

```

Remote Oracle Analytics Cloud MCP Server Connection Sample Code

You can use this sample code as a guideline for setting up your own remote MCP server connections to Oracle Analytics Cloud.

① Note

This code is dependent on you having an existing Oracle Analytics Cloud MCP server running and has been included for illustration.

langchain-mcp-adapters

Entry File

```
ACCESS_TOKEN = "enter_your_key"

from aidputils.agents.toolkit.agent_helper import init_oci_llm,
pre_invoke_setup
from aidputils.agents.toolkit.configs import OCIAIConf, ModelArgs
from aidp_flowutils.configs import AIDPToolConf, OCIAIConf, ModelArgs
from langgraph.prebuilt import create_react_agent
from langchain_core.messages import AIMessage, HumanMessage, SystemMessage
from langchain_core.messages import BaseMessage, AIMessage, HumanMessage,
SystemMessage
from aidp_flowutils.agent_helper import pre_tool_setup, post_tool_setup,
parse_stream_response
from typing import AsyncGenerator, Dict, Union
import logging
from langchain_mcp_adapters.client import MultiServerMCPClient
import uuid
import os
import asyncio
from concurrent.futures import ThreadPoolExecutor

logger = logging.getLogger('test')
checkpointer = globals().get("checkpointer", None)
_executor = ThreadPoolExecutor(1)

def async_to_sync(awaitable):
    loop = asyncio.new_event_loop()
    return _executor.submit(loop.run_until_complete, awaitable).result()

##### Guardrails Configuration #####
guardrails_config = {
    "name" : "Default Guardrails",
    "description" : "Default empty guardrails configuration",
    "policies" : [ ]
}
##### End Guardrails Configuration #####
##### Start tool List#####

MCP_URL = "https://<your-oac-instance-url>/api/mcp"
```

```

client = MultiServerMCPCClient({ "oac": { "transport":"streamable_http",
"url":MCP_URL, "headers":{ "Authorization": f"Bearer {ACCESS_TOKEN}",
"Content-Type": "application/json" }}})

tools_agent = async_to_sync(client.get_tools())
##### End tool List#####

model_args = {}

llm_conf = OCIAIConf(model_provider='generic',
                     compartment_id='<your-compartment-ocid>',
                     model_args=model_args,
                     endpoint='https://inference.generativeai.<oci-
region>.oci.oraclecloud.com',
                     model_id='xai.grok-code-fast-1',
                     guardrails_config=guardrails_config)

## Agent class definition
class Test:
    def __init__(self) -> None:
        self.agent = None
        """
        Setup for LangGraph agent.
        """
    def setup(self) -> None:
        logger.info(llm_conf)
        oci_llm = init_oci_llm(llm_conf)
        system_prompt = """
        Be a helpful and useful assistant. Use the Oracle Analytics MCP Tools
        to help answer data related information.
        """
        try:
            mem_url = os.getenv("MEMORY_SERVER_URL") or os.getenv("MEMORY_URL") or
"http://127.0.0.1:21100"
            from oracle_memory_clients.client import ProxyStoreClient,
AsyncProxyCheckpointClient # type: ignore
            checkpointer = AsyncProxyCheckpointClient(base_url=mem_url,
agent="basic demo agent")

            if checkpointer:
                self.agent = create_react_agent(model=oci_llm, tools=tools_agent,
prompt=system_prompt, debug=True, checkpointer= checkpointer)
            else:
                self.agent = create_react_agent(model=oci_llm, tools=tools_agent,
prompt=system_prompt, debug=True)
        except Exception as e:
            # Fallback compile without checkpointer if wiring fails
            self.agent = create_react_agent(model=oci_llm, tools=tools_agent,
prompt=system_prompt, debug=True)
            logger.warning(f"Checkpointer could not be initialized {e}")
            logger.info(f"Setup for agent completed {self.agent}")

        async def invoke(self, user_query: str, **kwargs) -> Union[Dict,
AsyncGenerator[BaseMessage, None]]:
            config = pre_invoke_setup(**kwargs)

```

```

        user_message = HumanMessage(content=user_query)
        message = {"messages": [dict(user_message)]}
        is_stream = bool(kwargs.get("stream", False))
        if is_stream:
            return self._stream_messages(input=message, config=config,
kwargs=kwargs)
        else:
            token = pre_tool_setup(**kwargs)
            try:
                agent_response = await self.agent.ainvoke(input=message, config =
config)
                final_response = {**agent_response, "messages":
agent_response.get("messages", [])[1:]}
                return final_response
            except Exception as e:
                logger.error(f"Exception while calling invoke {e}")
                raise

    async def _stream_messages(self, input, config, kwargs) ->
AsyncGenerator[BaseMessage, None]:
        """
        Stream messages as they are generated by the agent.
        Yields BaseMessage objects as new messages are added by nodes.
        Ensures the auth context (ContextVar) remains active during streaming and
is cleaned up after.
        """
        logger.info("Starting _stream_messages")
        token = pre_tool_setup(**kwargs)

        try:
            # Determine streaming mode based on availability of StreamingData class
            try:
                from aidp_auth.client.generative_ai_inference_v2_client import
StreamingData
                stream = self.agent.astream(input=input, config=config,
stream_mode="messages")
            except ImportError:
                stream = self.agent.astream(input=input, config=config)

            async for chunk in parse_stream_response(stream):
                yield chunk
        except Exception:
            logger.exception("Streaming error")
            raise

import asyncio
async def main():
    # Instantiate and initialize the agent
    demo_agent = Test()
    demo_agent.setup()
    # You can customize this user query or prompt for input
    user_query = "What datasets are available in my analytics instance?"
    #user_query = "Give me a summary of the recent Ice Cream Sales data,
providing highlights and lowlights in the summary."

    # Run the asynchronous invoke method and print the result

```

```
        result = await demo_agent.invoke(user_query, thread_id=uuid.uuid4().hex)
        print("Agent response:", result)

if __name__ == "__main__":
    asyncio.run(main())
```

Known Issues

This page lists current limitations, bugs, and behavior inconsistencies identified in Oracle AI Data Platform.

These issues are actively being tracked and will be addressed in future updates.

When reporting any issues regarding AI Data Platform resources (e.g. Workspace, Compute), you should also share the ID of the resource. You can find the resource ID by going to the page where the resource is listed, clicking ... **Actions** and clicking **Copy ID**.

Table 47-1 Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
1	AI Data Platform Instance	UPDATE	Updating tags on the AI Data Platform instance does not retroactively apply to sub resources that were already created in the same compartment by AI Data Platform.	Tag updates only apply to newly created resources in the customer tenancy after the tag change. If consistent tag values are required across all resources, manually update tags on existing sub resources.
2	AI Data Platform Instance	UPDATE	When a user selects or creates an external table using a bucket that has already been tagged/associated with another AI Data Platform instance the service blocks reuse. Deleting the external table from the original instance does not allow reusing the same bucket in the new AI Data Platform instance.	No current workaround. Once a bucket is chosen and the external table is created, the association is fixed. Customers must provision a new bucket if they need to create a new external table with a different AI Data Platform instance and copy the data to this new location.

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
3	AI Data Platform Instance	CREATE	A single bucket location cannot be referenced by multiple AI Data Platform instances for external tables at the same time. Attempting to reuse the same bucket results in validation failure.	No current workaround. Each AI Data Platform instance requires its own dedicated bucket for external tables. Customers must create separate buckets per instance for creating new external tables in a new AI Data Platform instance.
4	AI Data Platform Instance	UPDATE	Cannot edit tag namespaces after instance creation.	No current workaround
5	Auto-populate	CREATE	Auto-populate creation fails if the bucket has already been tagged with another AI Data Platform instance.	• If the tagged AI Data Platform is still present and active, the bucket can't be used in the new AI Data Platform instance's extraction. This is to prevent the two AI Data Platform instances from overwriting each other's data. • If the tagged AI Data Platform is deleted, the bucket can be reused for the new AI Data Platform. The governing AIDP Id tag needs to be removed first, then you can retry auto-populate creation.

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
6	Compute	UPDATE	Updating a cluster is currently only possible while it is active. Updating a stopped cluster is currently not possible.	To update a stopped cluster, you can either start the cluster and update it or delete and create that cluster again with the changes.
7	Workspace	MOVE	From a workspace page, moving multiple items by selecting multiple items at one go is currently not possible.	Users can move one item at a time by selecting that item first, then selecting Move from the Actions Menu.
8	External Catalog	REFRESH	When a user creates an external catalog, it runs a background job to harvest the metadata from the external source. The background job takes a long time for the metadata from the external system to reflect in the external catalog. Clicking on the refresh icon in the master catalog while the first job is still in progress causes problems in the processing of the background job.	You should not click the Refresh button until the job has finished. Users can see the progress of the background job in the History tab. When the job is complete, you can click Refresh if all the expected objects have not appeared in the external catalog.
9	External Catalog	REFRESH	When a user creates an external catalog, it runs a background job to harvest the metadata from the external source. If a column in a table in that external source has "\$" in the name it will not be harvested.	No current workaround

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
10	External Catalog	CREATE	Users cannot create external catalogs with regional wallet files.	You can create external catalogs using instance wallets or provide the instance details and create an external catalog.
11	External Catalog	LIST	Users cannot list or use synonyms in an external catalog when using an AI Lakehouse source.	No current workaround
12	Data Ingestion	CREATE	Users cannot write to an Oracle Autonomous AI Lakehouse table using a notebook without first creating an external catalog	Users are expected in general to first create an external Autonomous AI Lakehouse catalog and then write it to an Autonomous AI Lakehouse table.
13	External/Managed Table	CREATE	Users cannot create external tables when the input data is in multi-line JSON format.	No current workaround
14	Data Sharing	CONSUMPTION	Users cannot receive a share using Delta Sharing protocol.	Users can load the shared data into a data frame using the delta-sharing library's <code>load_as_spark(<<table_path>>)</code> .
15	Data Sharing	CONSUMPTION	Users cannot consume data shared by AI Data Platform in Autonomous AI Lakehouse	No current workaround
16	Auto Populate	CREATE	The Auto Populate feature does not support table creation when the target location contains delimited data files with a delimiter other than a comma.	No current workaround

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
17	Workflow	CREATE/ SCHEDULE	Directly accessing OCI Object Storage (or GenAI service) is currently not supported for scheduled jobs	Configure external volumes for Object Storage access in the workflow job. There is no current workaround to access generative AI models
18	Workspace configured for Private Network	Access AWS S3 Buckets	For workspaces configured for Private Network, currently there is a known issue about AWS S3 access from Notebooks and Workflows.	File a support ticket if you need to access AWS S3 buckets from notebooks and workflows for Workspaces configured for private networks. Please indicate the workspace key in that ticket. AI Data Platform team will enable a temporary workaround for your tenancy while we work on a permanent solution.
19	Compute	Monitor compute notebooks	All notebooks attached to a compute cluster as displayed are Active regardless of actual status	No current workaround AI Data Platform team is working on a fix. Notebooks will display the correct status when the fix is deployed.
20	AI Data Platform Instance	Network source and Admin recovery	Network source based access control does not currently cover the admin recovery related actions.	No current workaround AI Data Platform team is working on a fix. Until the fix is released, you can temporarily disable network source based access control, if you need to perform the infrequent actions related to admin recovery.

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
21	AI Data Platform Instance	Delta tables	Delta table writes and reads hit OCI Object Storage throttling errors if there are too many small files in the Delta tables.	No current workaround AI Data Platform team is working on automatic capabilities related to compaction. Until the fix is released, we recommend using the Optimize command frequently to avoid throttling-related errors.
22	Compute	Edit cluster	Changing shapes (changing compute architectures) while editing Spark clusters is currently not supported.	You can create a new cluster if you need a different compute shape. We are evaluating use-cases where changing shapes are critical. Please file a support request describing your use-case.
23	Data lineage	Lineage capture	Data lineage is not currently generated for Apache Spark Structured Streaming tasks run from notebooks or Workflow jobs. As a result, the lineage graph might not display relationships between source and target datasets.	No current workaround
24	AI Data Platform	Browser support	AI Data Platform is not supported on the Safari browser. Users may experience unsupported behavior or be unable to access some console functionality when using Safari.	Use a supported browser, such as Google Chrome or Mozilla Firefox.

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
25	Compute	Iceberg table	Iceberg features might not be available on a cluster that has not been restarted after Iceberg support was enabled.	Restart the cluster before using Iceberg features.
26	Iceberg table	Procedure	The <code>remove_orphan_files</code> procedure is not supported.	
27	Iceberg table	Procedure	The <code>rewrite_table_path</code> procedure is not supported.	
28	Iceberg table	Procedure	The location argument of the snapshot procedure is not supported.	
29	Iceberg table	Query	Queries against the snapshots, history, <code>metadata_log_entries</code> , manifests, files, <code>all_data_files</code> , partitions, and refs metadata tables require a fully qualified four-part name when the Iceberg table is in the default catalog and schema. For example, <code>SELECT * FROM default.<table_name>.snapshots</code> is not supported.	Use the complete three-part names: <code><catalog>.<schema>.<table>.<metadata_table></code> . For example: <code>SELECT * FROM default.default.<table_name>.snapshots</code> .
30	Iceberg table	Migrate	The <code>drop_backup</code> argument of the <code>migrate</code> procedure is not supported.	Run the <code>migrate</code> procedure without specifying the <code>drop_backup</code> argument.

Table 47-1 (Cont.) Known Issues

	Resource	Category	Error and How to Find the Problem	Workaround
31	Iceberg table	Migrate	Migrating external tables to Iceberg is not supported because AI Data Platform does not permit overlapping table storage locations.	No current workaround.
32	Iceberg table	Migrate	Renaming an existing Iceberg table is not supported.	Restart the cluster and recreate the table with the required name.