



Informatica® Big Data Management
10.2.2 Service Pack 1

User Guide

© Copyright Informatica LLC 2012, 2019

This software and documentation are provided only under a separate license agreement containing restrictions on use and disclosure. No part of this document may be reproduced or transmitted in any form, by any means (electronic, photocopying, recording or otherwise) without prior consent of Informatica LLC.

U.S. GOVERNMENT RIGHTS Programs, software, databases, and related documentation and technical data delivered to U.S. Government customers are "commercial computer software" or "commercial technical data" pursuant to the applicable Federal Acquisition Regulation and agency-specific supplemental regulations. As such, the use, duplication, disclosure, modification, and adaptation is subject to the restrictions and license terms set forth in the applicable Government contract, and, to the extent applicable by the terms of the Government contract, the additional rights set forth in FAR 52.227-19, Commercial Computer Software License.

Informatica, the Informatica logo, Big Data Management, and PowerExchange are trademarks or registered trademarks of Informatica LLC in the United States and many jurisdictions throughout the world. A current list of Informatica trademarks is available on the web at <https://www.informatica.com/trademarks.html>. Other company and product names may be trade names or trademarks of their respective owners.

Portions of this software and/or documentation are subject to copyright held by third parties. Required third party notices are included with the product.

The information in this documentation is subject to change without notice. If you find any problems in this documentation, report them to us at infa_documentation@informatica.com.

Informatica products are warranted according to the terms and conditions of the agreements under which they are provided. INFORMATICA PROVIDES THE INFORMATION IN THIS DOCUMENT "AS IS" WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED, INCLUDING WITHOUT ANY WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND ANY WARRANTY OR CONDITION OF NON-INFRINGEMENT.

Publication Date: 2019-05-09

Table of Contents

Preface	13
Informatica Resources.	13
Informatica Network.	13
Informatica Knowledge Base.	13
Informatica Documentation.	13
Informatica Product Availability Matrices.	14
Informatica Velocity.	14
Informatica Marketplace.	14
Informatica Global Customer Support.	14
 Chapter 1: Introduction to Informatica Big Data Management.....	 15
Informatica Big Data Management Overview.	15
Example.	15
Big Data Management Component Architecture.	16
Clients and Tools.	16
Application Services.	16
Repositories.	17
Hadoop Integration.	17
Databricks Integration.	19
Big Data Management Engines.	19
Run-time Process on the Blaze Engine.	20
Run-time Process on the Spark Engine.	21
Run-time Process on the Databricks Spark Engine.	22
Big Data Process.	23
Step 1. Collect the Data.	23
Step 2. Cleanse the Data.	24
Step 3. Transform the Data.	24
Step 4. Process the Data.	24
Step 5. Monitor Jobs.	24
Data Warehouse Optimization Mapping Example	25
 Chapter 2: Mappings.....	 27
Overview of Mappings.	27
Mapping Run-time Properties.	28
Validation Environments.	28
Execution Environment.	30
Updating Run-time Properties for Multiple Mappings.	33
Sqoop Mappings in a Hadoop Environment.	34
Sqoop Mapping-Level Arguments.	34
Incremental Data Extraction for Sqoop Mappings.	37

Configuring Sqoop Properties in the Mapping.	38
Configuring Parameters for Sqoop Arguments in the Mapping.	38
Rules and Guidelines for Mappings in a Non-native Environment.	39
Rules and Guidelines for Mappings on the Blaze Engine.	39
Rules and Guidelines for Mappings on the Spark Engine.	39
Rules and Guidelines for Mappings on the Databricks Spark Engine.	40
Workflows that Run Mappings in a Non-native Environment.	40
Configuring a Mapping to Run in a Non-native Environment.	41
Mapping Execution Plans.	42
Blaze Engine Execution Plan Details.	42
Spark Engine Execution Plan Details.	44
Databricks Spark Engine Execution Details.	44
Viewing the Execution Plan.	45
Optimization for the Hadoop Environment.	45
Blaze Engine High Availability.	47
Enabling Data Compression on Temporary Staging Tables.	47
Truncating Partitions in a Hive Target.	48
Scheduling, Queuing, and Node Labeling.	49
Big Data Job Recovery.	51
Spark Engine Optimization for Sqoop Pass-Through Mappings.	51
Troubleshooting Mappings in a Non-native Environment.	52
Mappings in the Native Environment.	53
Data Processor Mappings.	53
HDFS Mappings.	53
Hive Mappings.	55
Social Media Mappings.	55
Native Environment Optimization.	56
Chapter 3: Sources.	59
Overview of Sources.	59
PowerExchange Adapter Sources.	59
Sources on Databricks.	60
Complex File Sources on ADLS.	61
Complex File Sources on Azure Blob.	61
Rules and Guidelines for Databricks Sources.	61
File Sources on Hadoop.	62
Complex File Sources on Amazon S3.	62
Complex File Sources on ADLS.	63
Complex File Sources on Azure Blob.	64
Complex File Sources on MapR-FS.	65
Complex File Sources on HDFS.	65
Flat File Sources on Hadoop.	65
Relational Sources on Hadoop.	66

Hive Sources on Hadoop.	67
PreSQL and PostSQL Commands.	67
Rules and Guidelines for Hive Sources on the Blaze Engine.	68
Sqoop Sources on Hadoop.	70
Reading Data from Vertica Sources through Sqoop.	70
Rules and Guidelines for Sqoop Sources.	71
Rules and Guidelines for Sqoop Queries.	71
Chapter 4: Targets.	72
Overview of Targets.	72
PowerExchange Adapter Targets.	72
Targets on Databricks.	73
Complex File Targets on ADLS.	74
Complex File Targets on Azure Blob	74
Rules and Guidelines for Databricks Targets.	74
File Targets on Hadoop.	75
Complex File Targets on Amazon S3.	76
Complex File Targets on ADLS.	76
Complex File Targets on Azure Blob	77
Complex File Targets on MapR-FS.	78
Complex File Targets on HDFS	78
Flat File Targets on Hadoop.	79
Message Targets on Hadoop.	79
Relational Targets on Hadoop.	80
Hive Targets on Hadoop.	80
PreSQL and PostSQL Commands.	81
Truncating Hive Targets.	81
Updating Hive Targets with an Update Strategy Transformation.	82
Rules and Guidelines for Hive Targets on the Blaze Engine.	82
Sqoop Targets on Hadoop.	84
Rules and Guidelines for Sqoop Targets.	84
Chapter 5: Transformations.	85
Overview of Transformations.	85
Address Validator Transformation in a Non-native Environment.	87
Address Validator Transformation on the Blaze Engine.	87
Address Validator Transformation on the Spark Engine.	88
Aggregator Transformation in a Non-native Environment.	88
Aggregator Transformation on the Blaze Engine.	88
Aggregator Transformation on the Spark Engine.	89
Aggregator Transformation on the Databricks Spark Engine.	89
Case Converter Transformation in a Non-native Environment.	90
Classifier Transformation in a Non-native Environment.	90

Comparison Transformation in a Non-native Environment.	90
Consolidation Transformation in a Non-native Environment.	91
Consolidation Transformation on the Blaze Engine.	91
Consolidation Transformation on the Spark Engine.	91
Data Masking Transformation in a Non-native Environment.	91
Data Masking Transformation on the Blaze Engine.	91
Data Masking Transformation on the Spark Engine.	92
Data Processor Transformation in a Non-native Environment.	93
Data Processor Transformation on the Blaze Engine.	93
Decision Transformation in a Non-native Environment.	94
Decision Transformation on the Spark Engine.	94
Expression Transformation in a Non-native Environment.	94
Expression Transformation on the Blaze Engine.	94
Expression Transformation on the Spark Engine.	94
Expression Transformation on the Databricks Spark Engine.	95
Filter Transformation in a Non-native Environment.	95
Filter Transformation on the Blaze Engine.	95
Java Transformation in a Non-native Environment.	95
Java Transformation on the Blaze Engine.	96
Java Transformation on the Spark Engine.	96
Joiner Transformation in a Non-native Environment	97
Joiner Transformation on the Blaze Engine.	97
Joiner Transformation on the Spark Engine.	98
Joiner Transformation on the Databricks Spark Engine.	98
Key Generator Transformation in a Non-native Environment.	98
Labeler Transformation in a Non-native Environment.	99
Lookup Transformation in a Non-native Environment.	99
Lookup Transformation on the Blaze Engine.	99
Lookup Transformation on the Spark Engine.	99
Lookup Transformation on the Databricks Spark Engine.	101
Match Transformation in a Non-native Environment.	101
Match Transformation on the Blaze Engine.	101
Match Transformation on the Spark Engine.	102
Merge Transformation in a Non-native Environment.	102
Normalizer Transformation in a Non-native Environment.	102
Parser Transformation in a Non-native Environment.	102
Python Transformation in a Non-native Environment.	103
Python Transformation on the Spark Engine.	103
Rank Transformation in a Non-native Environment.	103
Rank Transformation on the Blaze Engine	103
Rank Transformation on the Spark Engine.	104
Rank Transformation on the Databricks Spark Engine.	104

Router Transformation in a Non-native Environment.	105
Sequence Generator Transformation in a Non-native Environment.	105
Sequence Generator Transformation on the Blaze Engine.	105
Sequence Generator Transformation on the Spark Engine.	105
Sorter Transformation in a Non-native Environment.	105
Sorter Transformation on the Blaze Engine.	106
Sorter Transformation on the Spark Engine.	106
Sorter Transformation on the Databricks Spark Engine.	107
Standardizer Transformation in a Non-native Environment.	108
Union Transformation in a Non-native Environment.	108
Union Transformation in a Streaming Mapping.	108
Update Strategy Transformation in a Non-native Environment.	108
Update Strategy Transformation on the Blaze Engine.	109
Update Strategy Transformation on the Spark Engine.	109
Weighted Average Transformation in a Non-native Environment.	111
Chapter 6: Cluster Workflows.	112
Cluster Workflows Overview.	112
Cluster Workflow Components.	113
Cluster Workflows Process	114
Create Cluster Task Properties.	115
Advanced Properties for Azure HDInsight.	116
Advanced Properties for Amazon EMR.	117
Advanced Properties for the Blaze Engine.	119
Advanced Properties for a Hive Metastore Database.	121
Advanced Properties for Databricks.	121
Mapping Task Properties.	123
Add a Delete Cluster Task.	123
Deploy and Run the Workflow.	124
Monitoring Azure HDInsight Cluster Workflow Jobs	124
Chapter 7: Profiles.	125
Profiles Overview.	125
Native Environment.	125
Hadoop Environment.	126
Column Profiles for Sqoop Data Sources.	126
Creating a Single Data Object Profile in Informatica Developer.	127
Creating an Enterprise Discovery Profile in Informatica Developer.	127
Creating a Column Profile in Informatica Analyst.	129
Creating an Enterprise Discovery Profile in Informatica Analyst.	130
Creating a Scorecard in Informatica Analyst.	131
Monitoring a Profile.	132
Profiling Functionality Support.	133

Troubleshooting.	133
Chapter 8: Monitoring.....	135
Overview of Monitoring.	135
Hadoop Environment Logs.	135
YARN Web User Interface.	136
Accessing the Monitoring URL.	137
Viewing Hadoop Environment Logs in the Administrator Tool.	137
Monitoring a Mapping.	138
Blaze Engine Monitoring.	141
Blaze Job Monitoring Application.	143
Blaze Summary Report.	144
Blaze Engine Logs.	147
Viewing Blaze Logs.	148
Orchestrator Sunset Time.	148
Troubleshooting Blaze Monitoring.	148
Spark Engine Monitoring.	149
Viewing Hive Tasks.	153
Spark Engine Logs.	153
Viewing Spark Logs	154
Troubleshooting Spark Engine Monitoring.	154
Chapter 9: Hierarchical Data Processing.....	155
Overview of Hierarchical Data Processing.	155
How to Develop a Mapping to Process Hierarchical Data.	156
Complex Data Types.	158
Array Data Type.	159
Map Data Type.	161
Struct Data Type.	161
Rules and Guidelines for Complex Data Types.	162
Complex Ports.	163
Complex Ports in Transformations.	164
Rules and Guidelines for Complex Ports.	164
Creating a Complex Port.	165
Complex Data Type Definitions.	165
Nested Data Type Definitions.	167
Rules and Guidelines for Complex Data Type Definitions.	167
Creating a Complex Data Type Definition.	167
Importing a Complex Data Type Definition.	168
Type Configuration.	170
Changing the Type Configuration for an Array Port.	170
Changing the Type Configuration for a Map Port.	172
Specifying the Type Configuration for a Struct Port.	173

Complex Operators.	174
Extracting an Array Element Using a Subscript Operator.	175
Extracting a Struct Element Using the Dot Operator.	175
Complex Functions.	176
Hierarchical Data Preview.	177
Hierarchical Data Preview Process.	178
Data Preview Interface.	178
Previewing Hierarchical Data.	180
Rules and Guidelines for Hierarchical Data Preview.	180
Chapter 10: Hierarchical Data Processing Configuration.....	181
Hierarchical Data Conversion.	181
Convert Relational or Hierarchical Data to Struct Data.	182
Creating a Struct Port.	182
Convert Relational or Hierarchical Data to Nested Struct Data.	184
Creating A Nested Complex Port.	185
Extract Elements from Hierarchical Data.	192
Extracting Elements from a Complex Port.	192
Flatten Hierarchical Data.	194
Flattening a Complex Port.	195
Chapter 11: Hierarchical Data Processing with Schema Changes.....	197
Overview of Hierarchical Data Processing with Schema Changes.	197
How to Develop a Dynamic Mapping to Process Schema Changes in Hierarchical Data.	198
Dynamic Complex Ports.	199
Dynamic Ports and Dynamic Complex Ports.	200
Dynamic Complex Ports in Transformations.	201
Input Rules for a Dynamic Complex Port.	201
Input Rule for a Dynamic Array.	202
Input Rules for a Dynamic Map.	203
Input Rules for a Dynamic Struct.	203
Port Selectors for Dynamic Complex Ports.	204
Dynamic Expressions.	204
Example - Dynamic Expression to Construct a Dynamic Struct.	205
Complex Operators.	206
Complex Functions.	206
Rules and Guidelines for Dynamic Complex Ports.	206
Optimized Mappings.	206
Chapter 12: Intelligent Structure Models.....	209
Overview of Intelligent Structure Models.	209
Intelligent Structure Discovery Process.	210
Use Case.	210

Using an Intelligent Structure Model in a Mapping.	211
Rules and Guidelines for Intelligent Structure Models.	211
Before You Begin.	212
How to Develop and Run a Mapping to Process Data with an Intelligent Structure Model	212
Mapping Example.	214
Creating an Informatica Intelligent Cloud Services Account.	215
Creating an Intelligent Structure Model.	215
Exporting an Intelligent Structure Model.	216
Checking for Data Loss.	216
Chapter 13: Stateful Computing.....	217
Overview of Stateful Computing.	217
Windowing Configuration.	218
Frame.	218
Partition and Order Keys.	219
Rules and Guidelines for Windowing Configuration.	221
Window Functions.	221
LEAD.	222
LAG.	222
Aggregate Functions as Window Functions.	222
Rules and Guidelines for Window Functions.	226
Windowing Examples.	226
Financial Plans Example.	226
GPS Pings Example.	228
Aggregate Function as Window Function Example.	230
Appendix A: Connections.....	233
Connections.	234
Cloud Provisioning Configuration.	234
AWS Cloud Provisioning Configuration Properties.	235
Azure Cloud Provisioning Configuration Properties.	236
Databricks Cloud Provisioning Configuration Properties.	239
Amazon Redshift Connection Properties.	240
Amazon S3 Connection Properties.	241
Cassandra Connection Properties.	243
Databricks Connection Properties.	244
Google Analytics Connection Properties.	246
Google BigQuery Connection Properties.	246
Google Cloud Spanner Connection Properties.	247
Google Cloud Storage Connection Properties.	248
Hadoop Connection Properties.	249
Hadoop Cluster Properties.	249
Common Properties.	251

Reject Directory Properties.	252
Hive Pushdown Configuration.	252
Blaze Configuration.	252
Spark Configuration.	253
HDFS Connection Properties.	254
HBase Connection Properties.	256
HBase Connection Properties for MapR-DB.	256
Hive Connection Properties.	257
JDBC Connection Properties.	260
JDBC Connection String.	262
Sqoop Connection-Level Arguments.	264
Kafka Connection Properties.	267
Microsoft Azure Blob Storage Connection Properties.	267
Microsoft Azure Cosmos DB SQL API Connection Properties.	269
Microsoft Azure Data Lake Store Connection Properties.	269
Microsoft Azure SQL Data Warehouse Connection Properties.	270
Snowflake Connection Properties.	272
Creating a Connection to Access Sources or Targets.	273
Creating a Hadoop Connection.	273
Configuring Hadoop Connection Properties.	274
Cluster Environment Variables.	275
Cluster Library Path.	275
Common Advanced Properties.	275
Blaze Engine Advanced Properties.	276
Spark Advanced Properties.	277
Appendix B: Data Type Reference.	280
Data Type Reference Overview.	280
Transformation Data Type Support in a Non-native Environment.	281
Complex File and Transformation Data Types.	282
Avro and Transformation Data Types.	282
JSON and Transformation Data Types.	283
Parquet and Transformation Data Types.	283
Hive Data Types and Transformation Data Types.	285
Hive Complex Data Types.	286
Sqoop Data Types.	287
Aurora Data Types.	287
IBM DB2 and DB2 for z/OS Data Types.	287
Greenplum Data Types.	288
Microsoft SQL Server Data Types.	288
Netezza Data Types.	289
Oracle Data Types.	289
Teradata Data Types.	290

Teradata Data Types with TDCH Specialized Connectors for Sqoop.	290
Vertica Data Types.	291
Appendix C: Function Reference.	293
Function Support in a Non-native Environment.	293
Function and Data Type Processing.	295
Rules and Guidelines for Spark Engine Processing.	295
Index.	298

Preface

The *Informatica Big Data Management® User Guide* provides information about configuring and running mappings in the native and non-native environments.

Informatica Resources

Informatica provides you with a range of product resources through the Informatica Network and other online portals. Use the resources to get the most from your Informatica products and solutions and to learn from other Informatica users and subject matter experts.

Informatica Network

The Informatica Network is the gateway to many resources, including the Informatica Knowledge Base and Informatica Global Customer Support. To enter the Informatica Network, visit <https://network.informatica.com>.

As an Informatica Network member, you have the following options:

- Search the Knowledge Base for product resources.
- View product availability information.
- Create and review your support cases.
- Find your local Informatica User Group Network and collaborate with your peers.

Informatica Knowledge Base

Use the Informatica Knowledge Base to find product resources such as how-to articles, best practices, video tutorials, and answers to frequently asked questions.

To search the Knowledge Base, visit <https://search.informatica.com>. If you have questions, comments, or ideas about the Knowledge Base, contact the Informatica Knowledge Base team at KB_Feedback@informatica.com.

Informatica Documentation

Use the Informatica Documentation Portal to explore an extensive library of documentation for current and recent product releases. To explore the Documentation Portal, visit <https://docs.informatica.com>.

Informatica maintains documentation for many products on the Informatica Knowledge Base in addition to the Documentation Portal. If you cannot find documentation for your product or product version on the Documentation Portal, search the Knowledge Base at <https://search.informatica.com>.

If you have questions, comments, or ideas about the product documentation, contact the Informatica Documentation team at infa_documentation@informatica.com.

Informatica Product Availability Matrices

Product Availability Matrices (PAMs) indicate the versions of the operating systems, databases, and types of data sources and targets that a product release supports. You can browse the Informatica PAMs at <https://network.informatica.com/community/informatica-network/product-availability-matrices>.

Informatica Velocity

Informatica Velocity is a collection of tips and best practices developed by Informatica Professional Services and based on real-world experiences from hundreds of data management projects. Informatica Velocity represents the collective knowledge of Informatica consultants who work with organizations around the world to plan, develop, deploy, and maintain successful data management solutions.

You can find Informatica Velocity resources at <http://velocity.informatica.com>. If you have questions, comments, or ideas about Informatica Velocity, contact Informatica Professional Services at ips@informatica.com.

Informatica Marketplace

The Informatica Marketplace is a forum where you can find solutions that extend and enhance your Informatica implementations. Leverage any of the hundreds of solutions from Informatica developers and partners on the Marketplace to improve your productivity and speed up time to implementation on your projects. You can find the Informatica Marketplace at <https://marketplace.informatica.com>.

Informatica Global Customer Support

You can contact a Global Support Center by telephone or through the Informatica Network.

To find your local Informatica Global Customer Support telephone number, visit the Informatica website at the following link:

<https://www.informatica.com/services-and-training/customer-success-services/contact-us.html>.

To find online support resources on the Informatica Network, visit <https://network.informatica.com> and select the eSupport option.

CHAPTER 1

Introduction to Informatica Big Data Management

This chapter includes the following topics:

- [Informatica Big Data Management Overview, 15](#)
- [Big Data Management Component Architecture, 16](#)
- [Big Data Management Engines, 19](#)
- [Big Data Process, 23](#)
- [Data Warehouse Optimization Mapping Example , 25](#)

Informatica Big Data Management Overview

Informatica Big Data Management enables your organization to process large, diverse, and fast changing data sets so you can get insights into your data. Use Big Data Management to perform big data integration and transformation without writing or maintaining external code.

Use Big Data Management to collect diverse data faster, build business logic in a visual environment, and eliminate hand-coding to get insights on your data. Consider implementing a big data project in the following situations:

- The volume of the data that you want to process is greater than 10 terabytes.
- You need to analyze or capture data changes in microseconds.
- The data sources are varied and range from unstructured text to social media data.

You can perform run-time processing in the native environment or in a non-native environment. The native environment is the Informatica domain where the Data Integration Service performs all run-time processing. Use the native run-time environment to process data that is less than 10 terabytes. A non-native environment is a distributed cluster outside of the Informatica domain, such as Hadoop or Databricks, where the Data Integration Service can push run-time processing. Use a non-native run-time environment to optimize mapping performance and process data that is greater than 10 terabytes.

Example

You are an investment banker who needs to calculate the popularity and risk of stocks and then match stocks to each customer based on the preferences of the customer. Your CIO wants to automate the process of calculating the popularity and risk of each stock, match stocks to each customer, and then send an email with a list of stock recommendations for all customers.

You consider the following requirements for your project:

- The volume of data generated by each stock is greater than 10 terabytes.
- You need to analyze the changes to the stock in microseconds.
- The stock is included in Twitter feeds and company stock trade websites, so you need to analyze these social media sources.

Based on your requirements, you work with the IT department to create mappings to determine the popularity of a stock. One mapping tracks the number of times the stock is included in Twitter feeds, and another mapping tracks the number of times customers inquire about the stock on the company stock trade website.

Big Data Management Component Architecture

The Big Data Management components include client tools, application services, repositories, and third-party tools that Big Data Management uses for a big data project. The specific components involved depend on the task you perform.

Clients and Tools

Based on your product license, you can use multiple Informatica tools and clients to manage big data projects.

Use the following tools to manage big data projects:

Informatica Administrator

Monitor the status of profile, mapping, and MDM Big Data Relationship Management jobs on the Monitoring tab of the Administrator tool. The Monitoring tab of the Administrator tool is called the Monitoring tool. You can also design a Vibe Data Stream workflow in the Administrator tool.

Informatica Analyst

Create and run profiles on big data sources, and create mapping specifications to collaborate on projects and define business logic that populates a big data target with data.

Informatica Developer

Create and run profiles against big data sources, and run mappings and workflows on the Hadoop cluster from the Developer tool.

Application Services

Big Data Management uses application services in the Informatica domain to process data.

Use the Administrator tool to create connections, monitor jobs, and manage application services that Big Data Management uses.

Big Data Management uses the following application services:

Analyst Service

The Analyst Service runs the Analyst tool in the Informatica domain. The Analyst Service manages the connections between service components and the users that have access to the Analyst tool.

Data Integration Service

The Data Integration Service can process mappings in the native environment or push the mapping for processing to a compute cluster in a non-native environment. The Data Integration Service also retrieves metadata from the Model repository when you run a Developer tool mapping or workflow. The Analyst tool and Developer tool connect to the Data Integration Service to run profile jobs and store profile results in the profiling warehouse.

Mass Ingestion Service

The Mass Ingestion Service manages and validates mass ingestion specifications that you create in the Mass Ingestion tool. The Mass Ingestion Service deploys specifications to the Data Integration Service. When a specification runs, the Mass Ingestion Service generates ingestion statistics.

Metadata Access Service

The Metadata Access Service allows the Developer tool to import and preview metadata from a Hadoop cluster.

The Metadata Access Service contains information about the Service Principal Name (SPN) and keytab information if the Hadoop cluster uses Kerberos authentication. You can create one or more Metadata Access Services on a node. Based on your license, the Metadata Access Service can be highly available.

HBase, HDFS, Hive, and MapR-DB connections use the Metadata Access Service when you import an object from a Hadoop cluster. Create and configure a Metadata Access Service before you create HBase, HDFS, Hive, and MapR-DB connections.

Model Repository Service

The Model Repository Service manages the Model repository. The Model Repository Service connects to the Model repository when you run a mapping, mapping specification, profile, or workflow.

REST Operations Hub

The REST Operations Hub Service is an application service in the Informatica domain that exposes Informatica product functionality to external clients through REST APIs.

Repositories

Big Data Management uses repositories and other databases to store data related to connections, source metadata, data domains, data profiling, data masking, and data lineage. Big Data Management uses application services in the Informatica domain to access data in repositories.

Big Data Management uses the following databases:

Model repository

The Model repository stores profiles, data domains, mapping, and workflows that you manage in the Developer tool. The Model repository also stores profiles, data domains, and mapping specifications that you manage in the Analyst tool.

Profiling warehouse

The Data Integration Service runs profiles and stores profile results in the profiling warehouse.

Hadoop Integration

Big Data Management can connect to clusters that run different Hadoop distributions. Hadoop is an open-source software framework that enables distributed processing of large data sets across clusters of machines. You might also need to use third-party software clients to set up and manage your Hadoop cluster.

Big Data Management can connect to the supported data source in the Hadoop environment, such as HDFS, HBase, or Hive, and push job processing to the Hadoop cluster. To enable high performance access to files across the cluster, you can connect to an HDFS source. You can also connect to a Hive source, which is a data warehouse that connects to HDFS.

It can also connect to NoSQL databases such as HBase, which is a database comprising key-value pairs on Hadoop that performs operations in real-time. The Data Integration Service can push mapping jobs to the Spark or Blaze engine, and it can push profile jobs to the Blaze engine in the Hadoop environment.

Big Data Management supports more than one version of some Hadoop distributions. By default, the cluster configuration wizard populates the latest supported version.

Hadoop Utilities

Big Data Management uses third-party Hadoop utilities such as Sqoop to process data efficiently.

Sqoop is a Hadoop command line program to process data between relational databases and HDFS through MapReduce programs. You can use Sqoop to import and export data. When you use Sqoop, you do not need to install the relational database client and software on any node in the Hadoop cluster.

To use Sqoop, you must configure Sqoop properties in a JDBC connection and run the mapping in the Hadoop environment. You can configure Sqoop connectivity for relational data objects, customized data objects, and logical data objects that are based on a JDBC-compliant database. For example, you can configure Sqoop connectivity for the following databases:

- Aurora
- Greenplum
- IBM DB2
- IBM DB2 for z/OS
- Microsoft SQL Server
- Netezza
- Oracle
- Teradata
- Vertica

The Model Repository Service uses JDBC to import metadata. The Data Integration Service runs the mapping in the Hadoop run-time environment and pushes the job processing to Sqoop. Sqoop then creates map-reduce jobs in the Hadoop cluster, which perform the import and export job in parallel.

Specialized Sqoop Connectors

When you run mappings through Sqoop, you can use the following specialized connectors:

OraOop

You can use OraOop with Sqoop to optimize performance when you read data from or write data to Oracle. OraOop is a specialized Sqoop plug-in for Oracle that uses native protocols to connect to the Oracle database.

You can configure OraOop when you run Sqoop mappings on the Spark engine.

Teradata Connector for Hadoop (TDCH) Specialized Connectors for Sqoop

You can use the following TDCH specialized connectors for Sqoop to read data from or write data to Teradata:

- Cloudera Connector Powered by Teradata

- Hortonworks Connector for Teradata (powered by the Teradata Connector for Hadoop)
- MapR Connector for Teradata

These connectors are specialized Sqoop plug-ins that Cloudera, Hortonworks, and MapR provide for Teradata. They use native protocols to connect to the Teradata database.

Informatica supports Cloudera Connector Powered by Teradata and Hortonworks Connector for Teradata on the Blaze and Spark engines. When you run Sqoop mappings on the Blaze engine, you must configure these connectors. When you run Sqoop mappings on the Spark engine, the Data Integration Service invokes these connectors by default.

Informatica supports MapR Connector for Teradata on the Spark engine. When you run Sqoop mappings on the Spark engine, the Data Integration Service invokes the connector by default.

Note: For information about running native Teradata mappings with Sqoop, see the *Informatica PowerExchange for Teradata Parallel Transporter API User Guide*.

Databricks Integration

Big Data Management can push mappings to the Azure Databricks environment. Azure Databricks is an analytics cloud platform that is optimized for the Microsoft Azure cloud services. It incorporates the open-source Apache Spark cluster technologies and capabilities.

The Data Integration Service automatically installs the binaries required to integrate the Informatica domain with the Databricks environment. The integration requires Informatica connection objects and cluster configurations. A cluster configuration is a domain object that contains configuration parameters that you import from the Databricks cluster. You then associate the cluster configuration with connections to access the Databricks environment.

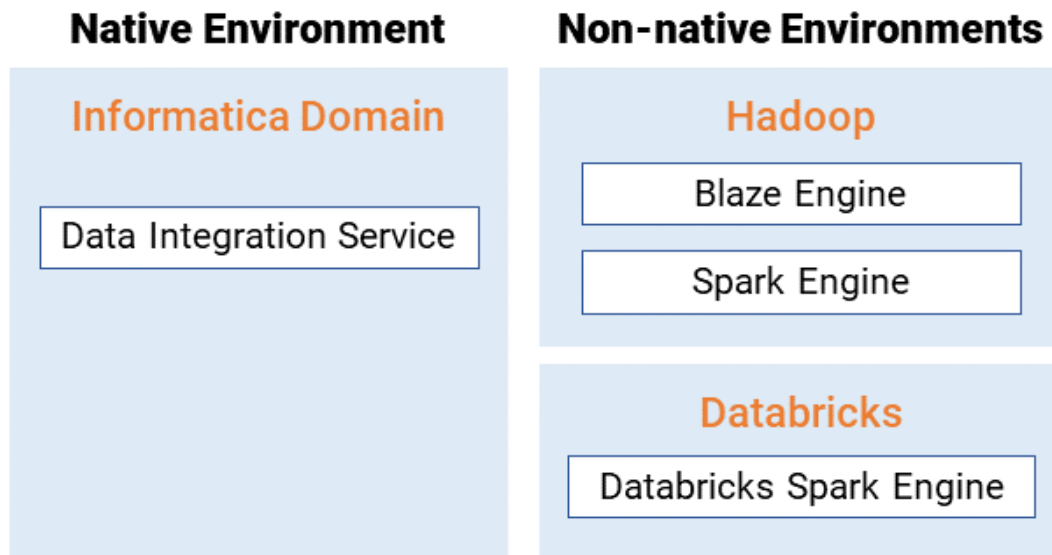
Big Data Management Engines

When you run a big data mapping, you can choose to run the mapping in the native environment or in a non-native environment, such as Hadoop or Databricks. When you validate a mapping, you can validate it against one or all of the engines. The Developer tool returns validation messages for each engine.

When you run a mapping in the native environment, the Data Integration Service in the Informatica domain runs the mapping. When you run the mapping in a non-native environment, the Data Integration Service pushes the run-time processing to a compute cluster in the non-native environment.

When you run the mapping in a non-native environment, the Data Integration Service uses a proprietary rule-based methodology to determine the best engine to run the mapping. The rule-based methodology evaluates the mapping sources and the mapping logic to determine the engine. The Data Integration Service translates the mapping logic into code that the engine can process, and it transfers the code to the engine.

The following image shows the processing environments and the run-time engines in the environments:



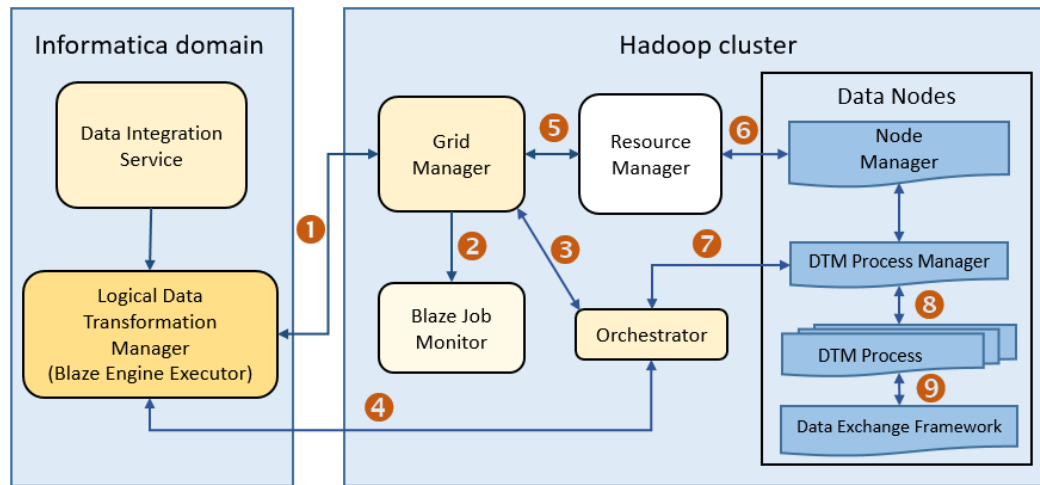
Run-time Process on the Blaze Engine

To run a mapping on the Informatica Blaze engine, the Data Integration Service submits jobs to the Blaze engine executor. The Blaze engine executor is a software component that enables communication between the Data Integration Service and the Blaze engine components on the Hadoop cluster.

The following Blaze engine components appear on the Hadoop cluster:

- Grid Manager. Manages tasks for batch processing.
- Orchestrator. Schedules and processes parallel data processing tasks on a cluster.
- Blaze Job Monitor. Monitors Blaze engine jobs on a cluster.
- DTM Process Manager. Manages the DTM Processes.
- DTM Processes. An operating system process started to run DTM instances.
- Data Exchange Framework. Shuffles data between different processes that process the data on cluster nodes.

The following image shows how a Hadoop cluster processes jobs sent from the Blaze engine executor:



The following events occur when the Data Integration Service submits jobs to the Blaze engine executor:

1. The Blaze Engine Executor communicates with the Grid Manager to initialize Blaze engine components on the Hadoop cluster, and it queries the Grid Manager for an available Orchestrator.
2. The Grid Manager starts the Blaze Job Monitor.
3. The Grid Manager starts the Orchestrator and sends Orchestrator information back to the LDTM.
4. The LDTM communicates with the Orchestrator.
5. The Grid Manager communicates with the Resource Manager for available resources for the Orchestrator.
6. The Resource Manager handles resource allocation on the data nodes through the Node Manager.
7. The Orchestrator sends the tasks to the DTM Processes through the DTM Process Manager.
8. The DTM Process Manager continually communicates with the DTM Processes.
9. The DTM Processes continually communicate with the Data Exchange Framework to send and receive data across processing units that run on the cluster nodes.

Application Timeline Server

The Hadoop Application Timeline Server collects basic information about completed application processes. The Timeline Server also provides information about completed and running YARN applications.

The Grid Manager starts the Application Timeline Server in the Yarn configuration by default.

The Blaze engine uses the Application Timeline Server to store the Blaze Job Monitor status. On Hadoop distributions where the Timeline Server is not enabled by default, the Grid Manager attempts to start the Application Timeline Server process on the current node.

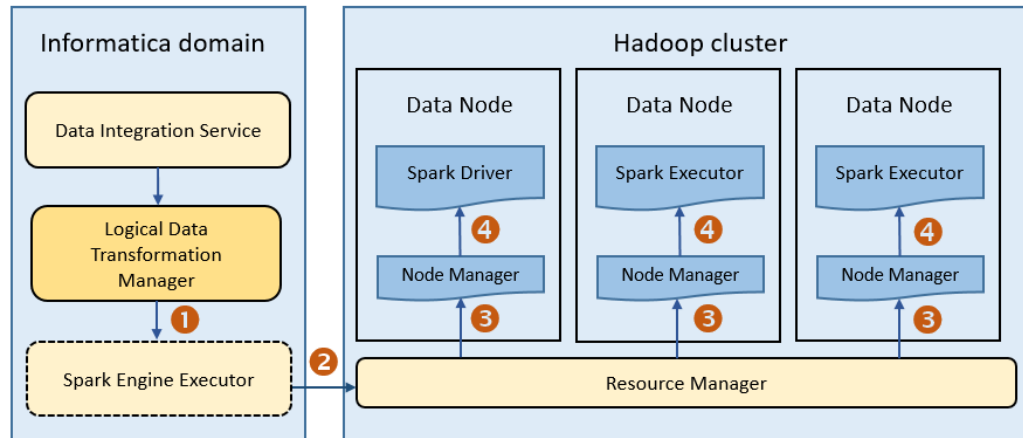
If you do not enable the Application Timeline Server on secured Kerberos clusters, the Grid Manager attempts to start the Application Timeline Server process in HTTP mode.

Run-time Process on the Spark Engine

The Data Integration Service can use the Spark engine on a Hadoop cluster to run Model repository mappings.

To run a mapping on the Spark engine, the Data Integration Service sends a mapping application to the Spark executor. The Spark executor submits the job to the Hadoop cluster to run.

The following image shows how a Hadoop cluster processes jobs sent from the Spark executor:



The following events occur when Data Integration Service runs a mapping on the Spark engine:

1. The Logical Data Transformation Manager translates the mapping into a Scala program, packages it as an application, and sends it to the Spark executor.
2. The Spark executor submits the application to the Resource Manager in the Hadoop cluster and requests resources to run the application.

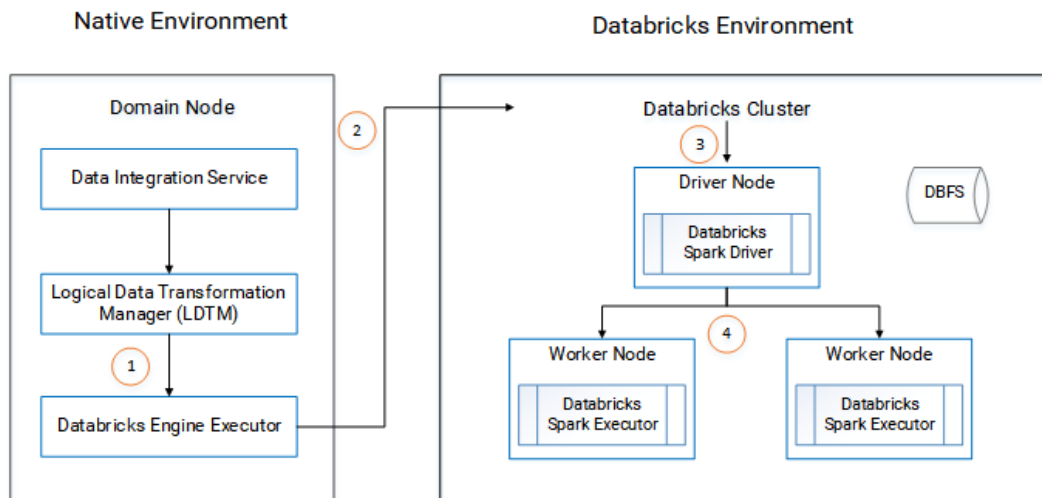
Note: When you run mappings on the HDInsight cluster, the Spark executor launches a spark-submit script. The script requests resources to run the application.

3. The Resource Manager identifies the Node Managers that can provide resources, and it assigns jobs to the data nodes.
4. Driver and Executor processes are launched in data nodes where the Spark application runs.

Run-time Process on the Databricks Spark Engine

When you run a job on the Databricks Spark engine, the Data Integration Service pushes the processing to the Databricks cluster, and the Databricks Spark engine runs the job.

The following image shows the components of the Informatica and the Databricks environments:



1. The Logical Data Transformation Manager translates the mapping into a Scala program, packages it as an application, and sends it to the Databricks Engine Executor on the Data Integration Service machine.
2. The Databricks Engine Executor submits the application through REST API to the Databricks cluster, requests to run the application, and stages files for access during run time.
3. The Databricks cluster passes the request to the Databricks Spark driver on the driver node.
4. The Databricks Spark driver distributes the job to one or more Databricks Spark executors that reside on worker nodes.
5. The executors run the job and stage run-time data to the Databricks File System (DBFS) of the workspace.

Big Data Process

As part of a big data project, you collect the data from diverse data sources. You can perform profiling, cleansing, and matching for the data. You build the business logic for the data and push the transformed data to the data warehouse. Then you can perform business intelligence on a view of the data.

Based on your big data project requirements, you can perform the following high-level tasks:

1. Collect the data.
2. Cleanse the data
3. Transform the data.
4. Process the data.
5. Monitor jobs.

Step 1. Collect the Data

Identify the data sources from which you need to collect the data.

Big Data Management provides several ways to access your data in and out of Hadoop based on the data types, data volumes, and data latencies in the data.

You can use PowerExchange adapters to connect to multiple big data sources. You can schedule batch loads to move data from multiple source systems to HDFS without the need to stage the data. You can move changed data from relational and mainframe systems into HDFS or the Hive warehouse. For real-time data feeds, you can move data off message queues and into HDFS.

You can collect the following types of data:

- Transactional
- Interactive
- Log file
- Sensor device
- Document and file
- Industry format

Step 2. Cleanse the Data

Cleanse the data by profiling, cleaning, and matching your data. You can view data lineage for the data.

You can perform data profiling to view missing values and descriptive statistics to identify outliers and anomalies in your data. You can view value and pattern frequencies to isolate inconsistencies or unexpected patterns in your data. You can drill down on the inconsistent data to view results across the entire data set.

You can automate the discovery of data domains and relationships between them. You can discover sensitive data such as social security numbers and credit card numbers so that you can mask the data for compliance.

After you are satisfied with the quality of your data, you can also create a business glossary from your data. You can use the Analyst tool or Developer tool to perform data profiling tasks. Use the Analyst tool to perform data discovery tasks. Use Metadata Manager to perform data lineage tasks.

Step 3. Transform the Data

You can build the business logic to parse data in the Developer tool. Eliminate the need for hand-coding the transformation logic by using pre-built Informatica transformations to transform data.

Step 4. Process the Data

Based on your business logic, you can determine the optimal run-time environment to process your data. If your data is less than 10 terabytes, consider processing your data in the native environment. If your data is greater than 10 terabytes, consider processing your data on a compute cluster.

Step 5. Monitor Jobs

Monitor the status of your processing jobs. You can view monitoring statistics for your processing jobs in the Monitoring tool. After your processing jobs complete you can get business intelligence and analytics from your data.

Data Warehouse Optimization Mapping Example

You can optimize an enterprise data warehouse with the Hadoop system to store more terabytes of data cheaply in the warehouse.

For example, you need to analyze customer portfolios by processing the records that have changed in a 24-hour time period. You can offload the data on Hadoop, find the customer records that have been inserted, deleted, and updated in the last 24 hours, and then update those records in your data warehouse. You can capture these changes even if the number of columns change or if the keys change in the source files.

To capture the changes, you can create the following mappings in the Developer tool:

Mapping_Day1

Create a mapping to read customer data from flat files in a local file system and write to an HDFS target for the first 24-hour period.

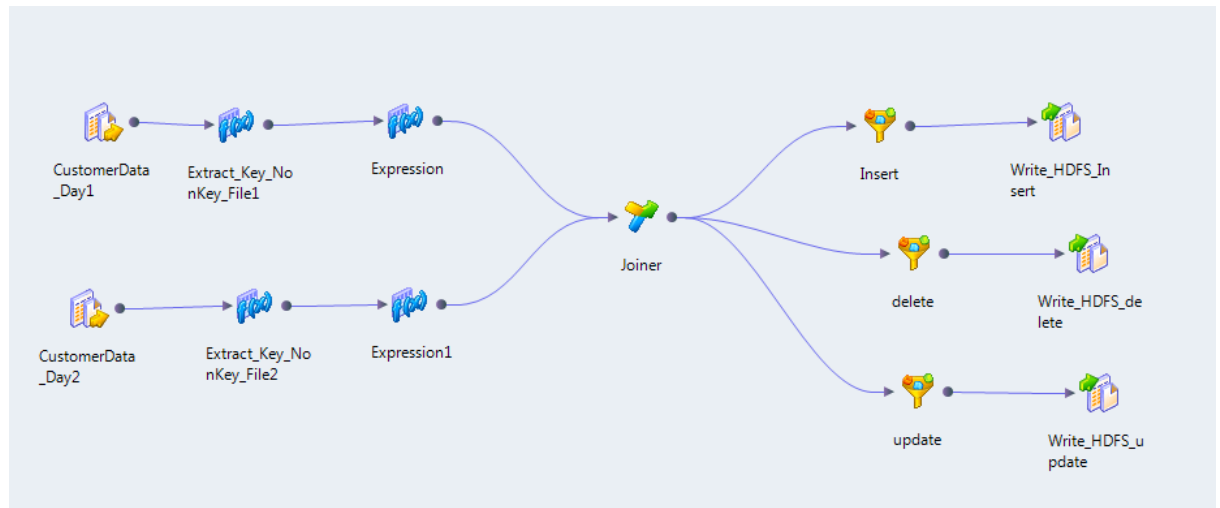
Mapping_Day2

Create a mapping to read customer data from flat files in a local file system and write to an HDFS target for the next 24-hour period.

m_CDC_DWHOptimization

Create a mapping to capture the changed data. The mapping reads data from HDFS and identifies the data that has changed. To increase performance, you configure the mapping to run on Hadoop cluster nodes in a Hadoop environment.

The following image shows the mapping m_CDC_DWHOptimization:



The mapping contains the following objects:

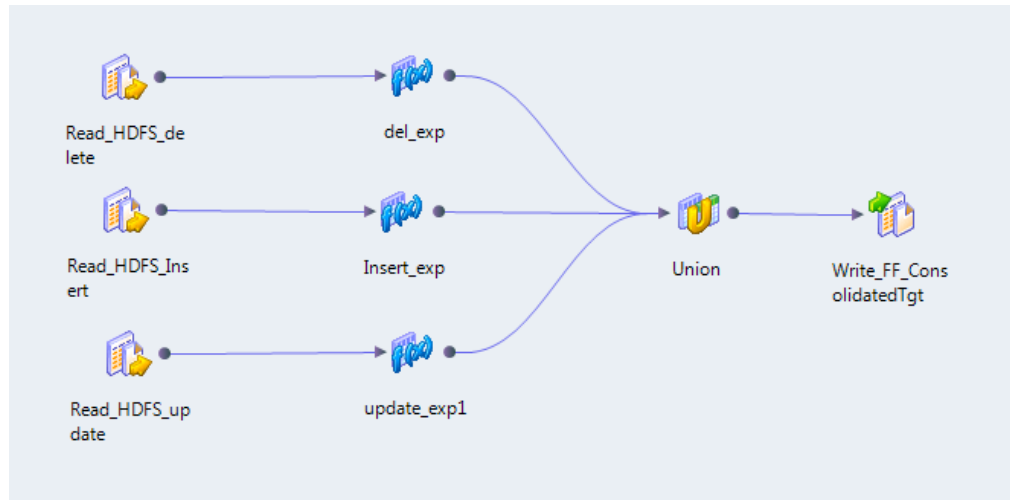
- Read transformations. Transformations that read data from HDFS files that were the targets of Mapping_Day1 and Mapping_Day2. The Data Integration Service reads all of the data as a single column.
- Expression transformations. Extract a key from the non-key values in the data. The expressions use the INSTR function and SUBSTR function to perform the extraction of key values.
- Joiner transformation. Performs a full outer join on the two sources based on the keys generated by the Expression transformations.
- Filter transformations. Use the output of the Joiner transformation to filter rows based on whether or not the rows should be updated, deleted, or inserted.

- Write transformations. Transformations that write the data to three HDFS files based on whether the data is inserted, deleted, or updated.

Consolidated_Mapping

Create a mapping to consolidate the data in the HDFS files and load the data to the data warehouse.

The following figure shows the mapping Consolidated_Mapping:

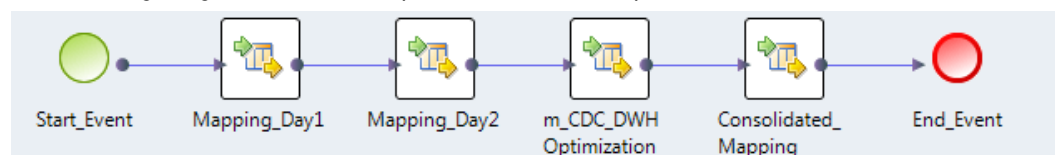


The mapping contains the following objects:

- Read transformations. Transformations that read data from HDFS files that were the target of the previous mapping are the sources of this mapping.
- Expression transformations. Add the deleted, updated, or inserted tags to the data rows.
- Union transformation. Combines the records.
- Write transformation. Transformation that writes data to the flat file that acts as a staging location on the local file system.

You can open each mapping and right-click to run the mapping. To run all mappings in sequence, use a workflow.

The following image shows the example Data Warehouse Optimization workflow:



To run the workflow, use the `infacmd wfs startWorkflow` command.

CHAPTER 2

Mappings

This chapter includes the following topics:

- [Overview of Mappings, 27](#)
- [Mapping Run-time Properties, 28](#)
- [Sqoop Mappings in a Hadoop Environment, 34](#)
- [Rules and Guidelines for Mappings in a Non-native Environment, 39](#)
- [Workflows that Run Mappings in a Non-native Environment, 40](#)
- [Configuring a Mapping to Run in a Non-native Environment, 41](#)
- [Mapping Execution Plans, 42](#)
- [Optimization for the Hadoop Environment, 45](#)
- [Troubleshooting Mappings in a Non-native Environment, 52](#)
- [Mappings in the Native Environment, 53](#)

Overview of Mappings

Configure the run-time environment in the Developer tool to optimize mapping performance and process data that is greater than 10 terabytes. You can choose to run a mapping in the native environment or in a non-native environment. When you run mappings in the native environment, the Data Integration Service processes and runs the mapping. When you run mappings in a non-native environment, the Data Integration Service pushes the processing to a compute cluster, such as Hadoop or Databricks.

You can run standalone mappings and mappings that are a part of a workflow in the non-native environment.

When you select the Hadoop environment, you can also select the engine to push the mapping logic to the Hadoop cluster. Based on the mapping logic, the Data Integration Service can push the mapping logic to one of the following engines in the Hadoop environment:

- Informatica Blaze engine. An Informatica proprietary engine for distributed processing on Hadoop.
- Spark engine. A high performance engine for batch processing that can run on a Hadoop cluster or on a Spark standalone mode cluster.

When you select the Databricks environment, the Integration Service pushes the mapping logic to the Databricks Spark engine, the Apache Spark engine packaged for Databricks.

When you select multiple engines, the Data Integration Service determines the best engine to run the mapping during validation. You can also choose to select which engine the Data Integration Service uses. You might select an engine based on whether an engine supports a particular transformation or based on the format in which the engine returns data.

When you run a mapping in a non-native environment, you must configure a connection to access the environment. You can set the run-time properties for the environment and for the engine that runs the mapping.

You can view the execution plan for a mapping to run in the non-native environment. View the execution plan for the engine that the Data Integration Service selects to run the mapping.

You can monitor Hive queries and Hadoop jobs in the Monitoring tool. Monitor the jobs on a Hadoop cluster with the YARN Web User Interface or the Blaze Job Monitor web application.

The Data Integration Service logs messages from the Blaze, Spark and Databricks Spark engines and the DTM.

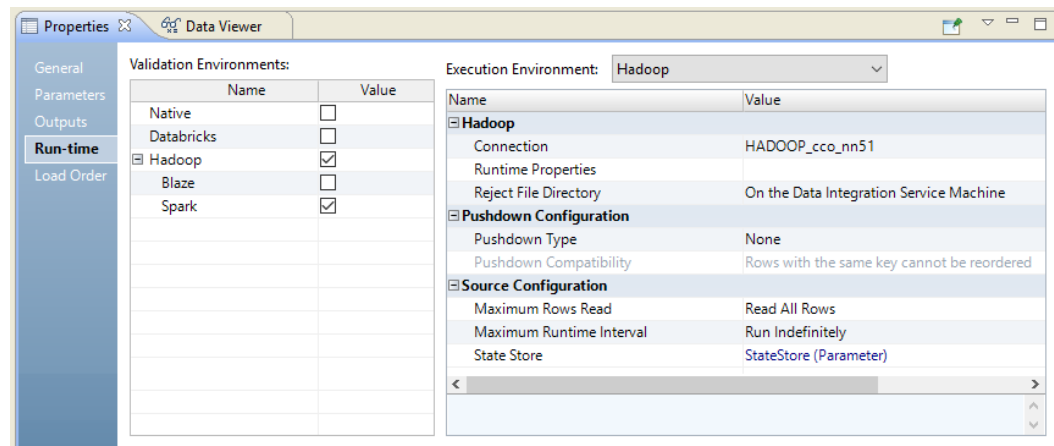
Mapping Run-time Properties

The mapping run-time properties depend on the environment that you select for the mapping.

The mapping properties contain the **Validation Environments** area and an **Execution Environment** area. The properties in the **Validation Environment** indicate whether the Developer tool validates the mapping definition for the native execution environment, the non-native execution environment, or both. When you run a mapping in the native environment, the Data Integration Service processes the mapping.

When you run a mapping in the non-native environment, the Data Integration Service pushes the mapping execution to the compute cluster through a runtime connection. The engine in the compute cluster processes the mapping.

The following image shows the mapping **Run-time** properties in a Hadoop environment:



Validation Environments

The properties in the **Validation Environments** indicate whether the Developer tool validates the mapping definition for the native or non-native execution environment.

You can configure the following properties for the **Validation Environments**:

Native

Default environment. The Data Integration Service runs the mapping in a native environment.

Hadoop

Run the mapping in the Hadoop environment. The Data Integration Service pushes the transformation logic to the Hadoop cluster through a Hadoop connection. Select the engine to process the mapping. You can select the Blaze or Spark engine.

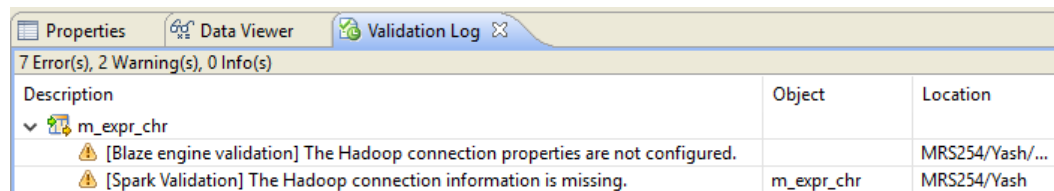
Databricks

Run the mapping in the Databricks environment. The Data Integration Service pushes the transformation logic to the Databricks cluster through a Databricks connection. The Databricks cluster processes the mapping on the Databricks Spark engine.

You can use a mapping parameter to indicate the execution environment for the mapping. When you select the execution environment, click **Assign Parameter**. Configure a string parameter. Set the default value to native, hadoop, or spark-databricks.

When you validate the mapping, validation occurs for each engine that you choose in the **Validation Environments**. The validation log might contain validation errors specific to each engine. If the mapping is valid for at least one mapping, the mapping is valid. The errors for the other engines appear in the validation log as warnings. If the mapping is valid for multiple engines, you can view the execution plan to determine which engine will run the job. You can view the execution plan in the **Data Viewer** view.

The following image shows a sample validation log:



7 Error(s), 2 Warning(s), 0 Info(s)		
Description	Object	Location
▼ m_expr_chr		
⚠ [Blaze engine validation] The Hadoop connection properties are not configured.		MRS254/Yash/...
⚠ [Spark Validation] The Hadoop connection information is missing.	m_expr_chr	MRS254/Yash

Execution Environment

Configure non-native properties, pushdown configuration properties, and source configuration properties in the **Execution Environment** area.

The following table describes properties that you can configure for the Hadoop and Databricks environments:

Name	Description
Connection	Configure for the Hadoop and Databricks environments. Defines the connection information that the Data Integration Service requires to push the mapping execution to the compute cluster. Select the non-native connection to run the mapping in the compute cluster. You can assign a user-defined parameter for the non-native connection.
Runtime Properties	Configure for the Hadoop environment. You can configure run-time properties for the Hadoop environment in the Data Integration Service, the Hadoop connection, and in the mapping. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Hadoop connection or in the mapping. The Data Integration Service processes property overrides based on the following priorities: 1. Mapping custom properties set using <code>infacmd ms runMapping</code> with the <code>-cp</code> option 2. Mapping run-time properties for the Hadoop environment 3. Hadoop connection advanced properties for run-time engines 4. Hadoop connection advanced general properties, environment variables, and classpaths 5. Data Integration Service custom properties
Reject File Directory	Configure for the Hadoop environment. The directory for Hadoop mapping files on HDFS when you run mappings in the Hadoop environment. The Blaze engine can write reject files to the Hadoop environment for flat file, HDFS, and Hive targets. The Spark engine can write reject files to the Hadoop environment for flat file and HDFS targets. Choose one of the following options: - On the Hadoop Cluster. The reject files are moved to the reject directory configured in the Hadoop connection. If the directory is not configured, the mapping will fail. - Defer to the Hadoop Connection. The reject files are moved based on whether the reject directory is enabled in the Hadoop connection properties. If the reject directory is enabled, the reject files are moved to the reject directory configured in the Hadoop connection. Otherwise, the Data Integration Service stores the reject files based on the <code>RejectDir</code> system parameter.

You can configure the following pushdown configuration properties:

Name	Description
Pushdown type	Configure for the Hadoop environment. Choose one of the following options: - None. Select no pushdown type for the mapping. - Source. The Data Integration Service tries to push down transformation logic to the source database. - Full. The Data Integration Service pushes the full transformation logic to the source database.
Pushdown Compatibility	Configure for the Hadoop environment. Optionally, if you choose full pushdown optimization and the mapping contains an Update Strategy transformation, you can choose a pushdown compatibility option or assign a pushdown compatibility parameter. Choose one of the following options: - Multiple rows do not have the same key. The transformation connected to the Update Strategy transformation receives multiple rows without the same key. The Data Integration Service can push the transformation logic to the target. - Multiple rows with the same key can be reordered. The target transformation connected to the Update Strategy transformation receives multiple rows with the same key that can be reordered. The Data Integration Service can push the Update Strategy transformation to the non-native environment. - Multiple rows with the same key cannot be reordered. The target transformation connected to the Update Strategy transformation receives multiple rows with the same key that cannot be reordered. The Data Integration Service cannot push the Update Strategy transformation to the non-native environment.

You can configure the following source properties for the Hadoop and Databricks environments:

Name	Description
Maximum Rows Read	Reserved for future use.
Maximum Runtime Interval	Reserved for future use.
State Store	Reserved for future use.

Parsing JSON Records on the Spark Engines

In the mapping run-time properties, you can configure how the Spark engine parses corrupt records and multiline records when it reads from JSON sources in a mapping.

Configure the following Spark run-time properties:

infaspark.json.parser.mode

Specifies the parser how to handle corrupt JSON records. You can set the value to one of the following modes:

- DROPMALFORMED. The parser ignores all corrupted records. Default mode.
- PERMISSIVE. The parser accepts non-standard fields as nulls in corrupted records.
- FAILFAST. The parser generates an exception when it encounters a corrupted record and the Spark application goes down.

infaspark.json.parser.multiLine

Specifies whether the parser can read a multiline record in a JSON file. You can set the value to true or false. Default is false. Applies only to non-native distributions that use Spark version 2.2.x and above.

Reject File Directory

You can write reject files to the Data Integration Service machine or to the Hadoop cluster. Or, you can defer to the Hadoop connection configuration. The Blaze engine can write reject files to the Hadoop environment for flat file, HDFS, and Hive targets. The Spark engine can write reject files to the Hadoop environment for flat file and HDFS targets.

If you configure the mapping run-time properties to defer to the Hadoop connection, the reject files for all mappings with this configuration are moved based on whether you choose to write reject files to Hadoop for the active Hadoop connection. You do not need to change the mapping run-time properties manually to change the reject file directory.

For example, if the reject files are currently moved to the Data Integration Service machine and you want to move them to the directory configured in the Hadoop connection, edit the Hadoop connection properties to write reject files to Hadoop. The reject files of all mappings that are configured to defer to the Hadoop connection are moved to the configured directory.

You might also choose to defer to the Hadoop connection when the connection is parameterized to alternate between multiple Hadoop connections. For example, the parameter might alternate between one Hadoop connection that is configured to move reject files to the Data Integration Service machine and another Hadoop connection that is configured to move reject files to the directory configured in the Hadoop connection. If you choose to defer to the Hadoop connection, the reject files are moved depending on the active Hadoop connection in the connection parameter.

You cannot create the reject file directory at run time. Depending on the option that you choose, the configured reject file directory must exist on the Data Integration Service machine or in the Hadoop environment before you can store reject files.

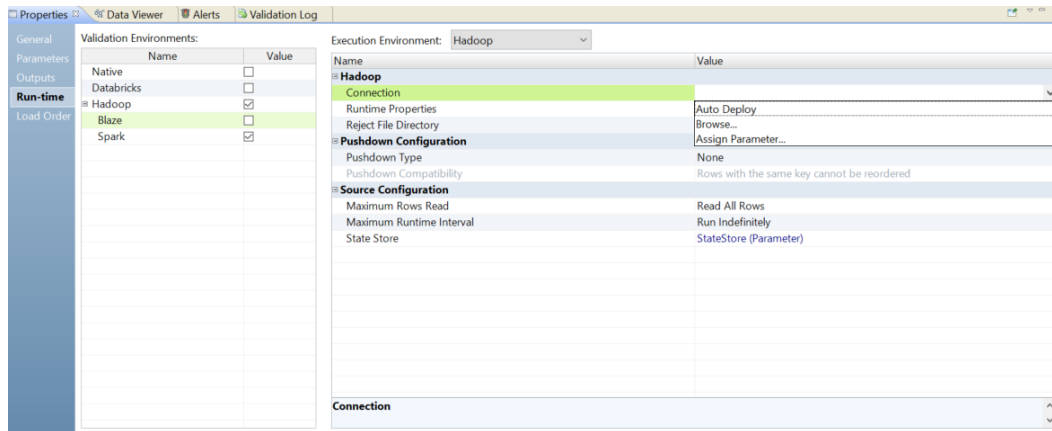
Changing the Compute Cluster for a Mapping Run

You can change the compute cluster that a mapping runs on.

You might want to change the compute cluster to run the mapping on a different compute cluster type or different version.

To change the compute cluster that a mapping runs on, click the connection where it appears in the mapping run-time properties. Then click **Browse**, and select a non-native connection that uses a cluster configuration for the cluster where you want to run the mapping.

The following image shows the **Browse** choice for a non-native connection:



After you change the connection that the mapping uses, you must restart the Data Integration Service for the change to take effect.

Updating Run-time Properties for Multiple Mappings

You can enable or disable the validation environment or set the execution environment for multiple mappings. You can update multiple mappings that you run from the Developer tool or mappings that are deployed to a Data Integration Service. Use the Command Line Interface to perform these updates.

The following table describes the commands to update mapping run-time properties:

Command	Description
dis disableMappingValidationEnvironment	Disables the mapping validation environment for mappings that are deployed to the Data Integration Service.
mrs disableMappingValidationEnvironment	Disables the mapping validation environment for mappings that you run from the Developer tool.
dis enableMappingValidationEnvironment	Enables a mapping validation environment for mappings that are deployed to the Data Integration Service.
mrs enableMappingValidationEnvironment	Enables a mapping validation environment for mappings that you run from the Developer tool.
dis setMappingExecutionEnvironment	Specifies the mapping execution environment for mappings that are deployed to the Data Integration Service.
mrs setMappingExecutionEnvironment	Specifies the mapping execution environment for mappings that you run from the Developer tool.

Sqoop Mappings in a Hadoop Environment

You can use a JDBC connection that is enabled for Sqoop connectivity to import a Sqoop source or Sqoop target and create a mapping. You can run Sqoop mappings on the Blaze and Spark engines.

If you use Cloudera Connector Powered by Teradata or Hortonworks Connector for Teradata, you can run mappings on the Blaze or Spark engines. If you use MapR Connector for Teradata, you can run mappings on the Spark engine.

In the mapping, you can specify additional Sqoop arguments and disable the Sqoop connector.

Note: If you add or delete a Type 4 JDBC driver .jar file required for Sqoop connectivity from the `externaljdbcjars` directory, changes take effect after you restart the Data Integration Service. If you run the mapping on the Blaze engine, changes take effect after you restart the Data Integration Service and Blaze Grid Manager. When you run the mapping for the first time, you do not need to restart the Data Integration Service and Blaze Grid Manager. You need to restart the Data Integration Service and Blaze Grid Manager only for the subsequent mapping runs.

Sqoop Mapping-Level Arguments

If a data object uses Sqoop, you can click the corresponding **Read** transformation or **Write** transformation in the Sqoop mapping to define the arguments that Sqoop must use to process the data. The Data Integration Service merges the additional Sqoop arguments that you specify in the mapping with the arguments that you specified in the JDBC connection and constructs the Sqoop command.

The Sqoop arguments that you specify in the mapping take precedence over the arguments that you specified in the JDBC connection. However, if you do not enable the Sqoop connector in the JDBC connection but enable the Sqoop connector in the mapping, the Data Integration Service does not run the mapping through Sqoop. The Data Integration Service runs the mapping through JDBC.

You can configure the following Sqoop arguments in a Sqoop mapping:

- `m` or `num-mappers`
- `split-by`
- `batch`
- `infaoptimize`
- `infaownername`
- `schema`
- `verbose`

For a complete list of the Sqoop arguments that you can configure, see the Sqoop documentation.

`m` or `num-mappers`

The `m` or `num-mappers` argument defines the number of map tasks that Sqoop must use to import and export data in parallel.

Use the following syntax:

```
-m <number of map tasks>
--num-mappers <number of map tasks>
```

If you configure the `m` argument or `num-mappers` argument, you must also configure the `split-by` argument to specify the column based on which Sqoop must split the work units.

Use the `m` argument or `num-mappers` argument to increase the degree of parallelism. You might have to test different values for optimal performance.

When you configure the `m` argument or `num-mappers` argument and run Sqoop mappings on the Spark or Blaze engines, Sqoop dynamically creates partitions based on the file size.

Note: If you configure the `num-mappers` argument to export data on the Blaze or Spark engine, Sqoop ignores the argument. Sqoop creates map tasks based on the number of intermediate files that the Blaze or Spark engine creates.

split-by

The `split-by` argument defines the column based on which Sqoop splits work units.

Use the following syntax:

```
--split-by <column_name>
```

You can configure the `split-by` argument to improve the performance. If the primary key does not have an even distribution of values between the minimum and maximum range, you can configure the `split-by` argument to specify another column that has a balanced distribution of data to split the work units.

If you do not define the `split-by` column, Sqoop splits work units based on the following criteria:

- If the data object contains a single primary key, Sqoop uses the primary key as the `split-by` column.
- If the data object contains a composite primary key, Sqoop defaults to the behavior of handling composite primary keys without the `split-by` argument. See the Sqoop documentation for more information.
- If the data object does not contain a primary key, the value of the `m` argument and `num-mappers` argument default to 1.

Rules and Guidelines for the split-by Argument

Consider the following restrictions when you configure the `split-by` argument:

- If you configure the `split-by` argument and the `split-by` column contains NULL values, Sqoop does not import the rows that contain NULL values. However, the mapping runs successfully and no error is written in the YARN log.
- If you configure the `split-by` argument and the `split-by` column contains special characters, the Sqoop import process fails.
- The `split-by` argument is required in the following scenarios:
 - You use Cloudera Connector Powered by Teradata or Hortonworks Connector for Teradata, and the Teradata table does not contain a primary key.
 - You create a custom query to override the default query when you import data from a Sqoop source.

batch

The `batch` argument indicates that Sqoop must export data in batches instead of exporting data row by row.

Use the following syntax:

```
--batch
```

You can configure the `batch` argument to improve the performance.

The `batch` argument is required when you configure Sqoop in the following scenarios:

- You write data to a Vertica target.
- You write data to an IBM DB2 for z/OS target.

- You write data to a Teradata target by running the mapping with a generic JDBC connector instead of a TDCH specialized connector for Sqoop. You can invoke the generic JDBC connector by configuring the `--driver` and `--connection-manager` Sqoop arguments in the JDBC connection.

infaoptimize

The `infaoptimize` argument defines whether you want to disable the performance optimization of Sqoop pass-through mappings on the Spark engine.

When you run a Sqoop pass-through mapping on the Spark engine, the Data Integration Service optimizes mapping performance in the following scenarios:

- You read data from a Sqoop source and write data to a Hive target that uses the Text format.
- You read data from a Sqoop source and write data to an HDFS target that uses the Flat, Avro, or Parquet format.

If you want to disable the performance optimization, set the `--infaoptimize` argument to `false`. For example, if you see data type issues after you run an optimized Sqoop mapping, you can disable the performance optimization.

Use the following syntax:

```
--infaoptimize false
```

infaownername

The `infaownername` argument indicates whether Sqoop must honor the owner name for a data object.

To configure Sqoop to honor the owner name for a data object, set the value of the `infaownername` argument to `true`.

Use the following syntax:

```
--infaownername true
```

Note: To honor the owner name for a Microsoft SQL Server database, you must configure the `--schema` argument and specify the owner name. Use the following syntax:

```
-- --schema <owner_name>
```

schema

The `schema` argument indicates whether Sqoop must honor the owner name for a data object that represents a Microsoft SQL Server database.

To honor the owner name for a Microsoft SQL Server source database, specify the owner name in the **Additional Sqoop Import Arguments** field in the Sqoop mapping. Use the following syntax:

```
-- --schema <owner_name>
```

To honor the owner name for a Microsoft SQL Server target database, you must leave the **Owner** field blank in the Sqoop target data object. Specify the owner name in the **Additional Sqoop Export Arguments** field in the Sqoop mapping. Use the following syntax:

```
-- --schema <owner_name>
```

verbose

Configure the `--verbose` argument when you want detailed information in the logs for debugging purposes.

Use the following syntax:

```
--verbose
```

Note: If you configure the `--verbose` argument, you might experience significant performance degradation, especially if you run the Sqoop mapping on the Spark engine.

Incremental Data Extraction for Sqoop Mappings

You can configure a Sqoop mapping to perform incremental data extraction based on an ID or timestamp based source column. With incremental data extraction, Sqoop extracts only the data that changed since the last data extraction. Incremental data extraction increases the mapping performance.

To perform incremental data extraction, configure arguments in the **Additional Sqoop Import Arguments** text box of the Read transformation in the Sqoop mapping.

You must configure all of the following arguments for incremental data extraction:

--infa-incremental-type

Indicates whether you want to perform the incremental data extraction based on an ID or timestamp column.

If the source table contains an ID column that acts as a primary key or unique key column, you can configure incremental data extraction based on the ID column. Set the value of the `--infa-incremental-type` argument as `ID`. Sqoop extracts rows whose IDs are greater than the last extracted ID.

If the source table contains a timestamp column that contains the last updated time for all rows, you can configure incremental data extraction based on the timestamp column. Set the value of the `--infa-incremental-type` argument as `timestamp`. Sqoop extracts rows whose timestamps are greater than the last read timestamp value or the maximum timestamp value.

Use the following syntax:

```
--infa-incremental-type <ID or timestamp>
```

--infa-incremental-key

Indicates the column name based on which Sqoop must perform the incremental data extraction.

Use the following syntax:

```
--infa-incremental-key <column_name>
```

--infa-incremental-value

Indicates the column value that Sqoop must use as the baseline value to perform the incremental data extraction. Sqoop extracts all rows that have a value greater than the value defined in the `--infa-incremental-value` argument.

Enclose the column value within double quotes.

Use the following syntax:

```
--infa-incremental-value "<column_value>"
```

--infa-incremental-value-format

Applicable if you configure incremental data extraction based on a timestamp column. Indicates the timestamp format of the column value defined in the `--infa-incremental-value` argument.

Enclose the format value within double quotes.

Use the following syntax:

```
--infa-incremental-value-format "<format>"
```

Default is "MM/dd/yyyy HH:mm:ss.SSSSSSSS".

Note: If you want to perform incremental data extraction but you do not configure all the required arguments, the Sqoop mapping fails. The argument values are not case sensitive.

Example

The Sales department in your organization stores information related to new customers in Salesforce. You want to create a Sqoop mapping to read customer records that were created after a certain date and write the data to SAP.

In the Sqoop mapping, you can define the Sqoop import arguments as follows:

```
--infa-incremental-type timestamp --infa-incremental-key CreatedDate --infa-incremental-value  
"10/20/2018 00:00:00.000000000" --infa-incremental-value-format "MM/dd/yyyy  
HH:mm:ss.SSSSSSSS"
```

The Data Integration Service reads all customer records whose `CreatedDate` column contains timestamp values greater than 10/20/2018 00:00:00.000000000. The Sqoop mapping uses the 10/20/2018 00:00:00.000000000 format while reading data.

Configuring Sqoop Properties in the Mapping

You can specify additional Sqoop arguments and disable the Sqoop connector at the mapping level. The Sqoop arguments that you specify at the mapping level take precedence over the arguments that you specified in the JDBC connection.

1. Open the mapping that contains the data object for which you want to configure Sqoop properties.
2. Select the Read or Write transformation that is associated with the data object.
3. Click the **Advanced** tab.
4. To disable the Sqoop connector for the data object, select the **Disable Sqoop Connector** check box.
5. Perform one of the following steps:
 - To specify additional Sqoop import arguments for the data object, enter the import arguments in the **Additional Sqoop Import Arguments** text box.
 - To specify additional Sqoop export arguments for the data object, enter the export arguments in the **Additional Sqoop Export Arguments** text box.

The Data Integration Service merges the additional Sqoop arguments that you specified in the mapping with the arguments that you specified in the JDBC connection and constructs the Sqoop command. The Data Integration Service then invokes Sqoop on a Hadoop node.

Configuring Parameters for Sqoop Arguments in the Mapping

In a Sqoop mapping, you can parameterize the Sqoop import and export arguments.

1. Open the Sqoop mapping where you want to parameterize the Sqoop arguments.
2. Click the **Parameters** tab, and define a parameter name and parameter value.

For example, enter the parameter name as `Param_splitByColumn`

3. Select the Read or Write transformation that is associated with the data object.
4. Click the **Advanced** tab.
5. Perform one of the following steps:
 - To parameterize Sqoop import arguments for the data object, click the **Additional Sqoop Import Arguments** text box and enter the parameter name you defined in the **Parameters** tab.
 - To parameterize Sqoop export arguments for the data object, click the **Additional Sqoop Export Arguments** text box and enter the parameter name you defined in the **Parameters** tab.

Use the following format to specify the parameter name for a Sqoop import or export argument:

```
--<Sqoop_argument_name> $<parameter_name>
```

For example, enter `--split-by $Param_splitByColumn`

Rules and Guidelines for Mappings in a Non-native Environment

When you run mappings in a non-native environment, some differences in processing and configuration apply. You can run mappings in a non-native environment on the Blaze, Spark, or Databricks Spark engine. Consider the following rules and guidelines when you run mappings in the native and non-native environments.

Rules and Guidelines for Mappings on the Blaze Engine

Consider the following run-time differences on the Blaze engine:

- In a Hadoop environment, sources that have data errors in a column result in a null value for the column. In the native environment, the Data Integration Service does not process the rows that have data errors in a column.
- When you cancel a mapping that reads from a flat file source, the file copy process that copies flat file data to HDFS may continue to run. The Data Integration Service logs the command to kill this process in the Hive session log, and cleans up any data copied to HDFS. Optionally, you can run the command to kill the file copy process.
- If you use Hive 2.0 or higher, the Blaze engine guarantees scale values. For example, when the Blaze engine processes the decimal 1.1234567 with scale 9 using Hive 2.0, the output is 1.123456700.

Rules and Guidelines for Mappings on the Spark Engine

Consider the following run-time differences on the Spark engine:

- Set the optimizer level to none or minimal if a mapping validates but fails to run. If you set the optimizer level to use cost-based or semi-join optimization methods, the Data Integration Service ignores this at run-time and uses the default.

- The run-time engine does not honor the early projection optimization method in all cases. If the Data Integration Service removes the links between unused ports, the run-time engine might reconnect the ports.
- In a Hadoop environment, sources that have data errors in a column result in a null value for the column. In the native environment, the Data Integration Service does not process the rows that have data errors in a column.
- When you cancel a mapping that reads from a flat file source, the file copy process that copies flat file data to HDFS may continue to run. The Data Integration Service logs the command to kill this process in the Hive session log, and cleans up any data copied to HDFS. Optionally, you can run the command to kill the file copy process.

When the Spark engine runs a mapping, it processes jobs on the cluster using HiveServer2 in the following cases:

- The mapping writes to a target that is a Hive table bucketed on fields of type char or varchar.
- The mapping reads from or writes to Hive transaction-enabled tables.
- The mapping reads from or writes to Hive tables where column-level security is enabled.
- The mapping writes to a Hive target and is configured to create or replace the table at run time.

Rules and Guidelines for Mappings on the Databricks Spark Engine

Consider the following run-time differences on the Databricks Spark engine:

- Set the optimizer level to none or minimal if a mapping validates but fails to run. If you set the optimizer level to use cost-based or semi-join optimization methods, the Data Integration Service ignores this at run-time and uses the default.
- The run-time engine does not honor the early projection optimization method in all cases. If the Data Integration Service removes the links between unused ports, the run-time engine might reconnect the ports.

Workflows that Run Mappings in a Non-native Environment

You can add a mapping that you configured to run in a non-native environment to a Mapping task in a workflow. When you deploy and run the workflow, the Mapping task runs the mapping.

You might decide to run a mapping from a workflow so that you can make decisions during the workflow run. You can configure a workflow to run multiple mappings in sequence or in parallel. You can configure a workflow to send emails that notify users about the status of the Mapping tasks.

If you add a mapping to a workflow that runs on the Spark or Databricks Spark engine, you can also configure mapping outputs. You can persist a mapping output in the Model repository to assign the persisted output to a Mapping task input, or you can bind a mapping output to a workflow variable to pass the value of the output to other tasks in the workflow.

Note: You cannot use the SUM aggregation type in mapping outputs on the Spark or Databricks Spark engine.

For more information about mapping outputs, see the "Mapping Task" chapter in the *Developer Workflow Guide*.

When a Mapping task runs a mapping configured to run on the Blaze engine, do not assign the Mapping task outputs to workflow variables. Mappings that run on the Blaze engine do not provide the total number of target, source, and error rows. When a Mapping task includes a mapping that runs on the Blaze engine, the task outputs contain a value of zero (0).

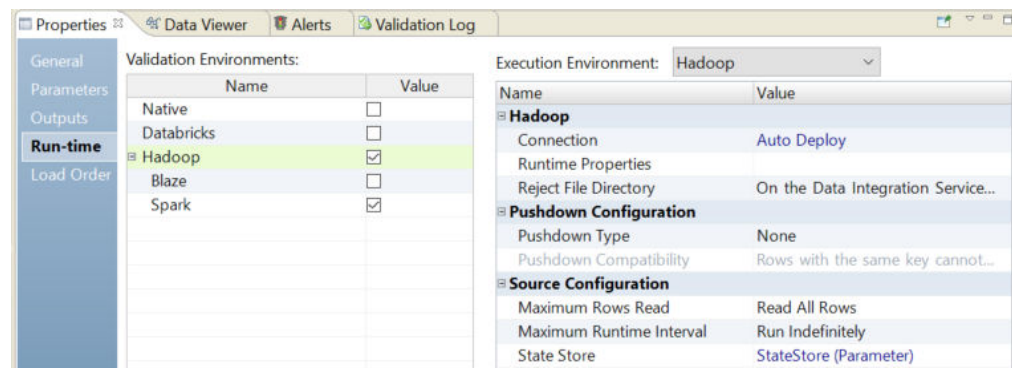
Configuring a Mapping to Run in a Non-native Environment

You can configure a mapping to run in a non-native environment. To configure a mapping, you must select a validation environment and an execution environment.

1. Select a mapping from a project or folder from the **Object Explorer** view to open in the editor.
2. In the **Properties** view, select the **Run-time** tab.
3. Select **Hadoop**, **Databricks**, or both as the value for the validation environment.

When you select the Hadoop environment, the Blaze and Spark engines are selected by default. Disable the engines that you do not want to use.

4. Select **Hadoop** or **Databricks** for the execution environment.



5. Select **Connection** and use the drop down in the value field to browse for a connection or to create a connection parameter:
 - To select a connection, click **Browse** and select a connection.
 - To create a connection parameter, click **Assign Parameter**.
6. Configure the rest of the properties for the execution environment.
7. Right-click an empty area in the editor and click **Validate**.
The Developer tool validates the mapping.
8. View validation errors on the **Validation Log** tab.
9. Click the **Data Viewer** view.
10. Click **Show Execution Plan** to view the execution plan for the mapping.

Mapping Execution Plans

The Data Integration Service generates an execution plan to run mappings on a Blaze, Spark, or Databricks Spark engine. The Data Integration Service translates the mapping logic into code that the run-time engine can execute. You can view the plan in the Developer tool before you run the mapping and in the Administrator tool after you run the mapping.

The Data Integration Service generates mapping execution plans to run on the following engines:

Informatica Blaze engine

The Blaze engine execution plan simplifies the mapping into segments. It contains tasks to start the mapping, run the mapping, and clean up the temporary tables and files. It contains multiple tasklets and the task recovery strategy. It also contains pre- and post-grid task preparation commands for each mapping before running the main mapping on a compute cluster. A pre-grid task can include a task such as copying data to HDFS. A post-grid task can include tasks such as cleaning up temporary files or copying data from HDFS.

Spark engine

The Spark execution plan shows the run-time Scala code that runs the mapping logic. A translation engine translates the mapping into an internal representation of the logic. The internal representation is rendered into Scala code that accesses the Spark API. You can view the Scala code in the execution plan to debug the logic.

Databricks Spark engine

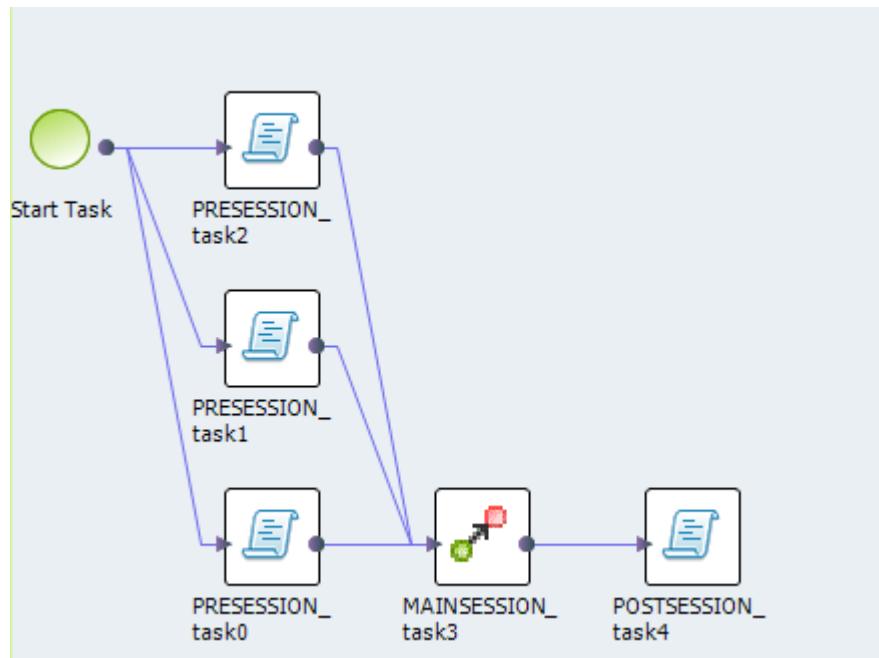
The Databricks Spark execution plan shows the run-time Scala code that runs the mapping logic. A translation engine translates the mapping into an internal representation of the logic. The internal representation is rendered into Scala code that accesses the Spark API. You can view the Scala code in the execution plan to debug the logic.

Blaze Engine Execution Plan Details

You can view details of the Blaze engine execution plan in the Administrator tool and Developer tool.

In the Developer tool, the Blaze engine execution plan appears as a workflow. You can click on each component in the workflow to get the details.

The following image shows the Blaze execution plan in the Developer tool:



The Blaze engine execution plan workflow contains the following components:

- Start task. The workflow start task.
- Command task. The pre-processing or post-processing task for local data.
- Grid mapping. An Informatica mapping that the Blaze engine compiles and distributes across a cluster of nodes.
- Grid task. A parallel processing job request sent by the Blaze engine executor to the Grid Manager.
- Grid segment. Segment of a grid mapping that is contained in a grid task.
- Tasklet. A partition of a grid segment that runs on a separate DTM.

In the Administrator tool, the Blaze engine execution plan appears as a script.

The following image shows the Blaze execution script:

Test - XraKb1DXEeWg		Properties	Blaze Execution Plan	Summary
Script Id	Script			
MAINSESSION_task3	Execution scriptStep [MAINSESSION_task3], type [GridTaskStepImpl]. With "from" step(s): PRESESSION_task0, PRESESSION_task1, PRESESSION_task2. With "to" step(s): POSTSESSION_task4. Grid mapping task has totally [3] substeps: Execution step [submapping-2], type [SegmentStepImpl]. With no "from" step. With "to" step(s): submapping-3. Included instances: Read_IN_OUT[SourceTx], DETarget_Joiner_G1[TargetTx], Execution step [submapping-1], type [SegmentStepImpl]. With no "from" step. With "to" step(s): submapping-3. Included instances: DETarget_Joiner_G0[TargetTx], Read_IN_OUT1[SourceTx], Execution step [submapping-3], type [SegmentStepImpl]. With "from" step(s): submapping-1, submapping-2. With no "to" step. Included instances: Write_IN_OUT[TargetTx], DESource_Joiner_G1[SourceTx], Joiner[JoinerTx], DESource_Joiner_G0[SourceTx],			

In the Administrator tool, the Blaze engine execution plan has the following details:

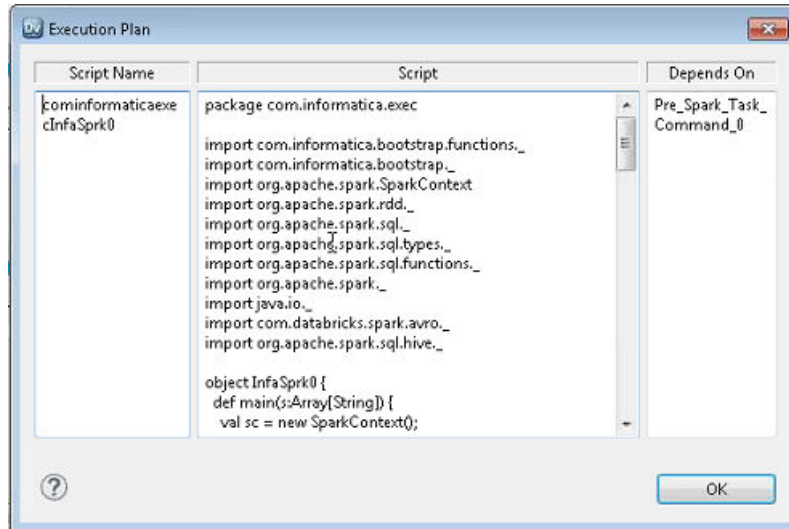
- Script ID. Unique identifier for the Blaze engine script.
- Script. Blaze engine script that the Data Integration Service generates based on the mapping logic.
- Depends on. Tasks that the script depends on. Tasks include other scripts and Data Integration Service tasks, like the Start task.

Spark Engine Execution Plan Details

You can view the details of a Spark engine execution plan from the Administrator tool or Developer tool.

The Spark engine execution plan shows the Scala code on the Databricks Spark engine.

The following image shows the execution plan for a mapping to run on the Spark engine:



The Spark engine execution plan has the following details:

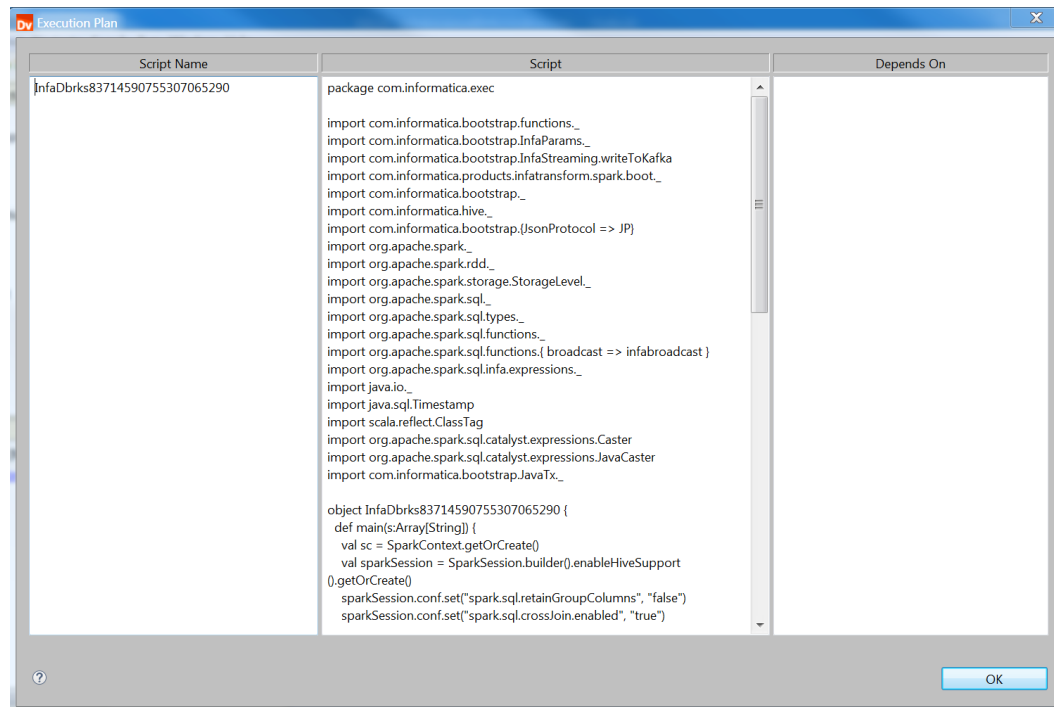
- Script ID. Unique identifier for the Spark engine script.
- Script. Scala code that the Data Integration Service generates based on the mapping logic.
- Depends on. Tasks that the script depends on. Tasks include other scripts and Data Integration Service tasks.

Databricks Spark Engine Execution Details

You can view the details of a Databricks Spark engine execution plan from the Administrator tool or Developer tool.

The Databricks Spark engine execution plan shows the Scala code to run on the Databricks Spark engine.

The following image shows the execution plan for a mapping to run on the Databricks Spark engine:



The Databricks Spark engine execution plan has the following details:

- Script ID. Unique identifier for the Databricks Spark engine script.
- Script. Scala code that the Data Integration Service generates based on the mapping logic.
- Depends on. Tasks that the script depends on. Tasks include other scripts and Data Integration Service tasks.

Viewing the Execution Plan

You can view the engine execution plan for a mapping. You do not have to run the mapping to view the execution plan in the Developer tool.

Note: You can also view the execution plan in the Administrator tool.

1. To view the execution plan in the Developer tool, select the **Data Viewer** view for the mapping and click **Show Execution Plan**.
2. Select the **Data Viewer** view.
3. Select **Show Execution Plan**.

The **Data Viewer** view shows the details for the execution plan.

Optimization for the Hadoop Environment

Optimize the Hadoop environment to increase performance.

You can optimize the Hadoop environment in the following ways:

Configure a highly available Hadoop cluster.

You can configure the Data Integration Service and the Developer tool to read from and write to a highly available Hadoop cluster. The steps to configure a highly available Hadoop cluster depend on the type of Hadoop distribution. For more information about configuration steps for a Hadoop distribution, see the *Informatica Big Data Management Integration Guide*.

Compress data on temporary staging tables.

You can enable data compression on temporary staging tables to increase mapping performance.

Run mappings on the Blaze engine.

Run mappings on the highly available Blaze engine. The Blaze engine enables restart and recovery of grid tasks and tasklets by default.

Perform parallel sorts.

When you use a Sorter transformation in a mapping, the Data Integration Service enables parallel sorting by default when it pushes the mapping logic to the Hadoop cluster. Parallel sorting improves mapping performance.

Partition Joiner transformations.

When you use a Joiner transformation in a Blaze engine mapping, the Data Integration Service can apply map-side join optimization to improve mapping performance. The Data Integration Service applies map-side join optimization if the master table is smaller than the detail table. When the Data Integration Service applies map-side join optimization, it moves the data to the Joiner transformation without the cost of shuffling the data.

Truncate partitions in a Hive target.

You can truncate partitions in a Hive target to increase performance. To truncate partitions in a Hive target, you must choose to both truncate the partition in the Hive target and truncate the target table.

Assign resources on Hadoop clusters.

You can use schedulers to assign resources on a Hadoop cluster. You can use a capacity scheduler or a fair scheduler depending on the needs of your organization.

Configure YARN queues to share resources on Hadoop clusters.

You can configure YARN queues to redirect jobs on the Hadoop cluster to specific queues. The queue where a job is assigned defines the resources that are allocated to perform the job.

Label nodes in a Hadoop cluster.

You can label nodes in a Hadoop cluster to divide the cluster into partitions that have specific characteristics.

Optimize Sqoop mappings on the Spark engine.

The Data Integration Service can optimize the performance of Sqoop pass-through mappings that run on the Spark engine.

Enable big data job recovery.

You can enable big data job recovery to recover mapping jobs that the Data Integration Service pushes to the Spark engine for processing.

Blaze Engine High Availability

The Blaze engine is a highly available engine that determines the best possible recovery strategy for grid tasks and tasklets.

Based on the size of the grid task, the Blaze engine attempts to apply the following recovery strategy:

- No high availability. The Blaze engine does not apply a recovery strategy.
- Full restart. Restarts the grid task.

Enabling Data Compression on Temporary Staging Tables

To optimize performance when you run a mapping in the Hadoop environment, you can enable data compression on temporary staging tables. When you enable data compression on temporary staging tables, mapping performance might increase.

To enable data compression on temporary staging tables, complete the following steps:

1. Configure the Hadoop connection to use the codec class name that the Hadoop cluster uses to enable compression on temporary staging tables.
2. Configure the Hadoop cluster to enable compression on temporary staging tables.

Hadoop provides following compression libraries for the following compression codec class names:

Compression Library	Codec Class Name	Performance Recommendation
Zlib	org.apache.hadoop.io.compress.DefaultCodec	n/a
Gzip	org.apache.hadoop.io.compress.GzipCodec	n/a
Snappy	org.apache.hadoop.io.compress.SnappyCodec	Recommended for best performance.
Bz2	org.apache.hadoop.io.compress.BZip2Codec	Not recommended. Degrades performance.
LZO	com.hadoop.compression.lzo.LzoCodec	n/a

Step 1. Enable Data Compression in the Hadoop Connection

Use the Administrator tool or the Developer tool to configure the Hadoop connection to enable data compression on temporary staging tables.

1. In the Hadoop connection properties, edit the properties to run mappings in a Hadoop cluster.
2. Select **Temporary Table Compression Codec**.
3. Choose to select a predefined codec class name or enter a custom codec class name.
 - To select a predefined codec class name, select a compression library from the list.
 - To enter a custom codec class name, select custom from the list and enter the codec class name that matches the codec class name in the Hadoop cluster.

Step 2. Enable Data Compression on the Hadoop Environment

To enable compression on temporary staging tables, you must install a compression codec on the Hadoop cluster.

For more information about how to install a compression codec, refer to the Apache Hadoop documentation.

1. Verify that the native libraries for the compression codec class name are installed on every node on the cluster.
2. To include the compression codec class name that you want to use, update the property `io.compression.codecs` in `core-site.xml`. The value for this property is a comma separated list of all the codec class names supported on the cluster.
3. Verify that the Hadoop-native libraries for the compression codec class name that you want to use are installed on every node on the cluster.
4. Verify that the `LD_LIBRARY_PATH` variable on the Hadoop cluster includes the locations of both the native and Hadoop-native libraries where you installed the compression codec.

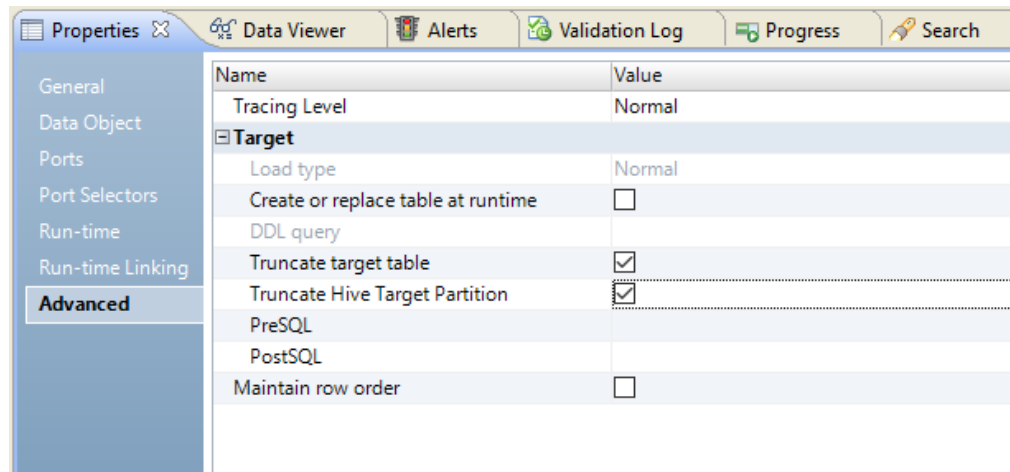
Truncating Partitions in a Hive Target

To truncate partitions in a Hive target, you must edit the write properties for the customized data object that you created for the Hive target in the Developer tool.

You can truncate partitions in a Hive target when you use the Blaze or Spark run-time engines to run the mapping.

1. Open the customized data object in the editor.
2. To edit write properties, select the **Input** transformation in the **Write** view, and then select the **Advanced** properties.

The following image shows the **Advanced** properties tab:



3. Select **Truncate Hive Target Partition**.
4. Select **Truncate target table**.

Scheduling, Queuing, and Node Labeling

You can use YARN schedulers, YARN queues, and node labels to optimize performance when you run a mapping in the Hadoop environment.

A YARN scheduler assigns resources to YARN applications on the Hadoop cluster while honoring organizational policies on sharing resources. You can configure YARN to use a fair scheduler or a capacity scheduler. A fair scheduler shares resources evenly among all jobs running on the cluster over time. A capacity scheduler allows multiple organizations to share a large cluster and distributes resources based on capacity allocations. The capacity scheduler guarantees each organization a certain capacity and distributes any excess capacity that is underutilized.

YARN queues are organizing structures for YARN schedulers and allow multiple tenants to share a cluster. The capacity of each queue specifies the percentage of cluster resources that are available for applications submitted to the queue. You can redirect Blaze, Spark, and Sqoop jobs to specific YARN queues.

Node labels allow YARN queues to run on specific nodes in a cluster. You can use node labels to partition a cluster into sub-clusters such that jobs run on nodes with specific characteristics. For example, you might label nodes that process data faster compared to other nodes. Nodes that are not labeled belong to the default partition. You can associate the node labels with capacity scheduler queues.

You can also use the node labels to configure the Blaze engine. When you use node labels to configure the Blaze engine, you can specify the nodes on the Hadoop cluster where you want the Blaze engine to run.

Enable Scheduling and Node Labeling

To enable scheduling and node labeling in the Hadoop environment, update the `yarn-site.xml` properties in the cluster configuration.

Configure the following properties:

yarn.resourcemanager.scheduler.class

Defines the YARN scheduler that the Data Integration Service uses to assign resources on the cluster.

```
<property>
  <name>yarn.resourcemanager.scheduler.class</name>
  <value>org.apache.hadoop.yarn.server.resourcemanager.scheduler.[Scheduler Type].
[Scheduler Type]Scheduler</value>
</property>
```

For example:

```
<property>
  <name>yarn.resourcemanager.scheduler.class</name>

  <value>org.apache.hadoop.yarn.server.resourcemanager.scheduler.capacity.CapacityScheduler</value>
</property>
```

yarn.node-labels.enabled

Enables node labeling.

```
<property>
  <name>yarn.node-labels.enabled</name>
  <value>TRUE</value>
</property>
```

yarn.node-labels.fs-store.root-dir

The HDFS location to update the node label dynamically.

```
<property>
  <name>yarn.node-labels.fs-store.root-dir</name>
  <value>hdfs://[Node name]:[Port]/[Path to store]/[Node labels]</value>
</property>
```

Define YARN Queues

You can define multiple YARN queues to redirect jobs to specific queues. You can redirect Blaze, Spark, and Sqoop jobs.

To redirect jobs on the Blaze engine to a specific queue, configure the following Blaze configuration property in the Hadoop connection:

Property	Description
YARN Queue Name	The YARN scheduler queue name used by the Blaze engine that specifies available resources on a cluster.

To redirect jobs on the Spark engine to a specific queue, configure the following Spark configuration property in the Hadoop connection:

Property	Description
YARN Queue Name	The YARN scheduler queue name used by the Spark engine that specifies available resources on a cluster. The name is case sensitive.

To redirect Sqoop jobs to a specific queue, configure the following JDBC connection property:

Property	Description
Sqoop Arguments	The Sqoop connection-level argument to direct a MapReduce job to a specific YARN queue. Use the following format: <code>-Dmapred.job.queue.name=<YARN queue name></code> If you do not direct the job to a specific queue, the Spark engine uses the default queue.

Configure the Blaze Engine to Use Node Labels

If you enable node labeling in the Hadoop environment, you can use node labels to run Blaze components as YARN applications on specific cluster nodes.

To run Blaze components on the labeled cluster nodes, specify the node labels when you configure the Blaze engine.

To specify node labels on the Blaze engine, list the node labels in the following Hadoop connection property:

Property	Description
Blaze YARN Node Label	Node label that determines the node on the Hadoop cluster where the Blaze engine runs. If you do not specify a node label, the Blaze engine runs on the nodes in the default partition. If the Hadoop cluster supports logical operators for node labels, you can specify a list of node labels. To list the node labels, use the operators <code>&&</code> (AND), <code> </code> (OR), and <code>!</code> (NOT).

Note: When the Blaze engine uses node labels, Blaze components might be redundant on the labeled nodes. If a node contains multiple labels and you specify the labels in different Hadoop connections, multiple Grid Manager, Orchestrator, or Job Monitor instances might run on the same node.

Big Data Job Recovery

An administrator can enable big data job recovery to recover a big data job configured to run on the Spark engine when a Data Integration Service node stops unexpectedly.

When a Data Integration Service node fails before a running job is complete, the Data Integration Service sends the job to another node, which resumes processing job tasks from the point at which the node failure occurred. Recovery occurs upon node startup.

To use big data recovery, you must configure jobs to run on the Spark engine and submit jobs from the infacmd client.

An administrator configures big data recovery in Data Integration Service properties. For more information about big data job recovery, see the *Big Data Management Administrator Guide*.

Spark Engine Optimization for Sqoop Pass-Through Mappings

When you run a pass-through mapping with a Sqoop source on the Spark engine, the Data Integration Service optimizes mapping performance in the following scenarios:

- You write data to a Hive target that uses the Text format.
- You write data to a Hive target that was created with a custom DDL query.
- You write data to an existing Hive target that is either partitioned with a custom DDL query or partitioned and bucketed with a custom DDL query.
- You write data to an existing Hive target that is both partitioned and bucketed.
- You write data to an HDFS target that uses the Flat, Avro, or Parquet format.

If you want to disable the performance optimization, set the `--infaoptimize` argument to `false` in the JDBC connection or Sqoop mapping. For example, if you see data type issues after you run an optimized Sqoop mapping, you can disable the performance optimization.

Use the following syntax:

```
--infaoptimize false
```

Rules and Guidelines for Sqoop Spark Engine Optimization

Consider the following rules and guidelines when you run Sqoop mappings on the Spark engine:

- The Data Integration Service does not optimize mapping performance in the following scenarios:
 - There are unconnected ports between the source and target in the mapping.
 - The data types of the source and target in the mapping do not match.
 - You write data to an existing Hive target table that is either partitioned or bucketed.
 - You run a mapping on an Azure HDInsight cluster that uses WASB to write data to an HDFS complex file target of the Parquet format.
- If you configure Hive-specific Sqoop arguments to write data to a Hive target, Sqoop ignores the arguments.
- If you configure a delimiter for a Hive target table that is different from the default delimiter, Sqoop ignores the delimiter.

Troubleshooting Mappings in a Non-native Environment

Consider troubleshooting tips for mappings in a non-native environment.

Hadoop Environment

When I run a mapping with a Hive source or a Hive target on a different cluster, the Data Integration Service fails to push the mapping to Hadoop with the following error: Failed to execute query [exec0_query_6] with error code [10], error message [FAILED: Error in semantic analysis: Line 1:181 Table not found customer_eur], and SQL state [42000]].

When you run a mapping in a Hadoop environment, the Hive connection selected for the Hive source or Hive target, and the mapping must be on the same Hive metastore.

When I run a mapping with SQL overrides concurrently, the mapping hangs.

There are not enough available resources because the cluster is being shared across different engines.

Configure YARN to use the capacity scheduler and use different YARN scheduler queues for Blaze and Spark.

When I configure a mapping to create a partitioned Hive table, the mapping fails with the error "Need to specify partition columns because the destination table is partitioned."

This issue happens because of internal Informatica requirements for a query that is designed to create a Hive partitioned table. For details and a workaround, see [Knowledge Base article 516266](#).

Databricks Environment

Mappings fail with the following error: SEVERE: Run with ID [1857] failed with state [INTERNAL_ERROR] and error message [Library installation timed out after 1800 seconds. Libraries that are not yet installed: jar: "dbfs:/tmp/DATABRICKS/sess6250142538173973565/staticCode.jar"]

This might happen when you run concurrent jobs. When Databricks does not have resources to process a job, it queues the job for a maximum of 1,800 seconds (30 minutes). If resources are not available in 30 minutes, the job fails. Consider the following actions to avoid timeouts:

- Configure preemption environment variables on the Databricks cluster to control the amount of resources that get allocated to each job. For more information about preemption, see the *Big Data Management Integration Guide*.
- Run cluster workflows to create ephemeral clusters. You can configure the workflow to create a cluster, run the job, and then delete the cluster. For more information about ephemeral clusters, see [Chapter 6, "Cluster Workflows" on page 112](#).

Important: Informatica integrates with Databricks, supporting standard concurrency clusters. Standard concurrency clusters have a maximum queue time of 30 minutes, and jobs fail when the timeout is reached. The maximum queue time cannot be extended. Setting the preemption threshold allows more jobs to run concurrently, but with a lower percentage of allocated resources, the jobs can take longer to run. Also, configuring the environment for preemption does not ensure that all jobs will run. In addition to configuring preemption, you might choose to run cluster workflows to create ephemeral clusters that create the cluster, run the job, and then delete the cluster. For more information about Databricks concurrency, contact Azure Databricks.

Mappings in the Native Environment

You can run a mapping in the native environment. In the native environment, the Data Integration Service runs the mapping from the Developer tool. You can run standalone mappings or mappings that are a part of a workflow.

In the native environment, you can read and process data from large unstructured and semi-structured files, Hive, or social media web sites. You can include the following objects in the mappings:

- Hive sources
- Flat file sources or targets in the local system or in HDFS
- Complex file sources in the local system or in HDFS
- Data Processor transformations to process unstructured and semi-structured file formats
- Social media sources

Data Processor Mappings

The Data Processor transformation processes unstructured and semi-structured file formats in a mapping. It converts source data to flat CSV records that MapReduce applications can process.

You can configure the Data Processor transformation to process messaging formats, relational data, HTML pages, XML, JSON, AVRO, Parquet, Cobol, Microsoft Excel, Microsoft Word, and PDF documents. You can also configure it to transform structured formats such as ACORD, HIPAA, HL7, EDI-X12, EDIFACT, and SWIFT.

For example, an application produces hundreds of data files per second and writes the files to a directory. You can create a mapping that extracts the files from the directory, passes them to a Data Processor transformation, and writes the data to a target.

HDFS Mappings

Create an HDFS mapping to read or write to HDFS.

You can read and write fixed-width and delimited file formats. You can read or write compressed files. You can read text files and binary file formats such as sequence file from HDFS. You can specify the compression format of the files. You can use the binary stream output of the complex file data object as input to a Data Processor transformation to parse the file.

You can define the following objects in an HDFS mapping:

- Flat file data object or complex file data object operation as the source to read data from HDFS.
- Transformations.
- Flat file data object as the target to write data to HDFS or any target.

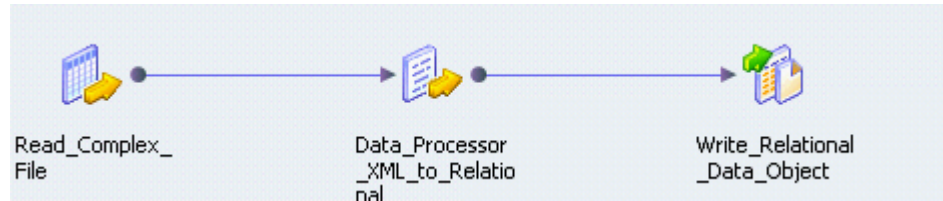
Validate and run the mapping. You can deploy the mapping and run it or add the mapping to a Mapping task in a workflow.

HDFS Data Extraction Mapping Example

Your organization needs to analyze purchase order details such as customer ID, item codes, and item quantity. The purchase order details are stored in a semi-structured compressed XML file in HDFS. The hierarchical data includes a purchase order parent hierarchy level and a customer contact details child hierarchy level. Create a mapping that reads all the purchase records from the file in HDFS. The mapping must convert the hierarchical data to relational data and write it to a relational target.

You can use the extracted data for business analytics.

The following figure shows the example mapping:



You can use the following objects in the HDFS mapping:

HDFS Input

The input object, Read_Complex_File, is a Read transformation that represents a compressed XML file stored in HDFS.

Data Processor Transformation

The Data Processor transformation, Data_Processor_XML_to_Relational, parses the XML file and provides a relational output.

Relational Output

The output object, Write_Relational_Data_Object, is a Write transformation that represents a table in an Oracle database.

When you run the mapping, the Data Integration Service reads the file in a binary stream and passes it to the Data Processor transformation. The Data Processor transformation parses the specified file and provides a relational output. The output is written to the relational target.

You can configure the mapping to run in a native or Hadoop run-time environment.

Complete the following tasks to configure the mapping:

1. Create an HDFS connection to read files from the Hadoop cluster.
2. Create a complex file data object read operation. Specify the following parameters:
 - The file as the resource in the data object.
 - The file compression format.
 - The HDFS file location.
3. Optionally, you can specify the input format that the Mapper uses to read the file.
4. Drag and drop the complex file data object read operation into a mapping.
5. Create a Data Processor transformation. Configure the following properties in the Data Processor transformation:
 - An input port set to buffer input and binary data type.
 - Relational output ports depending on the number of columns you want in the relational output. Specify the port size for the ports. Use an XML schema reference that describes the XML hierarchy. Specify the normalized output that you want. For example, you can specify PurchaseOrderNumber_Key as a generated key that relates the Purchase Orders output group to a Customer Details group.
 - Create a Streamer object and specify Streamer as a startup component.
6. Create a relational connection to an Oracle database.
7. Import a relational data object.
8. Create a write transformation for the relational data object and add it to the mapping.

Hive Mappings

Based on the mapping environment, you can read data from or write data to Hive.

In a native environment, you can read data from Hive. To read data from Hive, complete the following steps:

1. Create a Hive connection.
2. Configure the Hive connection mode to access Hive as a source or target.
3. Use the Hive connection to create a data object to read from Hive.
4. Add the data object to a mapping and configure the mapping to run in the native environment.

You can write to Hive in a Hadoop environment. To write data to Hive, complete the following steps:

1. Create a Hive connection.
2. Configure the Hive connection mode to access Hive as a source or target.
3. Use the Hive connection to create a data object to write to Hive.
4. Add the data object to a mapping and configure the mapping to run in the Hadoop environment.

You can define the following types of objects in a Hive mapping:

- A Read Transformation to read data from Hive
- Transformations
- A target or an SQL data service. You can write to Hive if you run the mapping in a Hadoop cluster.

Validate and run the mapping. You can deploy the mapping and run it or add the mapping to a Mapping task in a workflow.

Hive Mapping Example

Your organization, HypoMarket Corporation, needs to analyze customer data. Create a mapping that reads all the customer records. Create an SQL data service to make a virtual database available for end users to query.

You can use the following objects in a Hive mapping:

Hive input

The input file is a Hive table that contains the customer names and contact details.

Create a relational data object. Configure the Hive connection and specify the table that contains the customer data as a resource for the data object. Drag the data object into a mapping as a read data object.

SQL Data Service output

Create an SQL data service in the Developer tool. To make it available to end users, include it in an application, and deploy the application to a Data Integration Service. When the application is running, connect to the SQL data service from a third-party client tool by supplying a connect string.

You can run SQL queries through the client tool to access the customer data.

Social Media Mappings

Create mappings to read social media data from sources such as Facebook and LinkedIn.

You can extract social media data and load them to a target in the native environment only. You can choose to parse this data or use the data for data mining and analysis.

To process or analyze the data in Hadoop, you must first move the data to a relational or flat file target and then run the mapping in the Hadoop cluster.

You can use the following Informatica adapters in the Developer tool:

- PowerExchange for DataSift
- PowerExchange for Facebook
- PowerExchange for LinkedIn
- PowerExchange for Twitter
- PowerExchange for Web Content-Kapow Katalyst

Review the respective PowerExchange adapter documentation for more information.

Twitter Mapping Example

Your organization, Hypomarket Corporation, needs to review all the tweets that mention your product HypoBasket with a positive attitude since the time you released the product in February 2012.

Create a mapping that identifies tweets that contain the word HypoBasket and writes those records to a table.

You can use the following objects in a Twitter mapping:

Twitter input

The mapping source is a Twitter data object that contains the resource Search.

Create a physical data object and add the data object to the mapping. Add the Search resource to the physical data object. Modify the query parameter with the following query:

```
QUERY=HypoBasket:&since:2012-02-01
```

Sorter transformation

Optionally, sort the data based on the timestamp.

Add a Sorter transformation to the mapping. Specify the timestamp as the sort key with direction as ascending.

Mapping output

Add a relational data object to the mapping as a target.

After you run the mapping, Data Integration Service writes the extracted tweets to the target table. You can use text analytics and sentiment analysis tools to analyze the tweets.

Native Environment Optimization

You can optimize the native environment to increase performance. To increase performance, you can configure the Data Integration Service to run on a grid and to use multiple partitions to process data. You can also enable high availability to ensure that the domain can continue running despite temporary network, hardware, or service failures.

You can run profiles, sessions, and workflows on a grid to increase the processing bandwidth. A grid is an alias assigned to a group of nodes that run profiles, sessions, and workflows. When you enable grid, the Data Integration Service runs a service process on each available node of the grid to increase performance and scalability.

You can also run mapping with partitioning to increase performance. When you run a partitioned session or a partitioned mapping, the Data Integration Service performs the extract, transformation, and load for each partition in parallel.

You can configure high availability for the domain. High availability eliminates a single point of failure in a domain and provides minimal service interruption in the event of failure.

Processing Big Data on a Grid

You can run an Integration Service on a grid to increase the processing bandwidth. When you enable grid, the Integration Service runs a service process on each available node of the grid to increase performance and scalability.

Big data may require additional bandwidth to process large amounts of data. For example, when you run a Model repository profile on an extremely large data set, the Data Integration Service grid splits the profile into multiple mappings and runs the mappings simultaneously on different nodes in the grid.

When you run mappings on a grid, the Data Integration Service distributes the mappings to multiple DTM processes on nodes in the grid. When you run a profile on a grid, the Data Integration Service splits the profile into multiple mappings and distributes the mappings to multiple DTM processes on nodes in the grid. You can optimize the grid to increase performance and scalability of the Data Integration Service.

To optimize the grid, add nodes to the grid to increase processing bandwidth of the Data Integration Service.

For more information about the Data Integration Service grid, see the Informatica Administrator Guide.

Processing Big Data on Partitions

You can run a Model repository mapping with partitioning to increase performance. When you run a mapping configured with partitioning, the Data Integration Service performs the extract, transformation, and load for each partition in parallel.

Mappings that process large data sets can take a long time to process and can cause low data throughput. When you configure partitioning, the Data Integration Service uses additional threads to process the session or mapping which can increase performance.

If the nodes where mappings run have multiple CPUs, you can enable the Data Integration Service to maximize parallelism when it runs mappings. When you maximize parallelism, the Data Integration Service dynamically divides the underlying data into partitions and processes all of the partitions concurrently.

Optionally, developers can set a maximum parallelism value for a mapping in the Developer tool. By default, the maximum parallelism for each mapping is set to Auto. Each mapping uses the maximum parallelism value defined for the Data Integration Service. Developers can change the maximum parallelism value in the mapping run-time properties to define a maximum value for a particular mapping. When maximum parallelism is set to different integer values for the Data Integration Service and the mapping, the Data Integration Service uses the minimum value.

For more information, see the *Informatica Application Services Guide* and the *Informatica Developer Mapping Guide*.

Partition Optimization

You can optimize the partitioning of Model repository mappings to increase performance. You can add more partitions, select the best performing partition types, use more CPUs, and optimize the source or target database for partitioning.

To optimize partitioning, perform the following tasks:

Increase the number of partitions.

When you configure Model repository mappings, you increase the number of partitions when you increase the maximum parallelism value for the Data Integration Service or the mapping.

Increase the number of partitions to enable the Data Integration Service to create multiple connections to sources and process partitions of source data concurrently. Increasing the number of partitions increases the number of threads, which also increases the load on the Data Integration Service nodes. If

the Data Integration Service node or nodes contain ample CPU bandwidth, processing rows of data concurrently can increase performance.

Note: If you use a single-node Data Integration Service and the Data Integration Service uses a large number of partitions in a session or mapping that processes large amounts of data, you can overload the system.

Use multiple CPUs.

If you have a symmetric multi-processing (SMP) platform, you can use multiple CPUs to concurrently process partitions of data.

Optimize the source database for partitioning.

You can optimize the source database for partitioning. For example, you can tune the database, enable parallel queries, separate data into different tablespaces, and group sorted data.

Optimize the target database for partitioning.

You can optimize the target database for partitioning. For example, you can enable parallel inserts into the database, separate data into different tablespaces, and increase the maximum number of sessions allowed to the database.

High Availability

High availability eliminates a single point of failure in an Informatica domain and provides minimal service interruption in the event of failure. When you configure high availability for a domain, the domain can continue running despite temporary network, hardware, or service failures. You can configure high availability for the domain, application services, and application clients.

The following high availability components make services highly available in an Informatica domain:

- **Resilience.** An Informatica domain can tolerate temporary connection failures until either the resilience timeout expires or the failure is fixed.
- **Restart and failover.** A process can restart on the same node or on a backup node after the process becomes unavailable.
- **Recovery.** Operations can complete after a service is interrupted. After a service process restarts or fails over, it restores the service state and recovers operations.

When you plan a highly available Informatica environment, consider the differences between internal Informatica components and systems that are external to Informatica. Internal components include the Service Manager, application services, and command line programs. External systems include the network, hardware, database management systems, FTP servers, message queues, and shared storage.

High availability features for the Informatica environment are available based on your license.

CHAPTER 3

Sources

This chapter includes the following topics:

- [Overview of Sources, 59](#)
- [PowerExchange Adapter Sources, 59](#)
- [Sources on Databricks, 60](#)
- [File Sources on Hadoop, 62](#)
- [Relational Sources on Hadoop, 66](#)
- [Hive Sources on Hadoop, 67](#)
- [Sqoop Sources on Hadoop, 70](#)

Overview of Sources

You can push a mapping to a non-native environment that includes a source from the native environment or non-native environment. Some sources have restrictions when you reference them in a non-native environment.

Access to sources in a non-native environment require a PowerExchange adapter.

When a mapping runs in the Hadoop environment, an HDFS source cannot reside on a remote cluster.

PowerExchange Adapter Sources

PowerExchange sources can run in the Hadoop or Databricks environment, based on the adapter.

The following table lists PowerExchange sources and the non-native engines that support the sources:

PowerExchange Adapter	Supported Engines
Amazon Redshift	- Blaze - Spark
Amazon S3	- Blaze - Spark

PowerExchange Adapter	Supported Engines
Google Analytics	- Spark
Google BigQuery	- Spark
Google Cloud Spanner	- Spark
HBase	- Blaze - Spark
HDFS	- Blaze - Spark
Hive	- Blaze - Spark
MapR-DB	- Blaze - Spark
Microsoft Azure Blob Storage	- Spark - Databricks Spark
Microsoft Azure Cosmos DB	- Spark - Databricks Spark
Microsoft Azure Data Lake Store	- Spark - Databricks Spark
Microsoft Azure SQL Data Warehouse	- Spark - Databricks Spark
Snowflake	- Spark
Teradata Parallel Transporter API	- Blaze - Spark

Sources on Databricks

A mapping that runs in the Databricks environment can include file and database sources.

The following table lists the storage types and sources:

Storage Type	Source
File	- Microsoft Azure Blob Storage - Microsoft Azure Data Lake Store
Database	- Microsoft Azure Cosmos DB - Microsoft Azure SQL Data Warehouse

Complex File Sources on ADLS

Use a PowerExchange for Microsoft Azure Data Lake Store connection to read data from ADLS data objects.

The following table shows the complex files that a mapping can process within ADLS storage in the Databricks environment:

File Type	Format
Avro	- Flat - Hierarchical
JSON	- Flat - Hierarchical
Parquet	- Flat - Hierarchical
Note: The read operation for hierarchical files must be enabled to project columns as complex data type.	

Complex File Sources on Azure Blob

Use a PowerExchange for Microsoft Azure Blob Storage connection to read data from Blob Storage data objects. You can read compressed binary files.

The following table shows the complex files that a mapping can process within Azure Blob storage in the Databricks environment:

File Type	Format
Avro	- Flat - Hierarchical
JSON	- Flat - Hierarchical
Parquet	- Flat - Hierarchical
Note: The read operation for hierarchical files must be enabled to project columns as complex data type.	

Rules and Guidelines for Databricks Sources

File Sources

Consider the following general rules and guidelines for file sources:

- The flat file source must reside on Microsoft Azure Blob Storage or Microsoft Azure Data Lake Store.
- The row delimiter must be /n.
- The file cannot be fixed width.
- Multiple column delimiters are not supported.
- To read multiline data, set the text qualifier to single or double quotes and enclose the data in the quoted qualifier.

- Empty values only are treated as null values.

Null Processing

Consider the following rules and guidelines for null processing:

Unexpected values converted to nulls

The Databricks Spark engine generates null values for all fields in the same record if any field contains an unexpected value based on the following scenarios:

- Any type mismatch occurs, such as passing string data to a numeric column.
- Data is out of bounds, such as with bigint or int data types.

Consider using a Filter transformation to filter out null rows.

Date/time values converted to nulls

When the Databricks Spark engine reads date/time values, it uses the format configured in the Mapping properties for the run-time preferences of the Developer tool. If the date format read from the source does not match the format configured in the Developer tool, the Databricks Spark engine converts the date values to nulls. The default value in the Developer tool is MM/DD/YYYY HH24:MI:SS.

Double and Decimal Conversions

When the Databricks Spark engine reads from an Azure source, it converts double and decimal data types to scientific notation. When it converts that data back to a double or decimal to write to the target, it drops precision greater than 15 and maintains precision of 15 digits.

File Sources on Hadoop

A mapping that runs in the Hadoop environment can process complex files and flat files.

To read large volumes of data, you can connect a complex file source to read data from a directory of files that have the same format and properties. You can read compressed binary files.

You can read complex files from the following storage types in the Hadoop environment:

- Amazon Simple Storage Service (Amazon S3)
- Hadoop Distributed File System (HDFS)
- MapR File System (MapR-FS)
- Microsoft Azure Blob Storage (Azure Blob Storage)
- Microsoft Azure Data Lake Store (ADLS)

Complex File Sources on Amazon S3

Use a PowerExchange for Amazon S3 connection to read data from Amazon S3 data objects.

The following table shows the complex files that a mapping can process within Amazon S3 storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

Complex File Sources on ADLS

You can use a PowerExchange for HDFS or a PowerExchange for Microsoft Azure Data Lake Store connection to read data from ADLS data objects. If you use a PowerExchange for Microsoft Azure Data Lake Store connection, you cannot read compressed binary files from ADLS.

The following table shows the complex files that a PowerExchange for Microsoft Azure Data Lake Store connection can process within ADLS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ¹	Spark
JSON	- Flat - Hierarchical ¹	Spark
Parquet	- Flat - Hierarchical ¹	Spark
¹ To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

The following table shows the complex files that a PowerExchange for HDFS connection can process within ADLS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark

File Type	Supported Formats	Supported Engines
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

Complex File Sources on Azure Blob

You can use a PowerExchange for HDFS or a PowerExchange for Microsoft Azure Blob Storage connection to read data from Azure Blob data objects.

The following table shows the complex files that a mapping can process within Azure Blob storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ¹	Spark
JSON	- Flat - Hierarchical ¹	Spark
Parquet	- Flat - Hierarchical ¹	Spark
¹ To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

The following table shows the complex files that a PowerExchange for HDFS connection can process within Azure Blob Storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

Complex File Sources on MapR-FS

Use a PowerExchange for HDFS connection to read data from MapR-FS data objects.

The following table shows the complex files that a mapping can process within MapR-FS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

Complex File Sources on HDFS

Use a PowerExchange for HDFS connection to read data from HDFS data objects.

The following table shows the complex files that a mapping can process within HDFS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

Flat File Sources on Hadoop

A mapping that is running in a Hadoop environment can read a flat file source from a native environment.

Consider the following limitations when you configure the mapping to read a flat file source:

- You cannot use an indirect source type.

- The row size in a flat file source cannot exceed 190 MB.
- You cannot use a command to generate or to transform flat file data and send the output to the flat file reader at run time.

Generate the Source File Name

You can generate the source file name for the flat file data object. The content of the file name column remains consistent across different modes of execution.

When you push processing to the specific engine for the required file types, the file name column returns the path based on the following formats:

Run-time Engine	Source File Location	Returned Path
Blaze	Local system	<code><staged path><local file path></code> For example, <code>hdfs://host name:port/hive/warehouse/home/devbld/Desktop/ff.txt</code>
Spark	HDFS	<code>hdfs://<host name>:<port>/<file name path></code> For example, <code>hdfs://host name:port/hive/warehouse/ff.txt</code>
Spark	Local system	<code><local file path></code> For example, <code>/home/devbld/Desktop/ff.txt</code>

The file name column returns the content in the following format for a high availability cluster: `hdfs://<host name>/<file name path>`

For example, `hdfs://irlldv:5008/hive/warehouse/ff.txt`

Relational Sources on Hadoop

You can run mappings with relational sources on the Blaze engine.

The Data Integration Service does not run pre-mapping SQL commands or post-mapping SQL commands against relational sources. You cannot validate and run a mapping with PreSQL or PostSQL properties for a relational source in a Hadoop environment.

The Data Integration Service can use multiple partitions to read from the following relational sources:

- IBM DB2
- Oracle

Hive Sources on Hadoop

You can include Hive sources in an Informatica mapping that runs in the Hadoop environment.

Consider the following limitations when you configure a Hive source in a mapping that runs in the Hadoop environment:

- A mapping fails to run when you have Unicode characters in a Hive source definition.
- The Hive source cannot reside on a remote cluster. A remote cluster is a cluster that is remote from the machine that the Hadoop connection references in the mapping.
- The third-party Hive JDBC driver does not return the correct precision and scale values for the Decimal data type. As a result, when you import Hive tables with a Decimal data type into the Developer tool, the Decimal data type precision is set to 38 and the scale is set to 0. Consider the following configuration rules and guidelines based on the version of Hive:
 - Hive 0.11. Accept the default precision and scale for the Decimal data type in the Developer tool.
 - Hive 0.12. Accept the default precision and scale for the Decimal data type in the Developer tool.
 - Hive 0.12 with Cloudera CDH 5.0. You can configure the precision and scale fields for source columns with the Decimal data type in the Developer tool.
 - Hive 0.13 and above. You can configure the precision and scale fields for source columns with the Decimal data type in the Developer tool.
 - Hive 0.14 or above. The precision and scale used for the Decimal data type in the Hive database also appears in the Developer tool.

A mapping that runs on the Spark engine can have partitioned Hive source tables and bucketed sources.

PreSQL and PostSQL Commands

You can create SQL commands for Hive sources. You can execute the SQL commands to execute SQL statements such as insert, update, and delete on the Hive source.

PreSQL is an SQL command that runs against the Hive source before the mapping reads from the source.

PostSQL is an SQL command that runs against the Hive source after the mapping writes to the target.

You can use PreSQL and PostSQL on the Spark engine. The Data Integration Service does not validate PreSQL or PostSQL commands for a Hive source.

Note: You can manually validate the SQL by running the following query in a Hive command line utility:

```
CREATE VIEW <table name> (<port list>) AS <SQL>
```

where:

- <table name> is a name of your choice
- <port list> is the comma-delimited list of ports in the source
- <SQL> is the query to validate

Pre-Mapping SQL Commands

PreSQL is an SQL command that runs against a Hive source before the mapping reads from the source.

For example, you might use a Hive source in a mapping. The data stored in the Hive source changes regularly and you must update the data in the Hive source before the mapping reads from the source to make sure that the mapping reads the latest records. To update the Hive source, you can configure a PreSQL command.

Post-Mapping SQL Commands

PostSQL is an SQL command that runs against a Hive source after the mapping writes to the target.

For example, you might use a Hive source in a mapping. After the mapping writes to a target, you might want to delete the stage records stored in the Hive source. You want to run the command only after the mapping writes the data to the target to make sure that the data is not removed prematurely from the Hive source. To delete the records in the Hive source table after the mapping writes to the target, you can configure a PostSQL command.

Rules and Guidelines for Pre- and Post-Mapping SQL Commands

Consider the following restrictions when you run PreSQL and PostSQL commands against Hive sources:

- When you create an SQL override on a Hive source, you must enclose keywords or special characters in backtick (``) characters.
- When you run a mapping with a Hive source in the Hadoop environment, references to a local path in pre-mapping SQL commands are relative to the Data Integration Service node. When you run a mapping with a Hive source in the native environment, references to local path in pre-mapping SQL commands are relative to the Hive server node.

Rules and Guidelines for Hive Sources on the Blaze Engine

You can include Hive sources in an Informatica mapping that runs on the Blaze engine.

Consider the following rules and guidelines when you configure a Hive source in a mapping that runs on the Blaze engine:

- Hive sources for a Blaze mapping include the TEXT, Sequence, Avro, RCfile, ORC, and Parquet storage formats.
- A mapping that runs on the Blaze engine can have bucketed
- Hive sources and Hive ACID tables.
- Hive ACID tables must be bucketed.
- The Blaze engine supports Hive tables that are enabled for locking.
- Hive sources can contain quoted identifiers in Hive table names, column names, and schema names.
- The TEXT storage format in a Hive source for a Blaze mapping can support ASCII characters as column delimiters and the newline characters as a row separator. You cannot use hex values of ASCII characters. For example, use a semicolon (;) instead of 3B.
- You can define an SQL override in the Hive source for a Blaze mapping.
- The Blaze engine can read from an RCFile as a Hive source. To read from an RCFile table, you must create the table with the `SerDe` clause.
- The Blaze engine can read from Hive tables that are compressed. To read from a compressed Hive table, you must set the `TBLPROPERTIES` clause.

RCFile as Hive Tables

The Blaze engine can read and write to RCFile as Hive tables. However, the Blaze engine supports only the ColumnarSerDe `SerDe`. In Hortonworks, the default `SerDe` for an RCFile is `LazyBinaryColumnarSerDe`. To read and write to an RCFile table, you must create the table by specifying the `SerDe` as

```
org.apache.hadoop.hive.serde2.columnar.ColumnarSerDe.
```

For example:

```
CREATE TABLE TEST_RCFile
(id int, name string)
ROW FORMAT SERDE
'org.apache.hadoop.hive.serde2.columnar.ColumnarSerDe' STORED AS RCFILE;
```

You can also set the default RCFile SerDe from the Ambari or Cloudera manager. Set the property `hive.default.rcfile.serde` to `org.apache.hadoop.hive.serde2.columnar.ColumnarSerDe`.

Compressed Hive Tables

The Blaze engine can read and write to Hive tables that are compressed. However, to read from a compressed Hive table or write to a Hive table in compressed format, you must set the `TBLPROPERTIES` clause as follows:

- When you create the table, set the table properties:

```
TBLPROPERTIES ('property_name'='property_value')
```

- If the table already exists, alter the table to set the table properties:

```
ALTER TABLE table_name SET TBLPROPERTIES ('property_name' = 'property_value');
```

The property name and value are not case sensitive. Depending on the file format, the table property can take different values.

The following table lists the property names and values for different file formats:

File Format	Table Property Name	Table Property Values
Avro	avro.compression	BZIP2, deflate, Snappy
ORC	orc.compress	Snappy, ZLIB
Parquet	parquet.compression	GZIP, Snappy
RCFile	rcfile.compression	Snappy, ZLIB
Sequence	sequencefile.compression	BZIP2, GZIP, LZ4, Snappy
Text	text.compression	BZIP2, GZIP, LZ4, Snappy

Note: The Blaze engine does not write data in the default ZLIB compressed format when it writes to a Hive target stored as ORC format. To write in a compressed format, alter the table to set the `TBLPROPERTIES` clause to use ZLIB or Snappy compression for the ORC file format.

The following text shows sample commands to create table and alter table:

- Create table:

```
create table CBO_3T_JOINS_CUSTOMER_HIVE_SEQ_GZIP
(C_CUSTKEY DECIMAL(38,0), C_NAME STRING, C_ADDRESS STRING,
C_PHONE STRING, C_ACCTBAL DECIMAL(10,2),
C_MKTSEGMENT VARCHAR(10), C_COMMENT VARCHAR(117))
partitioned by (C_NATIONKEY DECIMAL(38,0))
TBLPROPERTIES ('sequencefile.compression'='gzip')
stored as SEQUENCEFILE;
```

- Alter table:

```
ALTER TABLE table_name
SET TBLPROPERTIES ('avro.compression'='BZIP2');
```

Sqoop Sources on Hadoop

Sqoop sources are valid in mappings in a Hadoop environment.

You can include a JDBC-compliant database as a Sqoop source in an Informatica mapping that runs in a Hadoop environment.

For example, you can include the following sources in a Sqoop mapping:

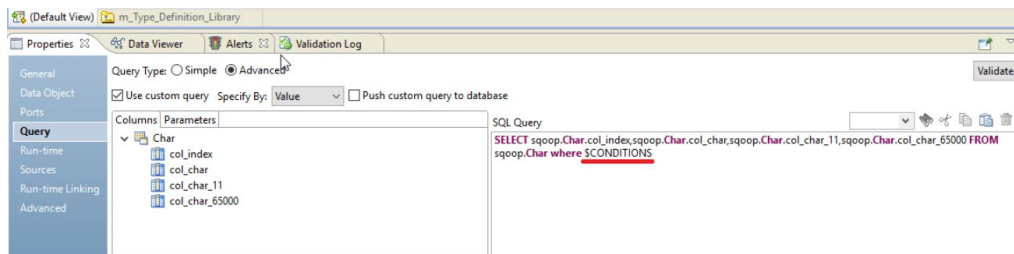
- Aurora
- Greenplum
- IBM DB2
- IBM DB2 for z/OS
- Microsoft SQL Server
- Netezza
- Oracle
- Teradata
- Vertica

Reading Data from Vertica Sources through Sqoop

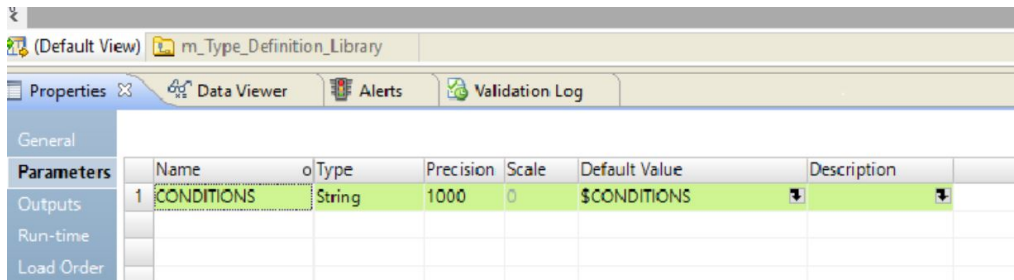
To read Vertica data through Sqoop connectivity, Sqoop requires a custom query.

Create the custom query with the `$CONDITIONS` parameter in the WHERE clause. Then, define a mapping parameter with the name as `CONDITIONS` and set its value to `$CONDITIONS`. At run time, Sqoop replaces `$CONDITIONS` with a unique condition expression.

1. Open the Sqoop data object and click the **Query** tab in the **Properties** view.
2. Define a custom query and include `$CONDITIONS` in the WHERE clause of the custom query as shown in the following image:



3. Open the Sqoop mapping and click the **Parameters** tab in the **Properties** view.
4. Define a mapping parameter with the name as `CONDITIONS` and set its value to `$CONDITIONS` as shown in the following image:



Rules and Guidelines for Sqoop Sources

Consider the following rules and guidelines when you configure a Sqoop source in a mapping:

- If you specify a sort condition in a mapping, the Data Integration Service ignores the Order By condition.
- You cannot sort columns in a Sqoop source.
- You cannot read distinct rows from a Sqoop source.
- When you read data from an Oracle source through Sqoop and run the mapping on the Blaze or Spark engine, Sqoop treats the owner name as case sensitive.
- Sqoop uses the values that you configure in the **User Name** and **Password** fields of the JDBC connection. If you configure the --username or --password argument in a JDBC connection or mapping, Sqoop ignores the arguments. If you create a password file to access a database, Sqoop ignores the password file.

Rules and Guidelines for Sqoop Queries

Consider the following rules and guidelines when you configure a Sqoop query in a mapping:

- To override the default query in a mapping with an advanced query, you must define a mapping parameter with the name as `CONDITIONS` and set its value to `$CONDITIONS`. You must then include `$CONDITIONS` in the WHERE clause of the custom query. For more information about configuring the custom query, see the following Informatica Knowledge Base article: [KB 564550](#)
- If you define a custom query, you must verify that the metadata of the custom query matches the metadata of the source object. Otherwise, Sqoop might write blank values to the target.
- When you enable OraOop and configure an advanced query to read data from an Oracle source through Sqoop, the mapping fails on the Spark engine.

CHAPTER 4

Targets

This chapter includes the following topics:

- [Overview of Targets, 72](#)
- [PowerExchange Adapter Targets, 72](#)
- [Targets on Databricks, 73](#)
- [File Targets on Hadoop, 75](#)
- [Message Targets on Hadoop, 79](#)
- [Relational Targets on Hadoop, 80](#)
- [Hive Targets on Hadoop, 80](#)
- [Sqoop Targets on Hadoop, 84](#)

Overview of Targets

You can push a mapping to a non-native environment that includes a target from the native environment or non-native environment. Some targets have restrictions when you reference them in a non-native environment.

Access to targets in a non-native environment might require a PowerExchange adapter.

PowerExchange Adapter Targets

PowerExchange targets can run in the Hadoop or Databricks environment, based on the adapter.

The following table lists PowerExchange targets and the non-native engines that support the targets:

PowerExchange Adapter	Supported Engines
Amazon Redshift	- Blaze - Spark
Amazon S3	- Blaze - Spark
Google Analytics	- Spark

PowerExchange Adapter	Supported Engines
Google BigQuery	- Spark
Google Cloud Spanner	- Spark
HBase	- Blaze - Spark
HDFS	- Blaze - Spark
Hive	- Blaze - Spark
MapR-DB	- Blaze - Spark
Microsoft Azure Blob Storage	- Spark - Databricks Spark
Microsoft Azure Cosmos DB	- Spark - Databricks Spark
Microsoft Azure Data Lake Store	- Blaze - Spark - Databricks Spark
Microsoft Azure SQL Data Warehouse	- Spark - Databricks Spark
Snowflake	- Spark
Teradata Parallel Transporter API	- Blaze - Spark

Targets on Databricks

A mapping that runs in the Databricks environment can include file and database targets.

The following table lists the storage types and targets:

Storage Type	Target
File	- Microsoft Azure Blob Storage - Microsoft Azure Data Lake Store
Database	- Microsoft Azure Cosmos DB - Microsoft Azure SQL Data Warehouse

Complex File Targets on ADLS

Use a PowerExchange for Microsoft Azure Data Lake Store connection to write data to ADLS data objects.

The following table shows the complex files that a mapping can process within ADLS storage in the Databricks environment:

File Type	Format
Avro	- Flat - Hierarchical
JSON	- Flat - Hierarchical
Parquet	- Flat - Hierarchical
Note: The write operation for hierarchical files must be enabled to project columns as complex data type.	

Complex File Targets on Azure Blob

Use a PowerExchange for Microsoft Azure Blob Storage connection to write data to Blob Storage data objects. You can write compressed binary files.

The following table shows the complex files that a mapping can process within Azure Blob storage in the Databricks environment:

File Type	Format
Avro	- Flat - Hierarchical
JSON	- Flat - Hierarchical
Parquet	- Flat - Hierarchical
Note: The write operation for hierarchical files must be enabled to project columns as complex data type.	

Rules and Guidelines for Databricks Targets

File Targets

Consider the following general rules and guidelines for file targets:

- The file targets must reside on Microsoft Azure Blob Storage or Microsoft Azure Data Lake Store.
- The row delimiter must be /n.
- The file cannot be fixed width.
- Multiple column delimiters are not supported.
- Empty values only are treated as null values.

Null Processing

Consider the following rules and guidelines for null processing:

Null value conversions

When the Databricks Spark engine writes to a target, it converts null values to empty strings (" "). For example, 12, AB,"",23p09udj.

Consider the null processing behavior based on the following Databricks targets:

- File target. The Databricks Spark engine writes all null values as empty strings.
- Azure SQL Data Warehouse. The Databricks Spark engine can write the empty strings to string columns, but when it tries to write an empty string to a non-string column, the mapping fails with a type mismatch.

To allow the Databricks Spark engine to convert the empty strings back to null values and write to the target, configure the following advanced property in the Databricks Spark connection:
`infaspark.flatfile.writer.nullValue=true`

Unexpected values converted to nulls

The Databricks Spark engine generates null values for all fields in the same record if any field contains an unexpected value based on the following scenarios:

- Any type mismatch occurs, such as passing string data to a numeric column.
- Data is out of bounds, such as with bigint or int data types.

Consider using a Filter transformation to filter out null rows.

Date/time values converted to nulls

When the Databricks Spark engine writes date/time values, it uses the format configured in the Mapping properties for the run-time preferences of the Developer tool. If the date format that passes to the target does not match the format configured in the Developer tool, the Databricks Spark engine writes null values. The default value in the Developer tool is MM/DD/YYYY HH24:MI:SS.

Double and Decimal Conversions

When the Databricks Spark engine reads from an Azure source, it converts double and decimal data types to scientific notation. When it converts that data back to a double or decimal to write to the target, it drops precision greater than 15 and maintains precision of 15 digits.

File Targets on Hadoop

A mapping that runs in the Hadoop environment can process complex files and flat files.

To write large volumes of data, you can connect a complex file target to write data to a directory of files that have the same format and properties. You can read compressed binary files.

You can write to complex files in the following storage types in the Hadoop environment:

- Amazon Simple Storage Service (Amazon S3)
- Hadoop Distributed File System (HDFS)
- MapR File System (MapR-FS)
- Microsoft Azure Blob Storage (Azure Blob Storage)
- Microsoft Azure Data Lake Store (ADLS)

Complex File Targets on Amazon S3

Use a PowerExchange for Amazon S3 connection to write data to Amazon S3 data objects.

The following table shows the complex files that a mapping can process within Amazon S3 storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file write operation must be enabled to project columns as complex data type.		

Complex File Targets on ADLS

You can use a PowerExchange for HDFS or a PowerExchange for Microsoft Azure Data Lake Store connection to write data to ADLS data objects. If you use a PowerExchange for Microsoft Azure Data Lake Store connection, you cannot write compressed binary files to ADLS.

The following table shows the complex files that a PowerExchange for Microsoft Azure Data Lake Store connection can process within ADLS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ¹	Spark
JSON	- Flat - Hierarchical ¹	Spark
Parquet	- Flat - Hierarchical ¹	Spark
¹ To run on the Spark engine, the complex file write operation must be enabled to project columns as complex data type.		

The following table shows the complex files that a PowerExchange for HDFS connection can process within ADLS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file write operation must be enabled to project columns as complex data type.		

Complex File Targets on Azure Blob

You can use a PowerExchange for HDFS or a PowerExchange for Microsoft Azure Blob Storage connection to write data to Azure Blob data objects.

The following table shows the complex files that a mapping can process within Azure Blob storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ¹	Spark
JSON	- Flat - Hierarchical ¹	Spark
Parquet	- Flat - Hierarchical ¹	Spark
¹ To run on the Spark engine, the complex file write operation must be enabled to project columns as complex data type.		

The following table shows the complex files that a PowerExchange for HDFS connection can process within Azure Blob Storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark

File Type	Supported Formats	Supported Engines
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file write operation must be enabled to project columns as complex data type.		

Complex File Targets on MapR-FS

Use a PowerExchange for HDFS connection to read data from MapR-FS data objects.

The following table shows the complex files that a mapping can process within MapR-FS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file read operation must be enabled to project columns as complex data type.		

Complex File Targets on HDFS

Use a PowerExchange for HDFS connection to write data to HDFS data objects.

The following table shows the complex files that a mapping can process within HDFS storage in the Hadoop environment:

File Type	Supported Formats	Supported Engines
Avro	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
JSON	- Flat ¹ - Hierarchical ^{1 2}	- Blaze - Spark
ORC	- Flat	- Spark

File Type	Supported Formats	Supported Engines
Parquet	- Flat - Hierarchical ^{1 2}	- Blaze - Spark
¹ To run on the Blaze engine, the complex file data object must be connected to a Data Processor transformation. ² To run on the Spark engine, the complex file write operation must be enabled to project columns as complex data type.		

Flat File Targets on Hadoop

Consider the following rules and guidelines for flat file targets in the Hadoop environment:

- A mapping that runs in the Hadoop environment can write to a flat file target in the native environment.
- The Data Integration Service truncates the target files and reject files before writing the data. When you use a flat file target, you cannot append output data to target files and reject files.
- The Data Integration Service can write to a file output for a flat file target. When you have a flat file target in a mapping, you cannot write data to a command.

Consider the following rules and guidelines for HDFS flat file targets:

- The Data Integration Service truncates the target files and reject files before writing the data. To append output data to HDFS target files and reject files, choose to append data if the HDFS target exists.
Data is appended to reject files only if the reject file directory is on the Data Integration Service machine. If the directory is in the Hadoop environment, rejected rows are overwritten.
- When you choose to append data if the HDFS target exists, the Data Integration Service appends the mapping execution ID to the names of the target files and reject files.
- When you use a HDFS flat file target in a mapping, you must specify the full path that includes the output file directory and file name. The Data Integration Service might generate multiple output files in the output directory when you run the mapping in a Hadoop environment.
- An HDFS target cannot reside on a remote cluster. A remote cluster is a cluster that is remote from the machine that the Hadoop connection references in the mapping.

Message Targets on Hadoop

You can write to Kafka messaging targets.

Apache Kafka is a publish-subscribe messaging system. Kafka runs as a cluster comprised of one or more servers each of which is called a broker. Kafka brokers process data in the form of messages, publishes the messages to a topic, subscribes the messages from a topic, and then writes it to the Kafka target.

You can use the Kafka data object write operation as a target. If you want to configure high availability for the mapping, ensure that the Kafka cluster is highly available.

For more information about Kafka clusters, Kafka brokers, and partitions, see <https://kafka.apache.org/documentation>.

Technical Preview: Support to write the batch data to a Kafka target in Big Data Management is available for technical preview. Technical preview functionality is supported but is unwarranted and is not production-ready. Informatica recommends that you use in non-production environments only.

Relational Targets on Hadoop

You can run mappings with relational targets on the Blaze engine.

The Data Integration Service does not run pre-mapping SQL commands or post-mapping SQL commands against relational targets in a Hadoop environment. You cannot validate and run a mapping with PreSQL or PostSQL properties for a relational target in a Hadoop environment.

The Data Integration Service can use multiple partitions to write to the following relational targets:

- IBM DB2
- Oracle

Hive Targets on Hadoop

A mapping that is running in the Hadoop environment can write to a Hive target.

A Hive target can be an internal table or an external table. Internal Hive tables are managed by Hive and are also known as managed tables. External Hive tables are managed by an external source such as HDFS, Amazon S3, or Microsoft Azure Blob Storage.

Consider the following restrictions when you configure a Hive target in a mapping that runs in the Hadoop environment:

- A mapping fails to run when you use Unicode characters in a Hive target definition.
- The Hive target cannot reside on a remote cluster. A remote cluster is a cluster that is remote from the machine that the Hadoop connection references in the mapping.
- When you set up a dynamic target for a partitioned Hive table, the value used for the partition is the final column in the table. If the table has a dynamic partition column, the final column of the table is the dynamic partition column. To use a different column for the partition, move it to the last column of the table. If the table has multiple partition columns, the dynamic partition values are selected from the last columns of the upstream transformation. You can use an Expression transformation to reorder the columns if necessary.

When a mapping creates or replaces a Hive table, the type of table that the mapping creates depends on the run-time engine that you use to run the mapping.

The following table shows the table type for each run-time engine:

Run-Time Engine	Resulting Table Type
Blaze	MANAGED_TABLE
Spark	EXTERNAL_TABLE

You can design a mapping to truncate an internal or external Hive table that is bucketed and partitioned.

In a mapping that runs on the Spark engine or the Blaze engine, you can create a custom DDL query that creates or replaces a Hive table at run time. However, with the Blaze engine, you cannot use a backtick (') character in the DDL query. The backtick character is required in HiveQL when you include special characters or keywords in a query.

The Spark engine can write to bucketed Hive targets. Bucketing and partitioning of Hive tables can improve performance by reducing data shuffling and sorting.

PreSQL and PostSQL Commands

You can create SQL commands for Hive targets. You can execute the SQL commands to execute SQL statements such as insert, update, and delete on the Hive target.

PreSQL is an SQL command that runs against the Hive target before the mapping reads from the source. PostSQL is an SQL command that runs against the Hive target after the mapping writes to the target.

You can use PreSQL and PostSQL on the Spark engine. The Data Integration Service does not validate PreSQL or PostSQL commands for a Hive target.

Pre-Mapping SQL Commands

PreSQL is an SQL command that runs against a Hive target before the mapping reads from a source.

For example, you might use a Hive target in a mapping. The data stored in the Hive target contains old records from the previous day and you must delete the old records in the Hive target before you run the mapping that writes new records to the target. To delete the old records in the Hive target, you can configure a PreSQL command.

Post-Mapping SQL Commands

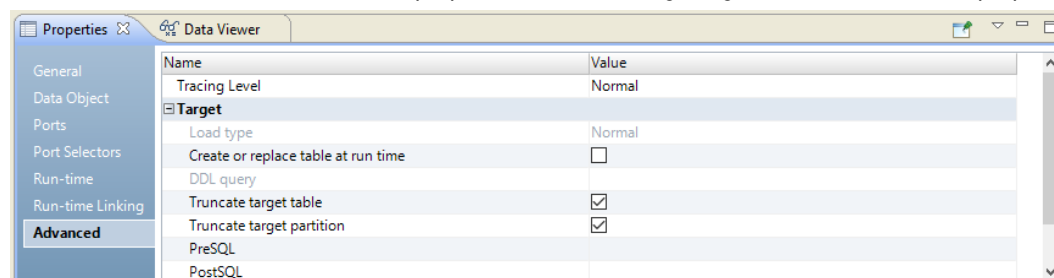
PostSQL is an SQL command that runs against a Hive source after the mapping writes to the target.

For example, you might use a Hive target in a mapping that runs in a test environment. After the mapping writes to the target, you might want to delete the records stored in the Hive target so that you can test the same Hive target again in the next mapping run. To delete the records in the Hive target, you can configure a PostSQL command.

Truncating Hive Targets

Truncate Hive target tables to delete the table contents. You can truncate internal and external Hive tables in the Hadoop environment.

Truncate a Hive table in the Hive table properties. The following image shows the Hive table properties:



To truncate the entire Hive table, choose the option to truncate the target table.

To truncate only the partitions in the Hive table for which the transformation received input data, you must choose to truncate the target table and to truncate the partition in the Hive target table.

Note: You must choose both "Truncate target table" and "Truncate target partition" options due to the possibility that input data might affect all partitions and cause the whole table to be truncated.

Consider the following restrictions when you truncate a Hive table in a mapping that runs in the Hadoop environment:

- The Data Integration Service can truncate the partition in the Hive target in which the data is being inserted.
- You must truncate the target table to overwrite data to a Hive table with Hive version 0.7. The Data Integration Service ignores write, update override, delete, insert, and update strategy properties when it writes data to a Hive target.

Updating Hive Targets with an Update Strategy Transformation

For mappings that run on the Spark engine, you can use Hive MERGE statements to perform Update Strategy tasks. When a query uses a MERGE statement instead of INSERT, UPDATE or DELETE statements, processing is more efficient.

To use Hive MERGE, select **true** for the option in the Advanced Properties of the Update Strategy transformation.

The mapping ignores the Hive MERGE option and the Data Integration Service uses INSERT, UPDATE and DELETE to perform the operation under the following scenarios:

- The mapping runs on the Blaze engine.
- Hive MERGE is restricted on the Hadoop distribution.

The mapping log contains results of the operation, including whether restrictions affected results.

When the update affects partitioning or bucketing columns, updates to the columns are omitted.

Note: The Developer tool and the Data Integration Service do not validate against this restriction. If the Update Strategy expression violates these restrictions, the mapping might produce unexpected results.

Rules and Guidelines for Hive Targets on the Blaze Engine

You can include Hive targets in an Informatica mapping that runs on the Blaze engine.

Consider the following rules and guidelines when you configure a Hive target in a mapping that runs on the Blaze engine:

- A mapping that runs on the Blaze engine can have partitioned and bucketed Hive tables as targets. However, if you append data to a bucketed table, the Blaze engine overwrites the data in the bucketed target.
- Mappings that run on the Blaze engine can read and write to sorted targets.
- The Blaze engine supports Hive tables that are enabled for locking.
- The Blaze engine can create or replace Hive target tables.
- A mapping that runs on the Blaze engine can write to Hive ACID tables. To write to a Hive ACID table, the mapping must contain an Update Strategy transformation connected to the Hive target. The update strategy expression must flag each row for insert.
- The Blaze engine can write to Hive tables that are compressed. To write to a Hive table in compressed format, you must set the `TBLPROPERTIES` clause.
- For a mapping that writes to a Hive target on the Blaze engine, you cannot use update or delete update strategy properties.

RCFile as Hive Tables

The Blaze engine can read and write to RCFile as Hive tables. However, the Blaze engine supports only the ColumnarSerDe SerDe. In Hortonworks, the default SerDe for an RCFile is LazyBinaryColumnarSerDe. To read and write to an RCFile table, you must create the table by specifying the SerDe as `org.apache.hadoop.hive.serde2.columnar.ColumnarSerDe`.

For example:

```
CREATE TABLE TEST_RCFile
(id int, name string)
ROW FORMAT SERDE
'org.apache.hadoop.hive.serde2.columnar.ColumnarSerDe' STORED AS RCFILE;
```

You can also set the default RCFile SerDe from the Ambari or Cloudera manager. Set the property `hive.default.rcfile.serde` to `org.apache.hadoop.hive.serde2.columnar.ColumnarSerDe`.

Compressed Hive Tables

The Blaze engine can read and write to Hive tables that are compressed. However, to read from a compressed Hive table or write to a Hive table in compressed format, you must set the `TBLPROPERTIES` clause as follows:

- When you create the table, set the table properties:
`TBLPROPERTIES ('property_name'='property_value')`
- If the table already exists, alter the table to set the table properties:
`ALTER TABLE table_name SET TBLPROPERTIES ('property_name' = 'property_value');`

The property name and value are not case sensitive. Depending on the file format, the table property can take different values.

The following table lists the property names and values for different file formats:

File Format	Table Property Name	Table Property Values
Avro	avro.compression	BZIP2, deflate, Snappy
ORC	orc.compress	Snappy, ZLIB
Parquet	parquet.compression	GZIP, Snappy
RCFile	rcfile.compression	Snappy, ZLIB
Sequence	sequencefile.compression	BZIP2, GZIP, LZ4, Snappy
Text	text.compression	BZIP2, GZIP, LZ4, Snappy

Note: The Blaze engine does not write data in the default ZLIB compressed format when it writes to a Hive target stored as ORC format. To write in a compressed format, alter the table to set the `TBLPROPERTIES` clause to use ZLIB or Snappy compression for the ORC file format.

The following text shows sample commands to create table and alter table:

- Create table:

```
create table CBO_3T_JOINS_CUSTOMER_HIVE_SEQ_GZIP
(C_CUSTKEY DECIMAL(38,0), C_NAME STRING, C_ADDRESS STRING,
C_PHONE STRING, C_ACCTBAL DECIMAL(10,2),
C_MKTSEGMENT VARCHAR(10), C_COMMENT VARCHAR(117))
partitioned by (C_NATIONKEY DECIMAL(38,0))
TBLPROPERTIES ('sequencefile.compression'='gzip')
stored as SEQUENCEFILE;
```

- Alter table:

```
ALTER TABLE table_name
SET TBLPROPERTIES('avro.compression'='BZIP2');
```

Sqoop Targets on Hadoop

Sqoop targets are valid in mappings in a Hadoop environment.

You can include a JDBC-compliant database as a Sqoop target in an Informatica mapping that runs in a Hadoop environment:

For example, you can include the following targets in a Sqoop mapping:

- Aurora
- Greenplum
- IBM DB2
- IBM DB2 for z/OS
- Microsoft SQL Server
- Netezza
- Oracle
- Teradata
- Vertica

You can insert data. You cannot update or delete data in a target. If you configure update arguments, Sqoop ignores them.

Rules and Guidelines for Sqoop Targets

Consider the following rules and guidelines when you configure a Sqoop target in a mapping:

- If a column name or table name contains a special character, the Sqoop export process fails.
- If you configure the **Maintain Row Order** property for a Sqoop target, the Data Integration Service ignores the property.
- If a mapping contains a Sqoop source, an Aggregator transformation, and a flat file target, you must disable the **Maintain Row Order** property for the target. Otherwise, the mapping fails.
- When you run a Sqoop mapping on the Blaze engine, verify that you have not deleted any target port from the mapping. Otherwise, the mapping fails.
- When you export null data to a Microsoft SQL Server column that is defined as not null, the Data Integration Service fails the Sqoop mapping on the Blaze engine instead of rejecting and writing the null data to the bad file.
- When you write data to an Oracle target through Sqoop and run the mapping on the Blaze or Spark engine, Sqoop treats the owner name as case sensitive.
- Sqoop uses the values that you configure in the **User Name** and **Password** fields of the JDBC connection. If you configure the `--username` or `--password` argument in a JDBC connection or mapping, Sqoop ignores the arguments. If you create a password file to access a database, Sqoop ignores the password file.
- When you write data to a Vertica target through Sqoop, the `--batch` argument is required.

CHAPTER 5

Transformations

This chapter includes information about transformation support in the non-native environment.

Overview of Transformations

Due to the differences between native environments and non-native environments, only certain transformations are valid or are valid with restrictions in a non-native environment. Some functions, expressions, data types, and variable fields are not valid in a non-native environment.

Consider the following processing differences that can affect whether transformations and transformation behavior are valid or are valid with restrictions in a non-native environment:

- Compute clusters use distributed processing and process data on different nodes. Each node does not have access to the data that is being processed on other nodes. As a result, the run-time engine might not be able to determine the order in which the data originated.
- Each of the run-time engines in the non-native environment can process mapping logic differently.
- Much of the processing behavior for batch mappings on the Spark engine also apply to streaming mappings.

The following table lists transformations and support for different engines in a non-native environment:

Transformation	Supported Engines
<i>Transformations not listed in this table are not supported in a non-native environment.</i>	
Address Validator	- Blaze - Spark
Aggregator	- Blaze - Spark* - Databricks Spark
Case Converter	- Blaze - Spark
Classifier	- Blaze - Spark
Comparison	- Blaze - Spark

Transformation	Supported Engines
Consolidation	- Blaze - Spark
Data Masking	- Blaze - Spark*
Data Processor	- Blaze
Decision	- Blaze - Spark
Expression	- Blaze - Spark* - Databricks Spark
Filter	- Blaze - Spark* - Databricks Spark
Java	- Blaze - Spark*
Joiner	- Blaze - Spark* - Databricks Spark
Key Generator	- Blaze - Spark
Labeler	- Blaze - Spark
Lookup	- Blaze - Spark* - Databricks Spark
Match	- Blaze - Spark
Merge	- Blaze - Spark
Normalizer	- Blaze - Spark* - Databricks Spark
Parser	- Blaze - Spark
Python	- Spark*
Rank	- Blaze - Spark* - Databricks Spark

Transformation	Supported Engines
Router	- Blaze - Spark* - Databricks Spark
Sequence Generator	- Blaze - Spark
Sorter	- Blaze - Spark* - Databricks Spark
Standardizer	- Blaze - Spark
Union	- Blaze - Spark* - Databricks Spark
Update Strategy	- Blaze - Spark
Weighted Average	- Blaze - Spark
Window	- Spark**
* Supported for both batch and streaming mappings. ** Supported for streaming mappings only. For more information, see the Big Data Streaming User Guide.	

Address Validator Transformation in a Non-native Environment

The Address Validator transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Address Validator Transformation on the Blaze Engine

The Address Validator transformation is supported with the following restrictions:

- The Address Validator transformation cannot generate a certification report.
- The Address Validator transformation fails validation when it is configured to run in Interactive mode or Suggestion List mode.

Address Validator Transformation on the Spark Engine

The Address Validator transformation is supported with the following restrictions:

- The Address Validator transformation cannot generate a certification report.
- The Address Validator transformation fails validation when it is configured to run in Interactive mode or Suggestion List mode.

Aggregator Transformation in a Non-native Environment

The Aggregator transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported with restrictions.

Aggregator Transformation on the Blaze Engine

Some processing rules for the Blaze engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- The transformation contains stateful variable ports.
- The transformation contains unsupported functions in an expression.

Aggregate Functions

If you use a port in an expression in the Aggregator transformation but you do not use the port within an aggregate function, the run-time engine might use any row in the port to process the expression.

The row that the run-time engine uses might not be the last row in the port. Processing is distributed, and thus the run-time engine might not be able to determine the last row in the port.

Data Cache Optimization

The data cache for the Aggregator transformation is optimized to use variable length to store binary and string data types that pass through the Aggregator transformation. The optimization is enabled for record sizes up to 8 MB. If the record size is greater than 8 MB, variable length optimization is disabled.

When variable length is used to store data that passes through the Aggregator transformation in the data cache, the Aggregator transformation is optimized to use sorted input and a pass-through Sorter transformation is inserted before the Aggregator transformation in the run-time mapping.

To view the Sorter transformation, view the optimized mapping or view the execution plan in the Blaze validation environment.

During data cache optimization, the data cache and the index cache for the Aggregator transformation are set to Auto. The sorter cache for the Sorter transformation is set to the same size as the data cache for the

Aggregator transformation. To configure the sorter cache, you must configure the size of the data cache for the Aggregator transformation.

Aggregator Transformation on the Spark Engine

Some processing rules for the Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- The transformation contains stateful variable ports.
- The transformation contains unsupported functions in an expression.

Aggregate Functions

If you use a port in an expression in the Aggregator transformation but you do not use the port within an aggregate function, the run-time engine might use any row in the port to process the expression.

The row that the run-time engine uses might not be the last row in the port. Processing is distributed, and thus the run-time engine might not be able to determine the last row in the port.

Data Cache Optimization

You cannot optimize the data cache for the transformation to store data using variable length.

Aggregator Transformation in a Streaming Mapping

Streaming mappings have additional processing rules that do not apply to batch mappings.

Mapping Validation

Mapping validation fails in the following situations:

- A streaming pipeline contains more than one Aggregator transformation.
- A streaming pipeline contains an Aggregator transformation and a Rank transformation.
- An Aggregator transformation is upstream from a Lookup transformation.
- An Aggregator transformation is in the same pipeline as a Lookup transformation with an inequality condition.

Aggregator Transformation on the Databricks Spark Engine

Some processing rules for the Databricks Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- The transformation contains stateful variable ports.
- The transformation contains unsupported functions in an expression.

Aggregate Functions

If you use a port in an expression in the Aggregator transformation but you do not use the port within an aggregate function, the run-time engine might use any row in the port to process the expression.

The row that the run-time engine uses might not be the last row in the port. Processing is distributed, and thus the run-time engine might not be able to determine the last row in the port.

Data Cache Optimization

You cannot optimize the data cache for the transformation to store data using variable length.

Case Converter Transformation in a Non-native Environment

The Case Converter transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Classifier Transformation in a Non-native Environment

The Classifier transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Comparison Transformation in a Non-native Environment

The Comparison transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Consolidation Transformation in a Non-native Environment

The Address Validator transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Consolidation Transformation on the Blaze Engine

The Consolidation transformation is supported with the following restrictions:

- The transformation might process records in a different order in each environment.
- The transformation might identify a different record as the survivor in each environment.

Consolidation Transformation on the Spark Engine

The Consolidation transformation is supported with the following restrictions:

- The transformation might process records in a different order in each environment.
- The transformation might identify a different record as the survivor in each environment.

Data Masking Transformation in a Non-native Environment

The Data Masking transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Not supported.

Deferment Notice: Effective in version 10.2.2, the Data Masking transformation is deferred for streaming mappings and is unavailable. Deferred functionality is not available in the current release. Informatica intends to reinstate it in an upcoming GA release, but might choose not to in accordance with changing market or technical circumstances.

Data Masking Transformation on the Blaze Engine

The Data Masking transformation is supported with the following restrictions.

Mapping validation fails in the following situations:

- The transformation is configured for repeatable expression masking.
- The transformation is configured for unique repeatable substitution masking.

You can use the following masking techniques on the Blaze engine:

Credit Card
Email
Expression
IP Address
Key
Phone
Random
SIN
SSN
Tokenization
URL
Random Substitution
Repeatable Substitution
Dependent with Random Substitution
Dependent with Repeatable Substitution

Data Masking Transformation on the Spark Engine

The Data Masking transformation is supported with the following restrictions.

Mapping validation fails in the following situations:

- The transformation is configured for repeatable expression masking.
- The transformation is configured for unique repeatable substitution masking.

You can use the following masking techniques on the Spark engine:

Credit Card
Email
Expression
IP Address
Key
Phone
Random
SIN
SSN
Tokenization
URL
Random Substitution
Repeatable Substitution

Dependent with Random Substitution

Dependent with Repeatable Substitution

To optimize performance of the Data Masking transformation, configure the following Spark engine configuration properties in the Hadoop connection:

spark.executor.cores

Indicates the number of cores that each executor process uses to run tasklets on the Spark engine.

Set to: `spark.executor.cores=1`

spark.executor.instances

Indicates the number of instances that each executor process uses to run tasklets on the Spark engine.

Set to: `spark.executor.instances=1`

spark.executor.memory

Indicates the amount of memory that each executor process uses to run tasklets on the Spark engine.

Set to: `spark.executor.memory=3G`

Data Masking Transformation in a Streaming Mapping

Deferment Notice: Effective in version 10.2.2, the Data Masking transformation is deferred for streaming mappings and is unavailable. Deferred functionality is not available in the current release. Informatica intends to reinstate it in an upcoming GA release, but might choose not to in accordance with changing market or technical circumstances.

Data Processor Transformation in a Non-native Environment

The Data Processor transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Not supported in batch or streaming mappings.
- Databricks Spark engine. Not supported.

Data Processor Transformation on the Blaze Engine

The Data Processor transformation is supported with the following restrictions:

- Mapping validation fails when the transformation **Data processor mode** is set to **Input Mapping** or **Service and Input Mapping**.

Decision Transformation in a Non-native Environment

The Decision transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported with restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Decision Transformation on the Spark Engine

Consider the following general restriction:

You must configure the Decision transformation properties to be partitionable.

Expression Transformation in a Non-native Environment

The Expression transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported with restrictions.

Expression Transformation on the Blaze Engine

Mapping validation fails in the following situations:

- The transformation contains stateful variable ports.
- The transformation contains unsupported functions in an expression.

An Expression transformation with a user-defined function returns a null value for rows that have an exception error in the function.

Expression Transformation on the Spark Engine

Mapping validation fails in the following situations:

- The transformation contains stateful variable ports.
- The transformation contains unsupported functions in an expression.

Note: If an expression results in numerical errors, such as division by zero or SQRT of a negative number, it returns a null value and rows do not appear in the output. In the native environment, the expression returns an infinite or an NaN value.

Expression Transformation in a Streaming Mapping

Streaming mappings have the same processing rules as batch mappings on the Spark engine.

Expression Transformation on the Databricks Spark Engine

Mapping validation fails in the following situations:

- The transformation contains stateful variable ports.
- The transformation contains unsupported functions in an expression.

Note: If an expression results in numerical errors, such as division by zero or SQRT of a negative number, it returns a null value and rows do not appear in the output. In the native environment, the expression returns an infinite or an NaN value.

Filter Transformation in a Non-native Environment

The Filter transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported without restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported without restrictions.

Filter Transformation on the Blaze Engine

When a mapping contains a Filter transformation on a partitioned column of a Hive source, the Blaze engine can read only the partitions that contain data that satisfies the filter condition. To push the filter to the Hive source, configure the Filter transformation to be the next transformation in the mapping after the source.

Java Transformation in a Non-native Environment

The Java transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Not supported.

Java Transformation on the Blaze Engine

To use external .jar files in a Java transformation, perform the following steps:

1. Copy external .jar files to the Informatica installation directory in the Data Integration Service machine at the following location: `<Informatica installation directory>/services/shared/jars`. Then recycle the Data Integration Service.
2. On the machine that hosts the Developer tool where you develop and run the mapping that contains the Java transformation:
 - a. Copy external .jar files to a directory on the local machine.
 - b. Edit the Java transformation to include an import statement pointing to the local .jar files.
 - c. Update the classpath in the Java transformation.
 - d. Compile the transformation.

Java Transformation on the Spark Engine

Some processing rules for the Spark engine differ from the processing rules for the Data Integration Service.

General Restrictions

The Java transformation is supported with the following restrictions on the Spark engine:

- The Java code in the transformation cannot write output to standard output when you push transformation logic to Hadoop. The Java code can write output to standard error which appears in the log files.
- For date/time values, the Spark engine supports the precision of up to microseconds. If a date/time value contains nanoseconds, the trailing digits are truncated.

Partitioning

The Java transformation has the following restrictions when used with partitioning:

- The Partitionable property must be enabled in the Java transformation. The transformation cannot run in one partition.
- The following restrictions apply to the Transformation Scope property:
 - The value Transaction for transformation scope is not valid.
 - If you enable an input port for partition key, the transformation scope must be set to All Input.
 - Stateless must be enabled if the transformation scope is row.

Mapping Validation

Mapping validation fails in the following situations:

- You reference an unconnected Lookup transformation from an expression within a Java transformation.
- You select a port of a complex data type as the partition or sort key.
- You enable nanosecond processing in date/time and the Java transformation contains a port of complex data type with an element of a date/time type. For example, a port of type `array<data/time>` is not valid if you enable nanosecond processing in date/time.
- When you enable high precision, a validation error occurs if the Java transformation contains an expression that uses a decimal port or a complex port with an element of a decimal data type, and the port is used with an operator.

For example, if the transformation contains a decimal port `decimal_port` and you use the expression `decimal_port + 1`, a validation error occurs.

The mapping fails in the following situation:

- The Java transformation and the mapping use different precision modes when the Java transformation contains a decimal port or a complex port with an element of a decimal data type.

Even if high precision is enabled in the mapping, the mapping processes data in low-precision mode in some situations, such as when the mapping contains a complex port with an element of a decimal data type, or the mapping is a streaming mapping. If high precision is enabled in both the Java transformation and the mapping, but the mapping processes data in low-precision mode, the mapping fails.

Using External .jar Files

To use external .jar files in a Java transformation, perform the following steps:

1. Copy external .jar files to the Informatica installation directory in the Data Integration Service machine at the following location:
`<Informatica installation directory>/services/shared/jars`
2. Recycle the Data Integration Service.
3. On the machine that hosts the Developer tool where you develop and run the mapping that contains the Java transformation:
 - a. Copy external .jar files to a directory on the local machine.
 - b. Edit the Java transformation to include an import statement pointing to the local .jar files.
 - c. Update the classpath in the Java transformation.
 - d. Compile the transformation.

Java Transformation in a Streaming Mapping

Streaming mappings have additional processing rules that do not apply to batch mappings.

The Data Integration Service ignores the following properties:

- Partitionable. The Data Integration Service ignores the property and can process the transformation with multiple threads.
- Stateless. The Data Integration Service ignores the property and maintains the row order of the input data.

Joiner Transformation in a Non-native Environment

The Joiner transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported with restrictions.

Joiner Transformation on the Blaze Engine

Mapping validation fails in the following situations:

- The transformation contains an inequality join and map-side join is disabled.

- The Joiner transformation expression references an unconnected Lookup transformation.

Map-side join is disabled when the Joiner transformation is configured for detail outer join or full outer join.

Joiner Transformation on the Spark Engine

Mapping validation fails in the following situations:

- Case sensitivity is disabled.
- The join condition is of binary data type or contains binary expressions.

Joiner Transformation in a Streaming Mapping

Streaming mappings have additional processing rules that do not apply to batch mappings.

Mapping Validation

Mapping validation fails in the following situations:

- A Joiner transformation is downstream from an Aggregator transformation.
- A Joiner transformation is downstream from a Rank transformation.
- A streaming pipeline contains more than one Joiner transformation.
- A Joiner transformation joins data from streaming and non-streaming pipelines.

Joiner Transformation on the Databricks Spark Engine

Mapping validation fails in the following situations:

- Case sensitivity is disabled.
- The join condition is of binary data type or contains binary expressions.

Key Generator Transformation in a Non-native Environment

The Key Generator transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Labeler Transformation in a Non-native Environment

The Labeler transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Lookup Transformation in a Non-native Environment

The Lookup transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported with restrictions.

Lookup Transformation on the Blaze Engine

Some processing rules for the Blaze engine differ from the processing rules for the Data Integration Service.

Mapping validation fails in the following situations:

- The cache is configured to be shared, named, persistent, dynamic, or uncached. The cache must be a static cache.

If you add a data object that uses Sqoop as a Lookup transformation in a mapping, the Data Integration Service does not run the mapping through Sqoop. It runs the mapping through JDBC.

Lookup Transformation on the Spark Engine

Some processing rules for the Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- Case sensitivity is disabled.
- The lookup condition in the Lookup transformation contains binary data type.
- The cache is configured to be shared, named, persistent, dynamic, or uncached. The cache must be a static cache.

The mapping fails in the following situation:

- The transformation is unconnected and used with a Joiner or Java transformation.

Multiple Matches

When you choose to return the first, last, or any value on multiple matches, the Lookup transformation returns any value.

If you configure the transformation to report an error on multiple matches, the Spark engine drops the duplicate rows and does not include the rows in the logs.

Note: If an HBase lookup does not result in a match, it generates a row with null values for all columns. You can add a Filter transformation after the Lookup transformation to filter out null rows.

Lookup Transformation in a Streaming Mapping

Streaming mappings have additional processing rules that do not apply to batch mappings.

Mapping Validation

Mapping validation fails in the following situations:

- The lookup is a data object.
- An Aggregator transformation is in the same streaming pipeline as a passive Lookup transformation configured with an inequality lookup condition.
- A Rank transformation is in the same streaming pipeline as a passive Lookup transformation configured with an inequality lookup condition.
- A pipeline contains more than one passive Lookup transformation configured with an inequality condition.

The mapping fails in the following situations:

- The transformation is unconnected.

General Guidelines

Consider the following general guidelines:

- Using a float data type to look up data can return unexpected results.
- Use a Lookup transformation to look up data in a flat file, HDFS, Hive, relational, and HBase data.
- To avoid cross join of DataFrames, configure the Lookup transformation to ignore null values that match.

HBase Lookups

To use a Lookup transformation on uncached HBase tables, perform the following steps:

1. Create an HBase data object. When you add an HBase table as the resource for a HBase data object, include the ROW ID column.
2. Create a HBase read data operation and import it into the streaming mapping.
3. When you import the data operation to the mapping, select the **Lookup** option.
4. On the Lookup tab, configure the following options:
 - Lookup column. Specify an equality condition on ROW ID
 - Operator. Specify =
5. Verify that format for any date value in the HBase tables is of a valid Java date format. Specify this format in the **Date Time Format** property of the **Advanced Properties** tab of the data object read operation.

Note: If an HBase lookup does not result in a match, it generates a row with null values for all columns. You can add a Filter transformation after the Lookup transformation to filter out null rows.

Mapping validation fails in the following situations:

- The condition does not contain a ROW ID.
- The transformation contains an inequality condition.
- The transformation contains multiple conditions.
- An input column is of a date type.

Lookup Transformation on the Databricks Spark Engine

Some processing rules for the Databricks Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- Case sensitivity is disabled.
- The lookup condition in the Lookup transformation contains binary data type.
- The cache is configured to be shared, named, persistent, dynamic, or uncached. The cache must be a static cache.
- The lookup source is not Microsoft Azure SQL Data Warehouse.

The mapping fails in the following situation:

- The transformation is unconnected and used with a Joiner transformation.

Multiple Matches

When you choose to return the first, last, or any value on multiple matches, the Lookup transformation returns any value.

If you configure the transformation to report an error on multiple matches, the Spark engine drops the duplicate rows and does not include the rows in the logs.

Note: If an HBase lookup does not result in a match, it generates a row with null values for all columns. You can add a Filter transformation after the Lookup transformation to filter out null rows.

Match Transformation in a Non-native Environment

The Match transformation processing in the non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Match Transformation on the Blaze Engine

Mapping validation fails when the Match transformation is configured to write identity index data to database tables.

A Match transformation generates cluster ID values differently in native and non-native environments. In a non-native environment, the transformation appends a group ID value to the cluster ID.

Match Transformation on the Spark Engine

Mapping validation fails when the Match transformation is configured to write identity index data to database tables.

A Match transformation generates cluster ID values differently in native and non-native environments. In a non-native environment, the transformation appends a group ID value to the cluster ID.

Merge Transformation in a Non-native Environment

The Merge transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Normalizer Transformation in a Non-native Environment

The Normalizer transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported without restrictions.

Parser Transformation in a Non-native Environment

The Parser transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch and streaming mappings.
- Databricks Spark engine. Not supported.

Python Transformation in a Non-native Environment

The Python transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Not supported.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Not supported.

Python Transformation on the Spark Engine

Mapping validation fails if a user-defined default value is assigned to an output port.

The mapping fails in the following situations:

- An output port is not assigned a value in the Python code.
- The data types in corresponding input and output ports are not the same, and the Python code does not convert the data type in the input port to the data type in the output port.
- The Python transformation contains decimal ports and high precision is enabled in the mapping.

Note: The Data Integration Service does not validate Python code.

Python Transformation in a Streaming Mapping

Streaming mappings have additional processing rules that do not apply to batch mappings.

When you close a Jep instance, you might not be able to call CPython modules.

Rank Transformation in a Non-native Environment

The Rank transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported with restrictions.

Rank Transformation on the Blaze Engine

Some processing rules for the Blaze engine differ from the processing rules for the Data Integration Service.

The data cache for the Rank transformation is optimized to use variable length to store binary and string data types that pass through the Rank transformation. The optimization is enabled for record sizes up to 8 MB. If the record size is greater than 8 MB, variable length optimization is disabled.

When variable length is used to store data that passes through the Rank transformation in the data cache, the Rank transformation is optimized to use sorted input and a pass-through Sorter transformation is inserted before the Rank transformation in the run-time mapping.

To view the Sorter transformation, view the optimized mapping or view the execution plan in the Blaze validation environment.

During data cache optimization, the data cache and the index cache for the Rank transformation are set to Auto. The sorter cache for the Sorter transformation is set to the same size as the data cache for the Rank transformation. To configure the sorter cache, you must configure the size of the data cache for the Rank transformation.

Rank Transformation on the Spark Engine

Some processing rules for the Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- Case sensitivity is disabled.
- The rank port is of binary data type.

Data Cache Optimization

You cannot optimize the data cache for the transformation to store data using variable length.

Rank Transformation in a Streaming Mapping

Streaming mappings have additional processing rules that do not apply to batch mappings.

Mapping Validation

Mapping validation fails in the following situations:

- A Rank transformation is in the same streaming pipeline as a passive Lookup transformation configured with an inequality lookup condition.
- A Rank transformation is upstream from a Joiner transformation.
- A streaming pipeline contains more than one Rank transformation.
- A streaming pipeline contains an Aggregator transformation and a Rank transformation.

Rank Transformation on the Databricks Spark Engine

Some processing rules for the Databricks Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- Case sensitivity is disabled.
- The rank port is of binary data type.

Data Cache Optimization

You cannot optimize the data cache for the transformation to store data using variable length.

Router Transformation in a Non-native Environment

The Router transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch and streaming mappings.
- Databricks Spark engine. Supported without restrictions.

Sequence Generator Transformation in a Non-native Environment

The Sequence Generator transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Sequence Generator Transformation on the Blaze Engine

A mapping with a Sequence Generator transformation consumes significant resources when the following conditions are true:

- You set the **Maintain Row Order** property on the transformation to *true*.
- The mapping runs in a single partition.

Sequence Generator Transformation on the Spark Engine

The Sequence Generator transformation does not maintain row order in output data. If you enable the **Maintain Row Order** property on the transformation, the Data Integration Service ignores the property.

Sorter Transformation in a Non-native Environment

The Sorter transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch and streaming mappings.

- Databricks Spark engine. Supported with restrictions.

Sorter Transformation on the Blaze Engine

Some processing rules for the Blaze engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails in the following situations:

- The target is configured to maintain row order and the Sorter transformation is not connected directly to a flat file target.

Parallel Sorting

The Data Integration Service enables parallel sorting with the following restrictions:

- The mapping does not include another transformation between the Sorter transformation and the target.
- The data type of the sort keys does not change between the Sorter transformation and the target.
- Each sort key in the Sorter transformation must be linked to a column in the target.

Global Sort

The Blaze engine can perform global sorts when the following conditions are true:

- The Sorter transformation is connected directly to flat file targets.
- The target is configured to maintain row order.
- The sort key is not a binary data type.

If any of the conditions are not true, the Blaze engine performs a local sort.

Data Cache Optimization

If a Sorter transformation is inserted before an Aggregator or Rank transformation to optimize the Aggregator or Rank data cache, the size of the sorter cache is the same size as the data cache for the Aggregator or Rank transformation. To configure the sorter cache, you must configure the size of the data cache for the Aggregator or Rank transformation.

Sorter Transformation on the Spark Engine

Some processing rules for the Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails when case sensitivity is disabled.

The Data Integration Service logs a warning and ignores the Sorter transformation in the following situations:

- There is a type mismatch in between the target and the Sorter transformation sort keys.
- The transformation contains sort keys that are not connected to the target.
- The Write transformation is not configured to maintain row order.
- The transformation is not directly upstream from the Write transformation.

Null Values

The Data Integration Service treats null values as high even if you configure the transformation to treat null values as low.

Data Cache Optimization

You cannot optimize the sorter cache to store data using variable length.

Parallel Sorting

The Data Integration Service enables parallel sorting with the following restrictions:

- The mapping does not include another transformation between the Sorter transformation and the target.
- The data type of the sort keys does not change between the Sorter transformation and the target.
- Each sort key in the Sorter transformation must be linked to a column in the target.

Sorter Transformation in a Streaming Mapping

Streaming mappings have the same processing rules as batch mappings on the Spark engine.

Mapping Validation

Mapping validation fails in the following situation:

- The Sorter transformation in a streaming mapping does not have an upstream Aggregator transformation.

Sorter Transformation on the Databricks Spark Engine

Some processing rules for the Databricks Spark engine differ from the processing rules for the Data Integration Service.

Mapping Validation

Mapping validation fails when case sensitivity is disabled.

The Data Integration Service logs a warning and ignores the Sorter transformation in the following situations:

- There is a type mismatch in between the target and the Sorter transformation sort keys.
- The transformation contains sort keys that are not connected to the target.
- The Write transformation is not configured to maintain row order.
- The transformation is not directly upstream from the Write transformation.

Null Values

The Data Integration Service treats null values as high even if you configure the transformation to treat null values as low.

Data Cache Optimization

You cannot optimize the sorter cache to store data using variable length.

Parallel Sorting

The Data Integration Service enables parallel sorting with the following restrictions:

- The mapping does not include another transformation between the Sorter transformation and the target.
- The data type of the sort keys does not change between the Sorter transformation and the target.
- Each sort key in the Sorter transformation must be linked to a column in the target.

Standardizer Transformation in a Non-native Environment

The Standardizer transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Union Transformation in a Non-native Environment

The Union transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Supported with restrictions in streaming mappings.
- Databricks Spark engine. Supported without restrictions.

Union Transformation in a Streaming Mapping

Streaming mappings have mapping validation restrictions.

Mapping Validation

Mapping validation fails in the following situation:

- The transformation is configured to merge data from streaming and non-streaming pipelines.

Update Strategy Transformation in a Non-native Environment

The Update Strategy transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported with restrictions.
- Spark engine. Supported with restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

Note: The Update Strategy transformation is supported only on Hadoop distributions that support Hive ACID.

Update Strategy Transformation on the Blaze Engine

You can use the Update Strategy transformation on the Hadoop distributions that support Hive ACID.

Some processing rules for the Blaze engine differ from the processing rules for the Data Integration Service.

General Restrictions

If the Update Strategy transformation receives multiple update rows for the same primary key value, the transformation selects one random row to update the target.

If multiple Update Strategy transformations write to different instances of the same target, the target data might be unpredictable.

The Blaze engine executes operations in the following order: deletes, updates, inserts. It does not process rows in the same order as the Update Strategy transformation receives them.

Hive targets always perform Update as Update operations. Hive targets do not support Update Else Insert or Update as Insert.

Mapping Validation and Compile Validation

Mapping validation fails in the following situations:

- The Update Strategy transformation is connected to more than one target.
- The Update Strategy transformation is not located immediately before the target.
- The Update Strategy target is not a Hive target.
- The Update Strategy transformation target is an external ACID table.
- The target does not contain a primary key.
- The Hive target property to truncate the target table at run time is enabled.
- The Hive target property to create or replace the target table at run time is enabled.

The mapping fails in the following situation:

- The target is not ORC bucketed.
- The Hive target is modified to have fewer rows than the actual table.

Compile validation errors occur and the mapping execution stops in the following situations:

- The Hive version is earlier than 0.14.
- The target table is not enabled for transactions.

Using Hive Target Tables

To use a Hive target table with an Update Strategy transformation, you must create the Hive target table with the following clause in the Hive Data Definition Language: `TBLPROPERTIES ("transactional"="true")`.

To use an Update Strategy transformation with a Hive target, verify that the following properties are configured in the `hive-site.xml` configuration set associated with the Hadoop connection:

```
hive.support.concurrency      true
hive.enforce.bucketing       true
hive.exec.dynamic.partition.mode nonstrict
hive.txn.manager              org.apache.hadoop.hive.ql.lockmgr.DbTxnManager
hive.compactor.initiator.on   true
hive.compactor.worker.threads 1
```

Update Strategy Transformation on the Spark Engine

You can use the Update Strategy transformation on the Hadoop distributions that support Hive ACID.

You can also use the Update Strategy transformation to write mapping results to JDBC-compliant relational targets.

Some processing rules for the Spark engine differ from the processing rules for the Data Integration Service.

General Restrictions for Hive Targets

The Update Strategy transformation does not forward rejected rows to the next transformation when the target is a Hive table or JDBC-compliant table.

If the Update Strategy transformation receives multiple update rows for the same primary key value, the transformation selects one random row to update the target.

If multiple Update Strategy transformations write to different instances of the same target, the target data might be unpredictable.

If the mapping runs on the Spark engine, you can choose the Use Hive Merge option. The option has the following restrictions:

- If you run the mapping with Blaze or Hive, the Hive MERGE option is ignored. The Data Integration Service runs the transformation with the previous implementation.
- A single row for delete or update cannot match multiple rows in the target. When the mapping violates this restriction, the mapping fails with a runtime error.
- If you configure the Update Strategy expression to update partitioning or bucketing columns, the mapping ignores the Hive MERGE option. In addition, the mapping does not update the columns.

Note: The Developer tool and the Data Integration Service do not validate against these restrictions. If the expression or the mapping violates these restrictions, the mapping might run, but the results will not be as expected.

Hive targets always perform Update as Update operations. Hive targets do not support Update Else Insert or Update as Insert.

Mapping Validation

Mapping validation fails in the following situations.

- The Update Strategy transformation is connected to more than one target.
- The Update Strategy transformation is not located immediately before the target.
- The Update Strategy transformation target is an external ACID table.
- The target does not contain a connected primary key.
- The property to enable truncation of the Hive or relational target table at run time is selected.
- One of the following target strategies for the Hive or relational target table at run time is selected:
 - Create or replace the target table
 - ApplyNewColumns
 - ApplyNewSchema
 - Fail

The mapping fails in the following situations when the target is a Hive target:

- The target table is not enabled for transactions.
- The target is not ORC bucketed.

Using Hive Target Tables

To use a Hive target table with an Update Strategy transformation, you must create the Hive target table with the following clause in the Hive Data Definition Language: `TBLPROPERTIES ("transactional"="true")`.

To use an Update Strategy transformation with a Hive target, verify that the following properties are configured in the `hive-site.xml` configuration set associated with the Hadoop connection:

```
hive.support.concurrency      true
hive.enforce.bucketing       true
hive.exec.dynamic.partition.mode nonstrict
hive.txn.manager              org.apache.hadoop.hive.ql.lockmgr.DbTxnManager
hive.compactor.initiator.on   true
hive.compactor.worker.threads 1
```

Using JDBC Target Tables

You can configure JDBC-compliant relational database targets for the Update Strategy transformation in a mapping to run on the Spark engine.

Weighted Average Transformation in a Non-native Environment

The Weighted Average transformation processing in a non-native environment depends on the engine that runs the transformation.

Consider the support for the following non-native run-time engines:

- Blaze engine. Supported without restrictions.
- Spark engine. Supported without restrictions in batch mappings. Not supported in streaming mappings.
- Databricks Spark engine. Not supported.

CHAPTER 6

Cluster Workflows

This chapter includes the following topics:

- [Cluster Workflows Overview, 112](#)
- [Cluster Workflow Components, 113](#)
- [Cluster Workflows Process , 114](#)
- [Create Cluster Task Properties, 115](#)
- [Mapping Task Properties, 123](#)
- [Add a Delete Cluster Task, 123](#)
- [Deploy and Run the Workflow, 124](#)

Cluster Workflows Overview

You can run a workflow to create a cluster that runs Mapping and other tasks on a cloud platform cluster.

Create cluster workflows to run on the Amazon EMR or Microsoft Azure HDInsight cloud platforms in the Hadoop environment. Create cluster workflows to run on the Databricks cloud platform in the Databricks environment.

The cluster workflow uses other elements that enable communication between the Data Integration Service and the cloud platform, such as a cloud provisioning configuration and a cluster connection.

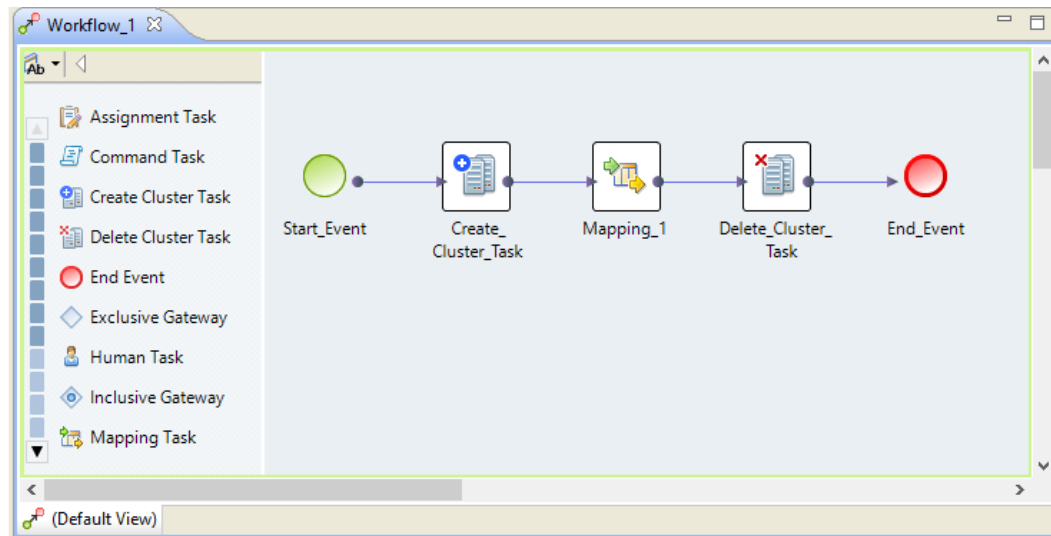
A cluster workflow contains a Create Cluster task that you configure with information about the cluster to create. If you want to create an ephemeral cluster, you can include a Delete Cluster task. An ephemeral cluster is a cloud platform cluster that you create to run mappings and other tasks, and then terminate when tasks are complete. Create ephemeral clusters to save cloud platform resources.

Note: To create a cluster on Cloudera Altus, you create a workflow with Command tasks that perform the tasks that a cluster workflow automates. For more information about creating a cluster on Cloudera Altus, see the article "Implementing Informatica Big Data Management with Ephemeral Clusters on a Cloudera Altus Cluster" on the Informatica Network.

Cluster Workflow Components

The cluster workflow that creates an ephemeral cluster includes a Create Cluster task, at least one Mapping task, and a Delete Cluster task.

The following image shows a sample cluster workflow:



A cluster workflow uses the following components:

Cloud provisioning configuration

The cloud provisioning configuration is associated with the Create Cluster task through the cluster connection.

Cluster connection

The cluster connection to use with a cluster workflow is associated with a cloud provisioning configuration. You can use a Hadoop or Databricks cluster connection. When you run the workflow the Data Integration Service creates a temporary cluster connection.

Create Cluster task

The Create Cluster task contains all the settings that the cloud platforms require to create a cluster with a master node and worker nodes. It also contains a reference to a cloud provisioning configuration. Include one Create Cluster task in a cluster workflow.

Mapping task

Add a big data mapping to the Mapping task. A cluster workflow can include more than one mapping task. You can run some mappings on an existing cluster and you can run some mappings on a cluster that the workflow creates. You configure the mappings and Mapping tasks based on where you want to run the task.

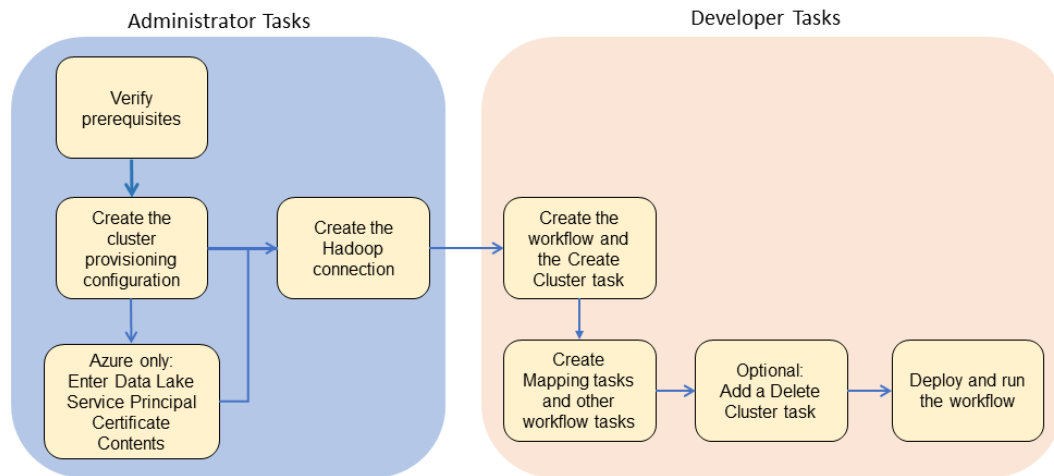
Delete Cluster task

The Delete Cluster task terminates the cluster and deletes the cluster and all resources that the workflow creates.

Cluster Workflows Process

Creation of a cluster workflow requires administrator and developer tasks.

The following image shows the process to create, configure, and run a cluster workflow:



The cluster workflow development process requires tasks from the following users:

Administrator

1. **Verify installation.** The domain can be on-premises or reside on the cloud platform.
2. **Create a cluster provisioning configuration.** Create a cloud provisioning configuration in the Administrator tool. The cloud provisioning configuration contains all of the information that the Data Integration Service requires to contact and create resources on the cloud platform.
3. **Create a cluster connection.** Create a dedicated cluster connection to associate with the cluster provisioning configuration.

For information about administrator tasks, see the *Big Data Management Administrator Guide*.

Developer

A developer completes the following tasks:

1. Create the workflow and the Create Cluster task.
2. Create Mapping tasks and other tasks as required.
3. Create a Delete Cluster task if you want to delete the cluster and resources when processing is complete.
4. Deploy and run the workflow.

Create Cluster Task Properties

When you configure a Create Cluster task, you configure options that are common for all cluster types, and you configure advanced options for specific clusters.

General Properties for the Create Cluster Task

The following table describes the General properties for the Create Cluster task:

Property	Description
Name	Task name.
Description	Optional description.
Connection Name	Name of the cloud provisioning configuration to use with the workflow.
Connection Type	Choose from the following options: - Amazon EMR. Create an Amazon EMR cluster in the Hadoop environment. - HDInsight. Create an Azure HDInsight cluster in the Hadoop environment. - Databricks. Create a Databricks cluster in the Databricks environment.

Input Properties for the Create Cluster Task

You can assign task input to workflow parameters and variables.

The following table describes the Input properties for the Create Cluster task:

Property	Description
Create Cluster Task Inputs	Parameters or variables to use for task input.

Output Properties for the Create Cluster Task

You can assign variables to the Create Cluster task output.

The following table describes the Output properties for the Create Cluster task:

Property	Description
Start Time	Date and time that the task started running.
End Time	Date and time that the task stopped running
Is Successful	Indicates whether the task ran successfully.
Cluster Identifier	The Cluster Identifier property of the Create Cluster task overrides the Cluster Identifier property of the Mapping task. Default is AutoDeployCluster.

Advanced Properties for the Create Cluster Task

Advanced properties for the Create Cluster task depend on the cluster type that you create.

For information about advanced properties for Microsoft Azure, see [“Advanced Properties for Azure HDInsight” on page 116](#).

For information about advanced properties for Amazon Web Services, see [“Advanced Properties for Amazon EMR” on page 117](#).

For information about advanced properties to run mappings on the Blaze engine, see [“Advanced Properties for the Blaze Engine” on page 119](#).

For information about advanced properties for Databricks, see [“Advanced Properties for Databricks” on page 121](#).

Advanced Properties for Azure HDInsight

Set the advanced properties for an Azure HDInsight cluster.

The following table describes the Advanced properties for a Microsoft Azure HDInsight cluster:

Property	Description
Cluster Name	Name of the cluster to create.
Azure Cluster Type	Type of the cluster to be created. Choose one of the options in the drop-down list. Default is Hadoop.
HDInsight version	HDInsight version to run on the cluster. Enter the HDInsight version tag string to designate the version. Default is the latest version supported.
Azure Cluster Location	Use the drop-down list to choose the location in which to create the cluster.
Head Node VM Size	Size of the head node instance to create. Default is Standard_D12_v2.
Number of Worker Node Instances	Number of worker node instances to create in the cluster. Default is 2.
Worker Node VM Size	Size of the worker node instance to create. Default is Standard_D13_v2.
Default Storage Type	Primary storage type to be used for the cluster. Choose one of the following options: - Azure Data Lake Store - Azure BLOB storage account Default is BLOB storage
Default Storage Container or Root Mount Path	Default container for data. Type one of the following paths: - For ADLS storage, type the path to the storage. For example, you can type <code>storage-name</code> or <code>storage-name/folder-name</code>. - For blob storage, type the path to the container. Format: <code>/path/</code>
Log Location	Optional. Path to the directory to store workflow event logs. Default is <code>/app-logs</code> .

Property	Description
Attach External Hive Metastore	If you select this option, the workflow attaches an external Hive metastore to the cluster if you configured an external Hive metastore in the cloud provisioning configuration.
Bootstrap JSON String	JSON statement to run during cluster creation. You can use this statement to configure cluster details. For example, you could configure Hadoop properties on the cluster, add tags to cluster resources, or run script actions. Choose one of the following methods to populate the property: - Type the JSON statement. Use the following format: <pre>{ "core-site" : { "<sample_property_key1>": "<sample_property_val1>", "<sample_property_key2>": "<sample_property_val2>" }, "tags": { "<tag_key>": "<tag_val>" }, "scriptActions": [{ "name": "setenvironmentvariable", "uri": "scriptActionUri", "parameters": "headnode" }] }</pre> - Provide a path to a file that contains a JSON statement. Format: <pre>file://<path_to_bootstrap_file></pre>

Advanced Properties for Amazon EMR

Set the advanced properties for an Amazon EMR cluster.

General Options

The following table describes general options that you can set for an EMR cluster:

Property	Description
Cluster Name	Name of the cluster to create.
Release Version	EMR version to run on the cluster. Enter the AWS version tag string to designate the version. For example: <code>emr-5.8.0</code> Default is Latest version supported.
Connection Name	Name of the Hadoop connection that you configured for use with the cluster workflow.
S3 Log URI	Optional. S3 location of logs for cluster creation. Format: <pre>s3://<bucket name>/<folder name></pre> If you do not supply a location, no cluster logs will be stored.

Master Instance Group Options

The following table describes master instance group options that you can set for an EMR cluster:

Property	Description
Master Instance Type	Master node EC2 instance type. You can specify any available EC2 instance type. Default is m4.4xlarge.
Master Instance Maximum Spot Price	Maximum spot price for the master node. Setting this property changes the purchasing option of the master instance group to Spot instead of On-demand.

Core Instance Group Options

The following table describes core instance group options that you can set for an EMR cluster:

Property	Description
Core Instance Type	Core node EC2 instance type. You can specify any available EC2 instance type. Default is m4.4xlarge.
Core Instance Count	Number of core EC2 instances to create in the cluster. Default is 2.
Core Instance Maximum Spot Price	Maximum spot price for core nodes. Setting this property changes the purchasing option of the core instance group to Spot instead of On-demand.
Core Auto-Scaling Policy	Optional. Auto-scaling policy for core instances. Type the policy JSON statement here, or provide a path to a file that contains a JSON statement. Format: <code>file:\\<path_to_policy_config_file></code>

Task Instance Group Options

The following table describes task instance group options that you can set for an EMR cluster:

Property	Description
Task Instance Type	Task node EC2 instance type. You can specify any available EC2 instance type. Default is m4.4xlarge.
Task Instance Count	Number of task EC2 instances to create in the cluster. Default is 2.

Property	Description
Task Instance Maximum Spot Price	Maximum spot price for task nodes. Setting this property changes the purchasing option of the task instance group to Spot instead of On-demand.
Task Auto-Scaling Policy	Optional. Auto-scaling policy for task instances. Type the policy JSON statement here, or provide a path to a file that contains a JSON statement. Format: file:\\<path_to_policy_config_file>

Additional Options

The following table describes additional options that you can set for an EMR cluster:

Property	Description
Applications	Optional. Applications to add to the default applications that AWS installs. AWS installs certain applications when it creates an EMR cluster. In addition, you can specify additional applications. Select additional applications from the drop-down list. This field is equivalent to the Software Configuration list in the AWS EMR cluster creation wizard.
Tags	Optional. Tags to propagate to cluster EC2 instances. Tags assist in identifying EC2 instances. Format: TagName1=TagValue1, TagName2=TagValue2
Software Settings	Optional. Custom configurations to apply to the applications installed on the cluster. This field is equivalent to the Edit Software Settings field in the AWS cluster creation wizard. You can use this as a method to modify the software configuration on the cluster. Type the configuration JSON statement here, or provide a path to a file that contains a JSON statement. Format: file:\\<path_to_custom_config_file>
Steps	Optional. Commands to run after cluster creation. For example, you can use this to run Linux commands or HDFS or Hive Hadoop commands. This field is equivalent to the Add Steps field in the AWS cluster creation wizard. Type the command statement here, or provide a path to a file that contains a JSON statement. Format: file:\\<path_to_command_file>
Bootstrap Actions	Optional. Actions to perform after EC2 instances are running, and before applications are installed. Type the JSON statement here, or provide a path to a file that contains a JSON statement. Format: file:\\<path_to_policy_config_file>
Task Recovery Strategy	Choose from the following options: - Restart task - Skip task Default is Restart task.

Advanced Properties for the Blaze Engine

If you want to use the Blaze engine to run mappings on the cloud platform cluster, you must set cluster configuration properties in the Software Setting property of the Create Cluster task.

Configure the Create Cluster task to set configuration properties in *-site.xml files on the cluster. Hadoop clusters run based on these settings.

Use a text file to specify *-site.xml file settings, and specify the path to the file in the Software Settings property.

The following text shows sample configuration of the Software Settings property file:

```
[
  {
    "Classification": "yarn-site",
    "Properties": {
      "yarn.scheduler.minimum-allocation-mb": "250",
      "yarn.scheduler.maximum-allocation-mb": "8192",
      "yarn.nodemanager.resource.memory-mb": "16000",
      "yarn.nodemanager.resource.cpu-vcores": "12"
    }
  },
  {
    "Classification": "core-site",
    "Properties": {
      "hadoop.proxyuser.<DIS/OSPUSER>.groups": "<group names>",
      "hadoop.proxyuser.<DIS/OSPUSER>.hosts": "*"
    }
  }
]
```

yarn-site

yarn.scheduler.minimum-allocation-mb

The minimum RAM available for each container. Required for Blaze engine resource allocation.

yarn.scheduler.maximum-allocation-mb

The maximum RAM available for each container. Required for Blaze engine resource allocation.

yarn.nodemanager.resource.memory-mb

The maximum RAM available for each container. Set the maximum memory on the cluster to increase resource memory available to the Blaze engine.

yarn.nodemanager.resource.cpu-vcores

The number of virtual cores for each container. Required for Blaze engine resource allocation.

core-site

hadoop.proxyuser.<proxy user>.groups

Defines the groups that the proxy user account can impersonate. On a secure cluster the <proxy user> is the Service Principal Name that corresponds to the cluster keytab file. On a non-secure cluster, the <proxy user> is the system user that runs the Informatica daemon.

Set to group names of impersonation users separated by commas. If less security is preferred, use the wildcard " * " to allow impersonation from any group.

hadoop.proxyuser.<user>.hosts

Defines the host machines from which impersonated connections are allowed. On a secure cluster the <proxy user> is the Service Principal Name that corresponds to the cluster keytab file. On a non-secure cluster, the <proxy user> is the system user that runs the Informatica daemon.

Set the property to " * " to allow impersonation from any host. This is required to run a Blaze mapping on a cloud platform cluster.

Advanced Properties for a Hive Metastore Database

If you want to use a relational database on the cloud platform as the Hive metastore database for the cluster, you must set cluster configuration properties in the Software Setting property of the Create Cluster task.

Configure the Create Cluster task to set configuration properties in the hive-site.xml configuration file on the cluster. Use a text file to specify hive-site settings, and specify the path to the file in the Software Settings property.

Create the following properties in the Software Settings property file:

javax.jdo.option.ConnectionURL

JDBC connection string for the data store.

javax.jdo.option.ConnectionDriverName

JDBC driver class name for the data store. Specify a JDBC driver that is compatible with the cloud platform.

javax.jdo.option.ConnectionUserName

User name to use to connect to the database.

javax.jdo.option.ConnectionPassword

Password for the database user account.

The following text shows sample configuration of the Software Settings property file:

```
{
  "Classification": "hive-site",
  "Properties": {
    "javax.jdo.option.ConnectionURL": "jdbc:mysql://<RDS_HOST>:<PORT>\/
<USER_SCHEMA>?createDatabaseIfNotExist=true",
    "javax.jdo.option.ConnectionDriverName": "<JDBC driver name>",
    "javax.jdo.option.ConnectionUserName": "<USER>",
    "javax.jdo.option.ConnectionPassword": "<USER>"
  }
}
```

For example,

```
{
  "Classification": "hive-site",
  "Properties": {
    "javax.jdo.option.ConnectionURL": "jdbc:mysql://<host name>:<port number>\/
hive?createDatabaseIfNotExist=true",
    "javax.jdo.option.ConnectionDriverName": "org.mariadb.jdbc.Driver",
    "javax.jdo.option.ConnectionUserName": "hive",
    "javax.jdo.option.ConnectionPassword": "hive"
  }
}
```

Advanced Properties for Databricks

Set general and advanced options for advanced properties of the Databricks Create Cluster task.

General Options

The following table describes the general options that you can set for a Databricks cluster:

Property	Description
Cluster Name	Name of the cluster to create.
Databricks Runtime Version	The Databricks version to run on the cluster. Default is the latest supported version.
Driver Type	The type of node that you want to use for the driver node. Default is the worker type ID.
Worker Type	The type of node that you want to use for the worker node.
Workers	The number of worker nodes to create for the cluster. If you configure the cluster to scale automatically, this property is ignored. Default is 1.
Autoscale	Automatically scales the number of worker nodes based on workload.
Min Workers	The minimum number of worker nodes to use when the cluster is configured to scale automatically. Default is 0.
Max Workers	The maximum number of worker nodes to use when the cluster is configured to scale automatically. Default is 1.

Advanced Options

Configure advanced options such as environment variables and automatic termination.

The following table describes the advanced options that you can set for a Databricks cluster:

Property	Description
Auto Termination	Enables automatic termination of the cluster.
Auto Termination Time	Terminates the cluster after it is inactive for the specified number of minutes. Enter a value between 10 and 10,000. If you do not configure this, or if you set to 0, the cluster will not automatically terminate.
Cluster Log Conf	The location to deliver logs for long-term storage. If configured, the Databricks Spark engine will deliver the logs every five minutes. Provide the path to DBFS.
Init Scripts	The location to store init script logs. You can enter multiple destinations. The scripts are run sequentially in the order that you configure them. Use the follow format: <code><DBFS path>/script1, script2</code>
Cluster Tags	Labels that you can assign to resources for tracking purposes. Enter key-value pairs in the following format: "<key1>" = <value1> , "<key2>" = <value2>. You can also provide a path to a file that contains the key-value pairs.

Property	Description
Spark Configurations	Performance configurations for the Databricks Spark engine. Enter key-value pairs in the following format: "<key1>" = <value1> , "<key2>" = <value2>. You can also provide a path to a file that contains the key-value pairs.
Environment Variables	Environment variables that you can configure for the Databricks Spark engine. Enter key-value pairs in the following format: "<key1>" = <value1> , "<key2>" = <value2>.

Mapping Task Properties

Set mapping and Mapping task properties to specify where the workflow runs Mapping tasks.

When you include a Mapping task in a cluster workflow, you must identify the cluster where you want the mapping to run.

The following table describes the Cluster Identifier property values that you can configure in the Advanced properties of a Mapping task:

Cluster Identifier Property	Description
Auto Deploy	Runs the mapping on the cluster that the workflow creates. The Data Integration Service generates a temporary cluster connection based on the connection associated with the cluster provisioning configuration. The connection property in the mapping must also be set to Auto Deploy.
Blank	Runs the mapping on the cluster configured in the cluster connection associated with the mapping. The connection property in the mapping must be set to a cluster connection associated with a cluster configuration. Validation fails if it is associated with the cluster provisioning configuration.
Assign to task input	Accept input from a source other than the Create Cluster task. Use this option if you configure input properties to use parameters.

Note: If the Mapping task and the Create Cluster task contain different values for the cluster identifier, the Create Cluster task value takes precedence.

Add a Delete Cluster Task

To create an ephemeral cluster, add a Delete Cluster task.

The Delete Cluster task terminates the cluster and deletes the cluster and other resources that the cluster workflow creates.

If you do not add a Delete Cluster task, the cluster that the workflow creates remains running when the workflow ends. You can delete the cluster at any time.

1. Drag a Delete Cluster task to the workflow editor.

2. In the General properties, optionally rename the Delete Cluster task.
3. Connect the final task in the workflow to the Delete Cluster task, and connect the Delete Cluster task to the workflow End_Event.

You can also use `infacmd ccps deleteCluster` to delete a cloud cluster.

Deploy and Run the Workflow

After you complete the cluster workflow, deploy and run the workflow.

You can monitor AWS and Azure cluster workflows on the web console. If you configured a log location, view the logs at the location that you configured in the Create Cluster task properties.

You can monitor Databricks cluster workflows on the Databricks workspace.

You can also monitor Data Integration Service jobs in the Administrator tool.

Note: After the workflow begins executing tasks, the task to provision the cluster may take several minutes.

Monitoring Azure HDInsight Cluster Workflow Jobs

You can access mapping log URLs through the **Monitoring** tab in the Administrator tool to monitor workflow jobs that run on an Azure HDInsight cluster. The log location depends on the run-time engine that each mapping uses.

To access the monitoring URL for mappings that run on the Blaze or Spark engine, expand the workflow and the mapping on the **Monitoring** tab. Select the Grid Task and view the value for the Monitoring URL property in the lower pane. Use this path to find the log.

CHAPTER 7

Profiles

This chapter includes the following topics:

- [Profiles Overview, 125](#)
- [Native Environment, 125](#)
- [Hadoop Environment, 126](#)
- [Creating a Single Data Object Profile in Informatica Developer, 127](#)
- [Creating an Enterprise Discovery Profile in Informatica Developer, 127](#)
- [Creating a Column Profile in Informatica Analyst, 129](#)
- [Creating an Enterprise Discovery Profile in Informatica Analyst, 130](#)
- [Creating a Scorecard in Informatica Analyst, 131](#)
- [Monitoring a Profile, 132](#)
- [Profiling Functionality Support, 133](#)
- [Troubleshooting, 133](#)

Profiles Overview

You can create and run column profiles, enterprise discovery profiles, and scorecards in the native run-time environment or Hadoop run-time environment.

When you create or edit a profile or scorecard, you can choose the run-time environment. After you choose the run-time environment, the Developer tool or the Analyst tool sets the run-time environment in the profile definition. To process the profiles quickly, you can choose the Hadoop run-time environment.

Native Environment

In Informatica Developer, you can run single object profiles, multiple object profiles, and enterprise discovery profiles in the native environment. In Informatica Analyst, you can run column profiles, enterprise discovery profiles, and scorecards in the native environment.

When you run a profile in the native run-time environment, the Analyst tool or Developer tool submits the profile jobs to the Profiling Service Module. The Profiling Service Module then breaks down the profile jobs into a set of mappings. The Data Integration Service runs these mappings on the same machine where the

Data Integration Service runs and writes the profile results to the profiling warehouse. By default, all profiles run in the native run-time environment.

You can use native sources to create and run profiles in the native environment. A native data source is a non-Hadoop source, such as a flat file, relational source, or mainframe source. You can also run a profile on a mapping specification or a logical data source with a Hive or HDFS data source in the native environment.

Hadoop Environment

You can run profiles and scorecards in the Hadoop environment on the Blaze engine. You can choose the Hadoop option to run the profiles in the Hadoop run-time environment. After you choose the Hadoop option and select a Hadoop connection, the Data Integration Service pushes the profile logic to the Blaze engine on the Hadoop cluster to run profiles.

When you run a profile in the Hadoop environment, the Analyst tool or Developer tool submits the profile jobs to the Profiling Service Module. The Profiling Service Module then breaks down the profile jobs into a set of mappings. The Data Integration Service pushes the mappings to the Hadoop environment through the Hadoop connection. The Blaze engine processes the mappings and the Data Integration Service writes the profile results to the profiling warehouse.

In the Developer tool, you can run single object profiles and multiple object profiles, and enterprise discovery profiles on the Blaze engine. In the Analyst tool, you can run column profiles, enterprise discovery profiles, and scorecards on the Blaze engine.

Column Profiles for Sqoop Data Sources

You can run a column profile on data objects that use Sqoop. You can select the Hadoop run-time environment to run the column profiles.

When you run a column profile on a logical data object or customized data object, you can configure the num-mappers argument to achieve parallelism and optimize performance. You must also configure the split-by argument to specify the column based on which Sqoop must split the work units.

Use the following syntax:

```
--split-by <column_name>
```

If the primary key does not have an even distribution of values between the minimum and maximum range, you can configure the split-by argument to specify another column that has a balanced distribution of data to split the work units.

If you do not define the split-by column, Sqoop splits work units based on the following criteria:

- If the data object contains a single primary key, Sqoop uses the primary key as the split-by column.
- If the data object contains a composite primary key, Sqoop defaults to the behavior of handling composite primary keys without the split-by argument. See the Sqoop documentation for more information.
- If a data object contains two tables with an identical column, you must define the split-by column with a table-qualified name. For example, if the table name is CUSTOMER and the column name is FULL_NAME, define the split-by column as follows:

```
--split-by CUSTOMER.FULL_NAME
```
- If the data object does not contain a primary key, the value of the m argument and num-mappers argument default to 1.

When you use Cloudera Connector Powered by Teradata or Hortonworks Connector for Teradata and the Teradata table does not contain a primary key, the split-by argument is required.

Creating a Single Data Object Profile in Informatica Developer

You can create a single data object profile for one or more columns in a data object and store the profile object in the Model repository.

1. In the **Object Explorer** view, select the data object you want to profile.
2. Click **File > New > Profile** to open the profile wizard.
3. Select **Profile** and click **Next**.
4. Enter a name for the profile and verify the project location. If required, browse to a new location.
5. Optionally, enter a text description of the profile.
6. Verify that the name of the data object you selected appears in the **Data Objects** section.
7. Click **Next**.
8. Configure the profile operations that you want to perform. You can configure the following operations:
 - Column profiling
 - Primary key discovery
 - Functional dependency discovery
 - Data domain discovery

Note: To enable a profile operation, select **Enabled as part of the "Run Profile" action** for that operation. Column profiling is enabled by default.
9. Review the options for your profile.

You can edit the column selection for all profile types. Review the filter and sampling options for column profiles. You can review the inference options for primary key, functional dependency, and data domain discovery. You can also review data domain selection for data domain discovery.
10. Review the drill-down options, and edit them if necessary. By default, the **Enable Row Drilldown** option is selected. You can edit drill-down options for column profiles. The options also determine whether drill-down operations read from the data source or from staged data, and whether the profile stores result data from previous profile runs.
11. In the **Run Settings** section, choose a run-time environment. Choose **Native** or **Hadoop** as the run-time environment. After you choose the Hadoop option, you can select a Hadoop connection.
12. Click **Finish**.

Creating an Enterprise Discovery Profile in Informatica Developer

You can create a profile on multiple data sources under multiple connections. The Developer tool creates individual profile tasks for each source.

1. In the **Object Explorer** view, select multiple data objects you want to run a profile on.
2. Click **File > New > Profile** to open the profile wizard.
3. Select **Enterprise Discovery Profile** and click **Next**.

4. Enter a name for the profile and verify the project location. If required, browse to a new location.
5. Verify that the name of the data objects you selected appears within the **Data Objects** section. Click **Choose** to select more data objects, if required.
6. Click **Next**.

The **Add Resources to Profile Definition** pane appears. You can select multiple, external relational connections and data sources from this pane.
7. Click **Choose** to open the **Select Resources** dialog box.

The **Resources** pane lists all the internal and external connections and data objects under the Informatica domain.
8. Click **OK** to close the dialog box.
9. Click **Next**.
10. Configure the profile types that you want to run. You can configure the following profile types:
 - Data domain discovery
 - Column profile
 - Primary key profile
 - Foreign key profile

Note: Select **Enabled as part of "Run Enterprise Discovery Profile" action** for the profile types that you want to run as part of the enterprise discovery profile. Column profiling is enabled by default.
11. Review the options for the profile.

You can edit the sampling options for column profiles. You can also edit the inference options for data domain, primary key, and foreign key profiles.
12. Select **Create profiles**.

The Developer tool creates profiles for each individual data source.
13. Select **Run enterprise discovery profile on finish** to run the profile when you complete the profile configuration. If you enabled all the profiling operations, the Developer tool runs column, data domain, and primary key profiles on all selected data sources. Then, the Developer tool runs a foreign key profile across all the data sources.
14. Click **Finish**.

After you run an enterprise discovery profile, you need to refresh the Model Repository Service before viewing the results. This step is required as the import of metadata for external connections happens in the Model repository. You need to refresh the Model Repository Service so that the Developer tool reflects the changes to the Model repository.

Creating a Column Profile in Informatica Analyst

You can create a custom profile or default profile. When you create a custom profile, you can configure the columns, sample rows, and drill-down options. When you create a default profile, the column profile and data domain discovery runs on the entire data set with all the data domains.

1. In the **Discovery** workspace, click **Profile**, or select **New > Profile** from the header area.

Note: You can right-click on the data object in the **Library** workspace and create a profile. In this profile, the profile name, location name, and data object are extracted from the data object properties. You can create a default profile or customize the settings to create a custom profile.

The **New Profile** wizard appears.

2. The **Single source** option is selected by default. Click **Next**.
3. In the **Specify General Properties** screen, enter a name and an optional description for the profile. In the Location field, select the project or folder where you want to create the profile. Click **Next**.
4. In the **Select Source** screen, click **Choose** to select a data object, or click **New** to import a data object. Click **Next**.
 - In the **Choose Data Object** dialog box, select a data object. Click **OK**.
The Properties pane displays the properties of the selected data object. The Data Preview pane displays the columns in the data object.
 - In the **New Data Object** dialog box, you can choose a connection, schema, table, or view to create a profile on, select a location, and create a folder to import the data object. Click **OK**.
5. In the **Select Source** screen, select the columns that you want to run a profile on. Optionally, select **Name** to select all the columns. Click **Next**.

All the columns are selected by default. The Analyst tool lists column properties, such as the name, data type, precision, scale, nullable, and participates in the primary key for each column.

6. In the **Specify Settings** screen, choose to run a column profile, data domain discovery, or a column profile and data domain discovery. By default, column profile option is selected.
 - Choose **Run column profile** to run a column profile.
 - Choose **Run data domain discovery** to perform data domain discovery. In the **Data domain** pane, select the data domains that you want to discover, select a conformance criteria, and select the columns for data domain discovery in the **Edit columns selection for data domain discovery** dialog box.
 - Choose **Run column profile** and **Run data domain discovery** to run the column profile and data domain discovery. Select the data domain options in the **Data domain** pane.

Note: By default, the columns that you select is for column profile and data domain discovery. Click **Edit** to select or deselect columns for data domain discovery.
 - Choose **Data**, **Columns**, or **Data and Columns** to run data domain discovery on.
 - Choose a sampling option. You can choose **All rows (complete analysis)**, **Sample first**, **Random sample**, or **Random sample (auto)** as a sampling option in the **Run profile on** pane. This option applies to column profile and data domain discovery.
 - Choose a drilldown option. You can choose **Live** or **Staged** drilldown option, or you can choose **Off** to disable drilldown in the **Drilldown** pane. Optionally, click **Select Columns** to select columns to drill down on. You can choose to omit data type and data domain inference for columns with an approved data type or data domain.
 - Choose **Native** or **Hadoop** option as the run-time environment. If you choose the Hadoop option, click **Choose** to select a Hadoop connection in the **Select a Hadoop Connection** dialog box.

7. Click **Next**.
The **Specify Rules and Filters** screen opens.
8. In the **Specify Rules and Filters** screen, you can perform the following tasks:
 - Create, edit, or delete a rule. You can apply existing rules to the profile.
 - Create, edit, or delete a filter.

Note: When you create a scorecard on this profile, you can reuse the filters that you create for the profile.
9. Click **Save and Finish** to create the profile, or click **Save and Run** to create and run the profile.

Creating an Enterprise Discovery Profile in Informatica Analyst

You can run column profile and data domain discovery as part of enterprise discovery in Informatica Analyst.

1. In the **Discovery** workspace, select **New > Profile**.
The **New Profile** wizard appears.
2. Select **Enterprise Discovery**. Click **Next**.
The **Specify General Properties** tab appears.
3. In the **Specify General Properties** tab, enter a name for the enterprise discovery profile and an optional description. In the Location field, select the project or folder where you want to create the profile. Click **Next**.
The **Select Data Objects** tab appears.
4. In the **Select Data Objects** tab, click **Choose**.
The **Choose Data objects** dialog box appears.
5. In the **Choose Data objects** dialog box, choose one or more data objects to add to the profile. Click **Save**.
The data objects appear in the **Data Objects** pane.
6. Click **Next**.
The **Select Resources** tab appears.
7. In the **Select Resources** tab, click **Choose** to open the **Select Resources** tab.
You can import data from multiple relational data sources.
8. In the **Select Resources** tab, select the connections, schemas, tables, and views that you want to include in the profile. Click **Save**.
The left pane in the dialog box lists all the internal and external connections, schemas, tables, and views under the Informatica domain.
The resources appear in the **Resource** pane.
9. Click **Next**.
The **Specify Settings** tab appears.
10. In the **Specify Settings** tab, you can configure the column profile options and data domain discovery options. Click **Save and Finish** to save the enterprise discovery profile, or click **Save and Run** to run the profile.

You can perform the following tasks in the **Specify Settings** tab.

- Enable data domain discovery. Click **Choose** to select data domains that you want to discover from the **Choose Data Domains** dialog box. The selected data domains appear in the **Data Domains for Data Domain Discovery** pane.
- Run data domain on data, column name, or on both data and column name.
- Select all the rows in the data source, or choose a maximum number of rows to run domain discovery on.
- Choose a minimum conformance percentage or specify the minimum number of conforming rows for data domain discovery.
- Enable column profile settings and select all rows or first few rows in the data source for the column profile. You can exclude data type inference for columns with approved data types in the column profile.
- Choose **Native** or **Hadoop** as the run-time environment.

You can view the enterprise discovery results under the **Summary** and **Profiles** tabs.

Creating a Scorecard in Informatica Analyst

Create a scorecard and add columns from a profile to the scorecard. You must run a profile before you add columns to the scorecard.

1. In the **Library** workspace, select the project or folder that contains the profile.
2. Click the profile to open the profile.

The profile results appear in the summary view in the **Discovery** workspace.

3. Click **Actions > Add to scorecard**.

The **Add to Scorecard** wizard appears.

4. In the **Add to Scorecard** screen, you can choose to create a new scorecard, or edit an existing scorecard to add the columns to a predefined scorecard. The **New Scorecard** option is selected by default. Click **Next**.

5. In the **Step 2 of 8** screen, enter a name for the scorecard. Optionally, you can enter a description for the scorecard. Select the project and folder where you want to save the scorecard. Click **Next**.

By default, the scorecard wizard selects the columns and rules defined in the profile. You cannot add columns that are not included in the profile.

6. In the **Step 3 of 8** screen, select the columns and rules that you want to add to the scorecard as metrics. Optionally, click the check box in the left column header to select all columns. Optionally, select **Column Name** to sort column names. Click **Next**.

7. In the **Step 4 of 8** screen, you can add a filter to the metric.

You can apply the filter that you created for the profile to the metrics, or create a new filter. Select a metric in the **Metric Filters** pane, and click the **Manage Filters** icon to open the **Edit Filter: column name** dialog box. In the **Edit Filter: column name** dialog box, you can choose to perform one of the following tasks:

- Choose a filter that you created for the profile. Click **Next**.
- Select an existing filter. Click the edit icon to edit the filter in the **Edit Filter** dialog box. Click **Next**.
- Click the plus (+) icon to create filters in the **New Filter** dialog box. Click **Next**.

Optionally, you can choose to apply the selected filters to all the metrics in the scorecard.

The filter appears in the **Metric Filters** pane.

8. In the **Step 4 of 8** screen, click **Next**.
9. In the **Step 5 of 8** screen, select each metric in the **Metrics** pane to perform the following tasks:
 - Configure valid values. In the **Score using: Values** pane, select one or more values in the **Available Values** pane, and click the right arrow button to move them to the **Valid Values** pane. The total number of valid values for a metric appears at the top of the **Available Values** pane.
 - Configure metric thresholds. In the **Metric Thresholds** pane, set the thresholds for **Good**, **Acceptable**, and **Unacceptable** scores.
 - Configure the cost of invalid data. To assign a constant value to the cost for the metric, select **Fixed Cost**. To attach a numeric column as a variable cost to the metric, select **Variable Cost**, and click **Select Column** to select a numeric column. Optionally, click **Change Cost Unit** to change the unit of cost. If you do not want to configure the cost of invalid data for the metric, choose **None**.
10. Click **Next**.
11. In the **Step 6 of 8** screen, you can select a metric group to which you can add the metrics, or create a new metric group. To create a new metric group, click the group icon. Click **Next**.
12. In the **Step 7 of 8** screen, specify the weights for the metrics in the group and thresholds for the group.
13. In the **Step 8 of 8** screen, select **Native** or **Hadoop** run-time environment option to run the scorecard. If you choose the Hadoop option, click **Browse** to choose a Hadoop connection to run the profile on the Blaze engine.
14. Click **Save** to save the scorecard, or click **Save & Run** to save and run the scorecard.

The scorecard appears in the **Scorecard** workspace.

Monitoring a Profile

You can monitor a profile in the Administrator tool.

1. In the Administrator tool, click the **Monitor** tab.
2. In the Navigator workspace, select **Jobs**.
3. Select a profiling job.
4. In the **Summary Statistics** tab, you can view the general properties of the profile, summary statistics, and detailed statistics of the profile.
5. Click the **Execution Statistics** tab to view execution statistics for the profile.

Profiling Functionality Support

You can connect to the profiling warehouse through JDBC or native connectivity. You can perform specific functionality in profiling based on the profiling warehouse connection.

The following table lists the functionality that you can perform based on the type of profiling warehouse connection that you choose in Big Data Quality and Big Data Management:

Functionality	JDBC Connection	Native Connection
Single data object column profile	Supported	Supported
Single data object data domain discovery	Supported	Supported
Single data object profile with primary key discovery	Not supported	Not supported
Single data object profile with functional dependency discovery	Not supported	Not supported
Profile with <i>Sample first <number> rows</i> sampling	Supported	Supported
Profile with <i>Random sample <number> rows</i> sampling	Supported	Supported
Scorecard metrics and value cost	Not supported	Supported
Enterprise discovery profile with primary key and foreign key discovery	Not supported	Not supported
Enterprise discovery profile with join analysis and overlap discovery	Not supported	Not supported
Drill down on inferred values and value frequencies	Supported	Supported
Export profile results	Supported	Supported
Row count mapping*	Not supported	Supported
*Profiling runs Row Count Mapping when the stats helper fails to report the row count.		

Troubleshooting

Can I drill down on profile results if I run a profile in the Hadoop environment?

Yes, except for profiles in which you have set the option to drill down on staged data.

I get the following error message when I run a profile in the Hadoop environment: "[LDTM_1055] The Integration Service failed to generate a Hive workflow for mapping [Profile_CUSTOMER_INFO12_14258652520457390]." How do I resolve this?

This error can result from a data source, rule transformation, or run-time environment that is not supported in the Hadoop environment. For more information about objects that are not valid in the Hadoop environment, see the Mappings in a Hadoop Environment chapter.

You can change the data source, rule, or run-time environment and run the profile again. View the profile log file for more information on the error.

I see "N/A" in the profile results for all columns after I run a profile. How do I resolve this?

Verify that the profiling results are in the profiling warehouse. If you do not see the profile results, verify that the database path is accurate in the Cluster Environment Variables property of the Hadoop connection. You can also verify the database path from the Hadoop job tracker on the Monitoring tab of the Administrator tool.

After I run a profile on a Hive source, I do not see the results. When I verify the Hadoop job tracker, I see the following error when I open the profile job: "XML Parsing Error: no element found." What does this mean?

The Hive data source does not have any record and is empty. The data source must have a minimum of one row of data for successful profile run.

After I run a profile on a Hive source, I cannot view some of the column patterns. Why?

When you import a Hive source, the Developer tool sets the precision for string columns to 4000. The Developer tool cannot derive the pattern for a string column with a precision greater than 255. To resolve this issue, set the precision of these string columns in the data source to 255 and run the profile again.

When I run a profile on large Hadoop sources, the profile job fails and I get an "execution failed" error. What can be the possible cause?

One of the causes can be a connection issue. Perform the following steps to identify and resolve the connection issue:

1. Open the Hadoop job tracker.
2. Identify the profile job and open it to view the MapReduce jobs.
3. Click the hyperlink for the failed job to view the error message. If the error message contains the text "java.net.ConnectException: Connection refused", the problem occurred because of an issue with the Hadoop cluster. Contact your network administrator to resolve the issue.

CHAPTER 8

Monitoring

This chapter includes the following topics:

- [Overview of Monitoring, 135](#)
- [Hadoop Environment Logs, 135](#)
- [Blaze Engine Monitoring, 141](#)
- [Spark Engine Monitoring, 149](#)

Overview of Monitoring

On the Monitor tab of the Administrator tool, you can view statistics and log events for mappings run in the Hadoop environment. The Monitor tab displays current and historical information about mappings run on the Blaze and Spark engines.

Use the Summary Statistics view to view graphical summaries of object state and distribution across the Data Integration Services. You can also view graphs of the memory and CPU that the Data Integration Services used to run the objects.

When you run a mapping in the Hadoop environment, the Data Integration Service generates log events. You can view log events relating to different types of errors such as Hadoop connection failures, Hive query failures, or other Hadoop job failures.

With the REST Operations Hub, you can also get monitoring statistics for deployed mappings that run on the Data Integration Service or the Blaze or Spark engines.

Note: You cannot monitor the Databricks engine on the Monitoring tab.

Hadoop Environment Logs

The Data Integration Service generates log events when you run a mapping in the Hadoop environment.

You can view logs for the Blaze and Spark engines. You can view log events relating to different types of errors such as Hadoop connection or job failures.

When you run a mapping on the Spark engine, you can view the Scala code in logs that the Logical Data Translation Generator generates from the Informatica mapping.

You can view reject files in the reject file directory specified for the Data Integration Service.

YARN Web User Interface

You can view the applications that ran on a cluster in the YARN web user interface. Click the Monitoring URL for Blaze or Spark jobs to access the YARN web user interface.

Blaze and Spark engines run on the Hadoop cluster that you configure in the Hadoop connection. The YARN web user interface shows each job that the engine runs as a YARN application.

The following image shows the Application Monitoring page of the YARN web user interface:

The screenshot shows the Hadoop YARN Web User Interface. The sidebar on the left contains navigation links: Cluster, About, Nodes, Applications, NEW, NEW SAVING, SUBMITTED, ACCEPTED, RUNNING, FINISHED, FAILED, KILLED, Scheduler, and Tools. The main content area is titled 'All Applications' and displays 'Cluster Metrics' and 'User Metrics for dr.who'. The 'User Metrics' table shows application details including ID, User, Name, Application Type, Queue, StartTime, FinishTime, State, FinalStatus, Running Containers, and Allocated CPU VCoers.

ID	User	Name	Application Type	Queue	StartTime	FinishTime	State	FinalStatus	Running Containers	Allocated CPU VCoers
application_1463379223882_0568	ldmui	InfSprk0	SPARK	root.ldmui	Mon May 16 13:11:59 -0700 2016	Mon May 16 13:13:03 -0700 2016	FINISHED	SUCCEEDED	N/A	N/A
application_1463379223882_0567	ldmui	InfSprk0	SPARK	root.ldmui	Mon May 16 13:11:58 -0700 2016	Mon May 16 13:13:01 -0700 2016	FINISHED	SUCCEEDED	N/A	N/A
application_1463379223882_0566	ldmui	InfSprk0	SPARK	root.ldmui	Mon May 16 13:11:56 -0700 2016	Mon May 16 13:13:00 -0700 2016	FINISHED	SUCCEEDED	N/A	N/A
application_1463379223882_0565	ldmui	InfSprk0	SPARK	root.ldmui	Mon May 16 13:11:55 -0700 2016	Mon May 16 13:12:59 -0700 2016	FINISHED	SUCCEEDED	N/A	N/A
application_1463379223882_0564	ldmui	InfSprk0	SPARK	root.ldmui	Mon May 16 13:11:54 -0700 2016	Mon May 16 13:12:59 -0700 2016	FINISHED	SUCCEEDED	N/A	N/A

The **Application Type** indicates which engine submitted the YARN application.

The application ID is the unique identifier for the application. The application ID is a link to the application summary. The URL is the same as the Monitoring URL in the Administrator tool.

Click the **Logs** link in the application summary to view the application logs on the Hadoop cluster.

The amount of information in the application logs depends on the tracing level that you configure for a mapping in the Developer tool. The following table describes the amount of information that appears in the application logs for each tracing level:

Tracing Level	Messages
None	The log displays FATAL messages. FATAL messages include non-recoverable system failures that cause the service to shut down or become unavailable.
Terse	The log displays FATAL and ERROR code messages. ERROR messages include connection failures, failures to save or retrieve metadata, service errors.
Normal	The log displays FATAL, ERROR, and WARNING messages. WARNING errors include recoverable system failures or warnings.
Verbose initialization	The log displays FATAL, ERROR, WARNING, and INFO messages. INFO messages include system and service change messages.
Verbose data	The log displays FATAL, ERROR, WARNING, INFO, and DEBUG messages. DEBUG messages are user request logs.

Accessing the Monitoring URL

The Monitoring URL opens the Blaze Job Monitor web application or the YARN web user interface. Access the Monitoring URL from the **Execution Statistics** view in the Administrator tool.

1. In the **Monitor** tab of the Administrator tool, click the **Execution Statistics** view.
2. Select **Ad Hoc Jobs** or select a deployed mapping job or workflow from an application in the Navigator.
The list of jobs appears in the contents panel.
3. Select a mapping job and expand the mapping to select a grid task for the mapping.
The Monitoring URL appears in the **Properties** view.

The screenshot displays the 'Ad Hoc Jobs' interface in the Administrator tool. At the top, there's a header with 'Ad Hoc Jobs' and some navigation icons. Below it is a table listing jobs. The table has columns: Name, Type, State, Job ID, Started By, Start Time, Elapsed Time, and End Time. One job, 'MAINSESSION_task2', is highlighted in blue. Below the table, there's a section titled 'MAINSESSION_task2 - T2GJgmceEeWPuPKvpC0s5Q_MAINSESSION_task2'. This section shows a status message 'This grid task is completed.' and a 'General Properties' section. The 'General Properties' section lists various details about the task, including its name, type, started by, user security domain, start time, elapsed time, end time, percentage of task completed, and the Monitoring URL. The Monitoring URL is a long, complex string starting with 'http://psrhaqadn21.informatica.com:9080/Blaze?tasktype=gridtask&id=qtid-24-1-79555597-4&isParent=false'.

Name	Type	State	Job ID	Started By	Start Time	Elapsed Time	End Time
PassThrough	Mapping	Completed	T2GJgmceE...	Administrator	09/29/2015 19:52:38	00:01:32	09/29/2015 19:54:10
POSTSES...	Command ...	Completed	T2GJgmceE...	Administrator	09/29/2015 19:53:41	00:00:17	09/29/2015 19:53:58
MAINSESS...	Grid Task	Completed	T2GJgmceE...	Administrator	09/29/2015 19:52:57	00:00:43	09/29/2015 19:53:41
PRESESSI...	Command ...	Completed	T2GJgmceE...	Administrator	09/29/2015 19:52:38	00:00:02	09/29/2015 19:52:41

Showing 33 results. ☒ Receive New Job Notifications

MAINSESSION_task2 - T2GJgmceEeWPuPKvpC0s5Q_MAINSESSION_task2

☒ This grid task is completed.

General Properties

Name	MAINSESSION_task2
Type	Grid Task
Started By	Administrator
User Security Domain	Native
Start Time	09/29/2015 19:52:57
Elapsed Time	00:00:43
End Time	09/29/2015 19:53:41
% Task Completed	100
Monitoring URL	http://psrhaqadn21.informatica.com:9080/Blaze?tasktype=gridtask&id=qtid-24-1-79555597-4&isParent=false
Incoming Task Dependencies	, PRESESSION_task0PRESESSION_task1
Outgoing Task Dependencies	, POSTSESSION_task3

Viewing Hadoop Environment Logs in the Administrator Tool

You can view log events for a Blaze mapping from the Monitor tab of the Administrator tool.

1. In the Administrator tool, click the **Monitor** tab.
2. Select the **Execution Statistics** view.
3. In the Navigator, choose to open an ad hoc job, a deployed mapping job, or a workflow.
 - To choose an ad hoc job, expand a Data Integration Service and click **Ad Hoc Jobs**.
 - To choose a deployed mapping job, expand an application and click **Deployed Mapping Jobs**.
 - To choose a workflow, expand an application and click **Workflows**.

The list of jobs appears in the contents panel.

- Click **Actions > View Logs for Selected Object** to view the run-time logs for the mapping.

The log file shows the results of the Hive queries and Blaze engine queries run by the Data Integration Service. This includes the location of Hive session logs and Hive session history file.

Monitoring a Mapping

You can monitor a mapping that runs in the Hadoop environment.

- In the Administrator tool, click the **Monitor** tab.
- Select the **Execution Statistics** view.
- In the Navigator, choose to open an ad hoc job, a deployed mapping job, or a workflow.
 - To choose an ad hoc job, expand a Data Integration Service and click **Ad Hoc Jobs**.

Informatica Administrator

Administrator (Native) | Manage | Help

Manage | **Monitor** | Logs | Reports | Security | Cloud

Summary Statistics | **Execution Statistics**

Navigator | Actions

Ad Hoc Jobs

Name	Type	State	Job ID	Started By	Start Time	Elapsed Time	End Time
m_FF_FF	Mapping	Completed	xL8TmEEi6Q...	Administrator	04/06/2018 16:23:44	00:01:58	04/06/2018 16:25:43
InfoSpark0	Spark Applic...	Completed	xL8TmEEi6Q...	Administrator	04/06/2018 16:23:51	00:01:50	04/06/2018 16:25:42
m_FF_FF	Mapping	Completed	vWigGcmEEi6...	Administrator	04/06/2018 16:23:29	00:01:43	04/06/2018 16:25:12
m_FF_FF	Mapping	Completed	Ibq-RDmEEi6...	Administrator	04/06/2018 16:23:17	00:01:25	04/06/2018 16:24:42
m_FF_FF	Mapping	Completed	rduJzmEEi6...	Administrator	04/06/2018 16:23:04	00:01:38	04/06/2018 16:24:42
m_FF_FF	Mapping	Completed	kDdLTmEEi6...	Administrator	04/06/2018 16:00:48	00:02:41	04/06/2018 16:03:30
m_FF_FF	Mapping	Completed	p1MjUmEEi6...	Administrator	04/06/2018 15:54:14	00:03:08	04/06/2018 15:57:23
m_FF_FF	Mapping	Completed	QBUzmEEi6...	Administrator	04/06/2018 15:51:25	00:01:46	04/06/2018 15:53:12
m_FF_FF	Mapping	Failed	QAKYpTmEE...	Administrator	04/06/2018 15:51:07	00:00:08	04/06/2018 15:51:16
m_FF_FF	Mapping	Completed	D8hNocmEEi...	Administrator	04/06/2018 15:50:00	00:00:43	04/06/2018 15:50:43

Showing 26 results

m_FF_FF - m_FF_FF | **Properties** | Spark Execution Plan | Summary Statistics | Detailed Statistics | Historical Statistics

The Mapping job has completed.

General Properties

Name	m_FF_FF
Type	Mapping
Started By	Administrator
User Security Domain	Native
Start Time	04/06/2018 16:23:44
Elapsed Time	00:01:58
End Time	04/06/2018 16:25:43
Operating System Profile	OSP_ospUser4

- To choose a deployed mapping job, expand an application and click **Deployed Mapping Jobs**.

Informatica Administrator Administrator (Native) Manage Help

Manage Monitor Logs Reports Security Cloud

Summary Statistics Execution Statistics

Navigator Actions

Domain331

- DIS
- DIS_EMR_5_10
 - Ad Hoc Jobs
 - BDMUI
 - Spark_Tez_Monitoring
 - Deployed Mapping Jobs**
 - Logical Data Objects
 - REST Web Services
 - SOL Data Services
 - Web Services
 - Workflows
- EMR_5_10
- Test
 - Ad Hoc Jobs

Deployed Mapping Jobs

Name	Type	State	Job ID	Started By	Start Time	Elapsed Time	End Time
m_Ustx_Readback	Deployed Mapping	Completed	Egk2Jym0Eeil...	Administrator	04/06/2018 21:33:39	00:08:17	04/06/2018 21:41:57
▼ exec0	Script	Completed	Egk2Jym0Eeil...	Administrator	04/06/2018 21:35:53	00:02:26	04/06/2018 21:38:11
▼ exec0_qu...	Hive Query	Completed	svcoqa_2018...	Administrator	04/06/2018 21:36:34	00:01:06	04/06/2018 21:37:40
▼ exec0_qu...	Hive Query	Completed		Administrator	04/06/2018 21:36:30	00:00:04	04/06/2018 21:36:34
▼ TST_Monitoring...	Deployed Mapping	Completed	RrsjTmcEeil...	Administrator	04/06/2018 21:28:00	00:05:28	04/06/2018 21:33:28
▼ m_Ustx_Readback	Deployed Mapping	Completed	HOCMBDmy...	Administrator	04/06/2018 21:19:36	00:08:02	04/06/2018 21:27:38
▼ TST_Monitoring...	Deployed Mapping	Completed	TkRm0ZmwE...	Administrator	04/06/2018 21:13:52	00:05:36	04/06/2018 21:19:28
▼ m_Ustx_Readback	Deployed Mapping	Completed	FFQz3zmvE...	Administrator	04/06/2018 21:05:05	00:08:29	04/06/2018 21:13:34
▼ TST_Monitoring...	Deployed Mapping	Completed	OTqVzmvEeil...	Administrator	04/06/2018 20:59:02	00:05:54	04/06/2018 21:04:56
▼ hbm stress 100	Dashboard Monitor	Completed	KKTVZDm0F...	Administrator	04/06/2018 20:48:29	00:00:12	04/06/2018 20:48:41

Showing 54 results

Receive New Job Notifications

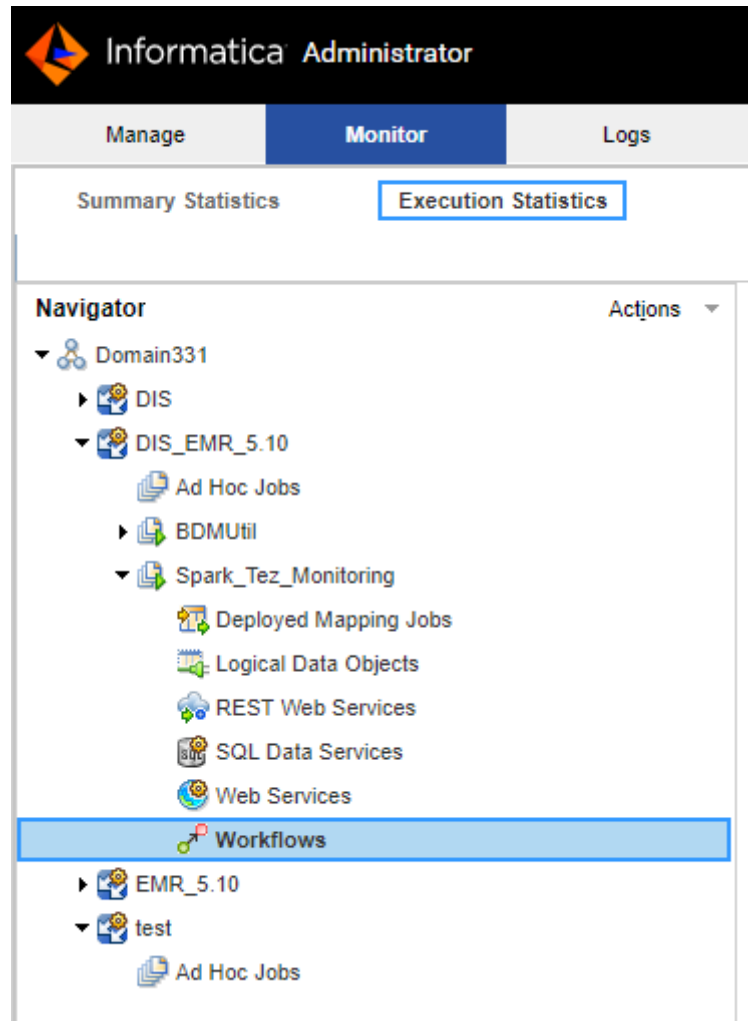
m_Ustx_Readback - Egk2Jym0Eeil... Properties Hive Execution Plan Summary Statistics Detailed Statistics Historical Statistics

The Deployed Mapping job has completed.

General Properties

Name	m_Ustx_Readback
Type	Deployed Mapping
Started By	Administrator
User Security Domain	Native
Start Time	04/06/2018 21:33:39
Elapsed Time	00:08:17
End Time	04/06/2018 21:41:57

- To choose a workflow, expand an application and click **Workflows**.



The list of jobs appears in the contents panel.

4. Click a job to view its properties.

The contents panel shows the default **Properties** view for the job. For a Blaze engine mapping, the Blaze engine monitoring URL appears in the general properties in the details panel. The monitoring URL is a link to the YARN web user interface for Spark jobs.

5. Choose a view in the contents panel to view more information about the job:
 - To view the execution plan for the mapping, select the **Execution Plan** view.
 - To view the summary statistics for a job, click the **Summary Statistics** view.
 - To view the detailed statistics for a job, click the **Detailed Statistics** view.

Note: You can view the number of rows processed in the Summary Statistics for a Hive source or target. The remaining values do not appear for Hive sources and targets.

The following table describes the mapping job states in the Administrator tool contents panel:

Job Status	Rules and Guidelines
Queued	The job is in the queue.
Running	The Data Integration Service is running the job.
Completed	The job ran successfully.
Aborted	The job was flushed from the queue at restart or the node shut down unexpectedly while the job was running.
Failed	The job failed while running or the queue is full.
Canceled	The job was deleted from the queue or cancelled while running.
Unknown	The job status is unknown.

Blaze Engine Monitoring

You can monitor statistics and view log events for a Blaze engine mapping job in the Monitor tab of the Administrator tool. You can also monitor mapping jobs for the Blaze engine in the Blaze Job Monitor web application.

The following image shows the Monitor tab in the Administrator tool:

The screenshot displays the Administrator tool interface with the Monitor tab selected. The interface includes a top navigation bar with tabs: Manage, Monitor (active), Logs, Reports, Security, and Cloud. Below this is a sub-navigation bar with 'Summary Statistics' and 'Execution Statistics' (active). On the left is a 'Navigator' pane showing a tree structure with 'domain776', 'DIS_HW', 'Ad Hoc Jobs', and 'Application_Mapping'. The main area shows a table of 'Ad Hoc Jobs' with columns: Name, Type, State, Job ID, Started By, Start Time, Elapsed Time, and End Time. The table lists several jobs, with the first one 'MappingSrcPa...' being highlighted. Below the table, there are tabs for 'MappingSrcPartitioned - zWp7Cmb...', 'Properties', 'Blaze Execution Plan', 'Summary Statistics', and 'Detailed Statistics'. The 'Properties' tab is active, showing a 'General Properties' section with details for the selected job. Numbered callouts 1 through 5 point to specific UI elements: 1. Navigator, 2. View in the tab, 3. Contents panel, 4. View in the contents panel, and 5. Details panel.

Name	Type	State	Job ID	Started By	Start Time	Elapsed Time	End Time
MappingSrcPa...	Mapping	Completed	zWp7Cmb...	Administrator	09/29/2015 11:20:46	00:00:49	09/29/2015 11:21:35
FFMapping	Mapping	Failed	ZsAcSGcN...	Administrator	09/29/2015 17:51:36	00:00:00	09/29/2015 17:51:36
OrderSupport...	Mapping	Failed	ZoSoFWbG...	Administrator	09/29/2015 09:23:21	00:03:33	09/29/2015 09:26:55
MappingSrcPa...	Mapping	Completed	zD02WGbc...	Administrator	09/29/2015 12:03:41	00:00:01	09/29/2015 12:03:42
MappingSrcPa...	Mapping	Completed	Zby2dWbdE...	Administrator	09/29/2015 12:07:58	00:00:00	09/29/2015 12:07:59

General Properties	
Name	MappingSrcPartitioned
Type	Mapping
Started By	Administrator
User Security Domain	Native
Start Time	09/29/2015 11:20:46
Elapsed Time	00:00:49
End Time	09/29/2015 11:21:35

1. Navigator
2. View in the tab
3. Contents panel
4. View in the contents panel
5. Details panel

The Monitor tab has the following views:

Summary Statistics

Use the **Summary Statistics** view to view graphical summaries of object states and distribution across the Data Integration Services. You can also view graphs of the memory and CPU that the Data Integration Services used to run the objects.

Execution Statistics

Use the **Execution Statistics** view to monitor properties, run-time statistics, and run-time reports. In the Navigator, you can expand a Data Integration Service to monitor **Ad Hoc Jobs** or expand an application to monitor deployed mapping jobs or workflows

When you select **Ad Hoc Jobs**, deployed mapping jobs, or workflows from an application in the Navigator of the **Execution Statistics** view, a list of jobs appears in the contents panel. The contents panel displays jobs that are in the queued, running, completed, failed, aborted, and cancelled state. The Data Integration Service submits jobs in the queued state to the cluster when enough resources are available.

The contents panel groups related jobs based on the job type. You can expand a job type to view the related jobs under it.

Access the following views in the **Execution Statistics** view:

Properties

The **Properties** view shows the general properties about the selected job such as name, job type, user who started the job, and start time of the job. You can also monitor jobs on the Hadoop cluster from the Monitoring URL that appears for the mapping in the general properties. The Monitoring URL opens the Blaze Job Monitor in a web page. The Blaze Job Monitor displays detailed monitoring statistics for a mapping such as the number of grid tasks, grid segments, or tasklets, and recovery attempts for each tasklet.

Blaze Execution Plan

The Blaze execution plan displays the Blaze engine script that the Data Integration Service generates based on the mapping logic. The execution plan includes the tasks that the script depends on. Each script has a unique identifier.

Summary Statistics

The **Summary Statistics** view appears in the details panel when you select a mapping job in the contents panel. The **Summary Statistics** view displays throughput and resource usage statistics for the job.

You can view the following throughput statistics for the job:

- Source. The name of the mapping source file.
- Target name. The name of the target file.
- Rows. The number of rows read for source and target. If the target is Hive, this is the only summary statistic available.
- Average Rows/Sec. Average number of rows read per second for source and target.
- Bytes. Number of bytes read for source and target.
- Average Bytes/Sec. Average number of bytes read per second for source and target.
- First Row Accessed. The date and time when the Data Integration Service started reading the first row in the source file.
- Dropped rows. Number of source rows that the Data Integration Service did not read.

Detailed Statistics

The **Detailed Statistics** view appears in the details panel when you select a mapping job in the contents panel. The **Detailed Statistics** view displays graphs of the throughput and resource usage statistics for the job run.

Blaze Job Monitoring Application

Use the Blaze Job Monitor application to monitor Blaze engine jobs on the Hadoop cluster.

You configure the host that starts the Blaze Job Monitor in the Hadoop connection properties. You might want to configure the Blaze Job Monitor address to avoid conflicts with other users on the same cluster, or if you have access to a limited number of nodes. If you do not configure the Blaze Job Monitor address, the Grid Manager starts the host on the first alphabetical cluster node with a default port of 9080.

The Blaze engine monitoring URL appears in the Monitor tab of the Administrator tool when you view a Blaze engine mapping job. When you click the URL, the Blaze engine monitoring application opens in a web page.

Note: You can also access the Blaze Job Monitor through the LDTM log. After the session load summary, the log displays a list of segments within the grid task. Each segment contains a link to the Blaze Job Monitor. Click on a link to see the execution details of that segment.

You configure the host that starts the Blaze Job Monitor in the Hadoop connection properties. The default address is <hostname>:9080.

The following image shows the Blaze Job Monitor:

Name	Start Time	End Time	Elapsed Time	Scale	Status	Host Name	Log
gms-499-1-42938595-32_s6_1_0_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40		Succeeded	psrhaqadr21.informatica.com	Log
gms-499-1-42938595-32_s6_1_1_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40		Succeeded	psrhaqadr26.informatica.com	Log
gms-499-1-42938595-32_s6_1_1_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:59 PM	0:1:52		Succeeded	psrhaqadr23.informatica.com	Log
gms-499-1-42938595-32_s6_1_2_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40		Succeeded	psrhaqadr26.informatica.com	Log
gms-499-1-42938595-32_s6_1_0_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40		Succeeded	psrhaqadr26.informatica.com	Log
gms-499-1-42938595-32_s6_1_0_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40		Succeeded	psrhaqadr21.informatica.com	Log
gms-499-1-42938595-32_s6_1_2_1	Mon Oct 31 2016 2:45:04 PM	Mon Oct 31 2016 2:45:47 PM	0:0:43		Succeeded	psrhaqadr21.informatica.com	Log
gms-499-1-42938595-32_s1_1_1_1	Mon Oct 31 2016 2:45:04 PM	Mon Oct 31 2016 2:45:07 PM	0:0:3		Succeeded	psrhaqadr26.informatica.com	Log

Use the **Task History** panel on the left to filter Blaze mapping jobs by the following criteria:

- Grid task. A parallel processing job request sent by the Blaze engine executor to the Grid Manager. You can further filter by all tasks, succeeded tasks, running tasks, or failed tasks.
- Grid segment. Part of a grid mapping that is contained in a grid task.
- Tasklet. A partition of a grid segment that runs on a separate DTM.
- Tasklet Attempts. The number of recovery attempts to restart a tasklet. Click **Log** to view the mapping grid task log.

The Blaze Job Monitor displays the task history for mapping jobs with the same namespace. You can monitor properties for a task such as start time, end time, elapsed time, or state of the task. You can also view log events. If you filter mapping jobs by grid segment, you can mouse over a grid segment to view the logical name of the segment.

By default, the Blaze Job Monitor automatically refreshes the list of tasks every five seconds and reverts to the first page that displays tasks. Disable auto refresh if you want to browse through multiple pages. To turn off automatic refresh, click **Action > Disable Auto Refresh**.

The Blaze Job Monitor displays the first 100,000 grid tasks run in the past seven days. The Blaze Job Monitor displays the grid segments, tasklets, and tasklet attempts for grid tasks that are running and grid tasks that were accessed in the last 30 minutes.

Blaze Summary Report

The Blaze Summary Report displays more detailed statistics about a mapping job. In the Blaze Job Monitor, a green summary report button appears beside the names of successful grid tasks. Click the button to open the Blaze Summary Report.

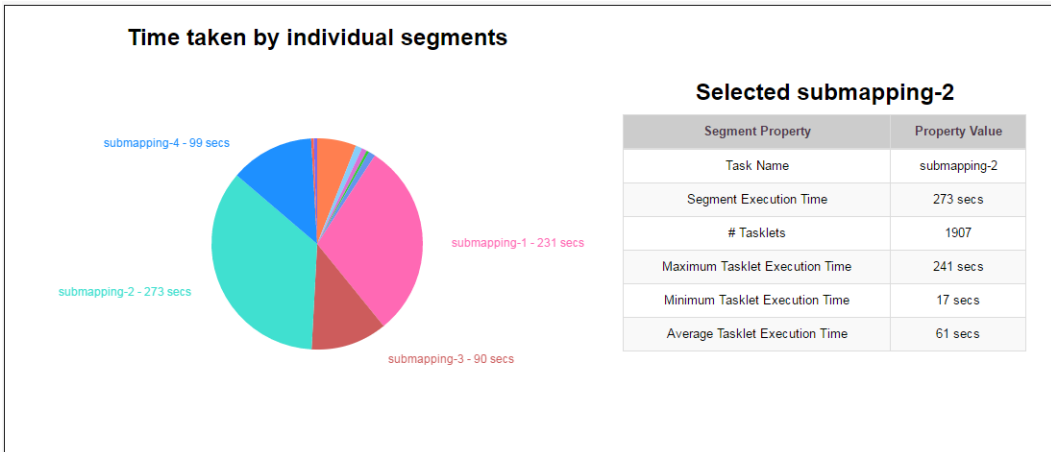
Note: The Blaze Summary Report is available for technical preview. Technical preview functionality is supported but is unwarranted and is not production-ready. Informatica recommends that you use in non-production environments only.

All Grid Tasks

Show 25 entries					
Name	Start Time	End Time	Elapsed Time	State	
gtid-476-1-36233666-2	Mon Oct 17 2016 1:32:42 PM	Mon Oct 17 2016 1:53:18 PM	0:20:36	Succeeded	
gtid-476-1-36233666-1	Summary Report 2016 1:30:51 PM	Mon Oct 17 2016 1:31:20 PM	0:0:28	Succeeded	
gtid-441-1-26795155-4	Mon Oct 17 2016 11:55:06 AM	Mon Oct 17 2016 11:59:44 AM	0:4:37	Succeeded	
gtid-441-1-26795155-3	Mon Oct 17 2016 11:47:10 AM	Mon Oct 17 2016 11:51:51 AM	0:4:40	Succeeded	
gtid-441-1-26795155-2	Mon Oct 17 2016 11:02:37 AM	Mon Oct 17 2016 11:06:18 AM	0:3:40	Succeeded	
gtid-441-1-26795155-1	Mon Oct 17 2016 10:53:35 AM	Mon Oct 17 2016 10:54:27 AM	0:0:51	Failed	
gtid-437-1-25270758-1	Mon Oct 17 2016 10:28:08 AM	Mon Oct 17 2016 10:28:56 AM	0:0:47	Failed	

Time Taken by Individual Segments

A pie chart visually represents the time taken by individual segments contained within a grid task.



When you click on a particular segment in the pie chart, the **Selected Submapping** table displays detailed information about that segment. The table lists the following segment statistics:

- Task Name. The logical name of the selected segment.
- Segment Execution Time. The time taken by the selected segment.
- # Tasklets. The number of tasklets in the selected segment.
- Minimum Tasklet Execution Time. The execution time of the tasklet within the selected segment that took the shortest time to run.

- Maximum Tasklet Execution Time. The execution time of the tasklet within the selected segment that took the longest time to run.
- Average Tasklet Execution Time. The average execution time of all tasklets within the selected segment.

Mapping Properties

The Mapping Properties table lists basic information about the mapping job.

Mapping Properties	
Mapping Property	Property Value
DIS Name	dis_cdh
Informatica Version	10.1.1
Mapping Name	q97_hive_mapping
Total Segments	13
Maximum Segment Execution Time	273 secs
Minimum Segment Execution Time	0 secs
Average Segment Execution Time	59 secs
Mapping Execution Time	0:10:14

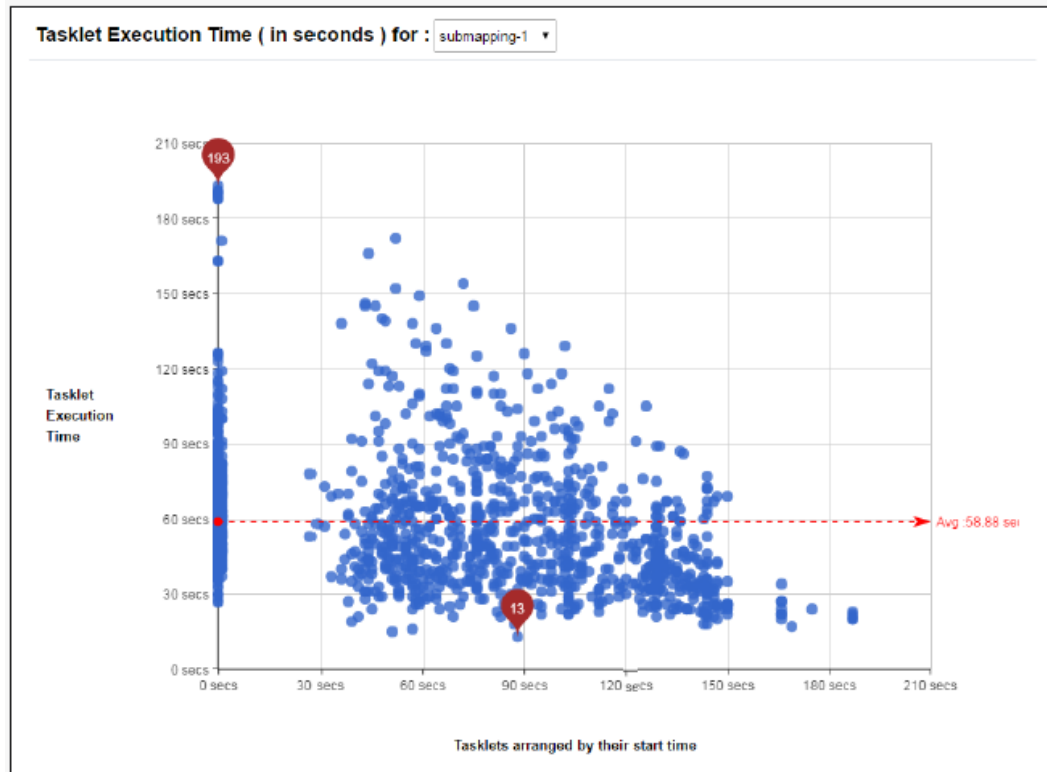
The Mapping Properties table displays the following information:

- The DIS name under which the mapping was run.
- The Informatica version.
- The name of the mapping.
- The total number of segments for the mapping.
- The execution time of the segment that took the longest time to run.
- The execution time of the segment that took the shortest time to run.
- The average execution time of all segments.
- The total mapping execution time.

Tasklet Execution Time

A time series graph displays the execution time of all tasklets within the selected segment.

The x-axis represents the tasklet start time and the y-axis represents the actual tasklet execution time. The red dashed line represents the average execution time for all tasklets, and the two red markers show the minimum and maximum execution times within the segment.



Selected Tasklet Information

When you select a tasklet from the **Tasklet Execution Time** graph, you can see more data about that individual tasklet. This data includes source and target row counts as well as cache information for any

cache-based transformation processed by the tasklet. Click the **Get Detailed Log** button to see a full log of the selected tasklet.

Selected tasklet `gtid-299-1-82064486-16_s8_t-394_1`

[Get Detailed Log](#)

Show entries Search:

Source Name	# Rows Processed
empty_source1	0
Read_catalog_sales	4,937,484

Showing 1 to 2 of 2 entries Previous Next

Row counts for all sources processed by tasklet

Show entries Search:

Target Name	# Rows Processed
DETarget_Aggregator1_G0	0

Showing 1 to 1 of 1 entries Previous Next

Row counts for all targets written by tasklet

Show entries Search:

Transformation Name	Index Cache (bytes)		Data Cache (bytes)	
	Configured	Used	Configured	Used
Joiner1	178,956,970	14,400	357,913,940	6,968

Showing 1 to 1 of 1 entries Previous Next

Cache information for any cache-based transformation processed by tasklet

Blaze Engine Logs

You can find information about the mapping run on the Blaze engine in the LDTM log file.

The LDTM logs the results of the mapping run on the Blaze engine. You can view the LDTM log from the Developer tool or the Monitoring tool for a mapping job.

You can configure the Data Integration Service to log details about the mapping execution to the session log. To enable logging of LDTM mapping execution details, set the log tracing level to verbose initialization or verbose data. For optimal logging performance, avoid setting the tracing level to verbose data. You can set the tracing level to verbose data if you need to debug a mapping run.

Logical mapping execution details include the following information:

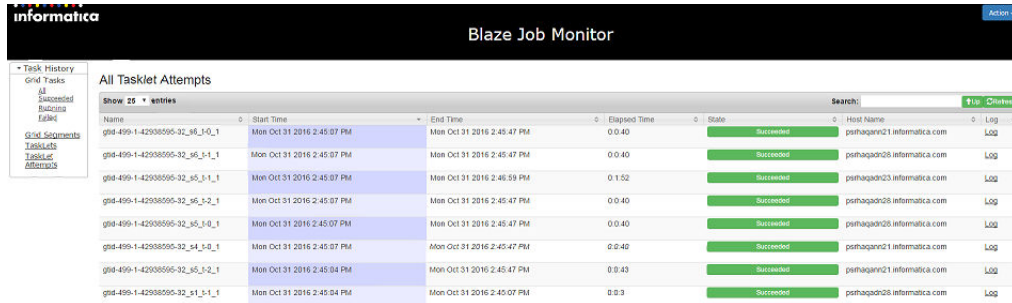
- Start time, end time, and state of each task
- Blaze Job Monitor URL
- Number of total, succeeded, and failed/cancelled tasklets
- Number of processed and rejected rows for sources and targets
- Data errors, if any, for transformations in each executed segment

The mapping execution details include information about the Blaze engine run-time processing, such as the partition count for each executed segment.

Viewing Blaze Logs

You can view logs for a Blaze mapping from the Blaze Job Monitor.

1. In the Blaze Job Monitor, select a job from the list of jobs.
2. In the row for the selected job, click the **Logs** link.



The screenshot shows the Informatica Blaze Job Monitor interface. On the left is a navigation menu with options like Task History, Grid Tasks, and Job Details. The main area is titled 'All Tasklet Attempts' and displays a table with columns: Name, Start Time, End Time, Elapsed Time, State, Host Name, and Log. The table lists several task attempts, all with a 'Succeeded' state. A search bar and a 'Show 25 entries' dropdown are at the top right of the table.

Name	Start Time	End Time	Elapsed Time	State	Host Name	Log
gms-409-1-42908095-32_r6_3-0_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40	Succeeded	psrhaqad21.informatica.com	Log
gms-409-1-42908095-32_r6_3-1_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40	Succeeded	psrhaqad26.informatica.com	Log
gms-409-1-42908095-32_r6_3-1_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:59 PM	0:1:52	Succeeded	psrhaqad23.informatica.com	Log
gms-409-1-42908095-32_r6_3-2_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40	Succeeded	psrhaqad26.informatica.com	Log
gms-409-1-42908095-32_r6_3-0_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:0:40	Succeeded	psrhaqad26.informatica.com	Log
gms-409-1-42908095-32_r4_1-0_1	Mon Oct 31 2016 2:45:07 PM	Mon Oct 31 2016 2:45:47 PM	0:2:40	Succeeded	psrhaqad21.informatica.com	Log
gms-409-1-42908095-32_r6_3-2_1	Mon Oct 31 2016 2:45:04 PM	Mon Oct 31 2016 2:45:47 PM	0:0:43	Succeeded	psrhaqad21.informatica.com	Log
gms-409-1-42908095-32_r1_1-1_1	Mon Oct 31 2016 2:45:04 PM	Mon Oct 31 2016 2:45:07 PM	0:0:3	Succeeded	psrhaqad26.informatica.com	Log

The log events appear in another browser window.

Orchestrator Sunset Time

Orchestrator sunset time is the maximum lifetime for an Orchestrator service. Sunset time determines the maximum amount of time that the Blaze engine can run a mapping job. The default sunset time is 24 hours. After 24 hours, the Orchestrator shuts down, which causes the Blaze Grid Manager to shut down.

You can configure the Orchestrator sunset time to be greater than or less than 24 hours. Configure the following property in the Hadoop connection:

Property	Description
infagrid.orchestrator.svc.sunset.time	Maximum lifetime for an Orchestrator service, in hours. Default is 24 hours.

You can also disable sunset by setting the property to 0 or a negative value. If you disable sunset, the Orchestrator never shuts down during a mapping run.

Troubleshooting Blaze Monitoring

When I run a mapping on the Blaze engine and try to view the grid task log, the Blaze Job Monitor does not fetch the full log.

The grid task log might be too large. The Blaze Job Monitor can only fetch up to 2 MB of an aggregated log. The first line of the log reports this information and provides the location of the full log on HDFS. Follow the link to HDFS and search for "aggregated logs for grid mapping." The link to the full log contains the grid task number.

The Blaze Job Monitor will not start.

Check the Hadoop environment logs to locate the issue. If you do not find an issue, stop the Grid Manager with the `infacmd stopBlazeService` command and run the mapping again.

The Monitoring URL does not appear in the Properties view of the Administrator tool.

Locate the URL in the YARN log.

When Blaze processes stop unexpectedly, Blaze does not save logs in the expected location.

When Blaze stops unexpectedly, you can access Blaze service logs through the YARN monitor. Use one of these methods:

- The Grid Manager log contains all Blaze job container IDs and identifies the host on which Blaze ran. Alter the Grid Manager log URL with the container ID and host name of the Blaze host.
- Run the command `yarn logs -applicationID <Blaze Grid Manager Application ID>`.

A Blaze Job Monitor that has been running for several days loses its connection to the Application Timeline Server on the Hortonworks cluster.

The Blaze engine requires a running Application Timeline Server on the cluster. When the Blaze engine starts a mapping run, the Blaze Job Monitor checks the state of the Application Timeline Server. The Grid Manager will start it if it is not running. When the connection to the Application Timeline Server is lost, the Blaze engine attempts to reconnect to it. If the Application Timeline Server stops during a Blaze mapping run, you can restart it by restarting the Grid Manager.

Note: When the Application Timeline Server is configured to run on the cluster by default, the cluster administrator must manually restart it on the cluster.

When a mapping takes more than 24 hours to execute, the mapping fails.

When mappings run on the Blaze engine for more than 24 hours, some mappings might fail because the Orchestrator service has a default sunset time of 24 hours. After 24 hours, the Orchestrator shuts down, which causes the Blaze Grid Manager to shut down.

To increase the sunset time to be more than 24 hours, configure the following property in the Hadoop connection advanced properties:

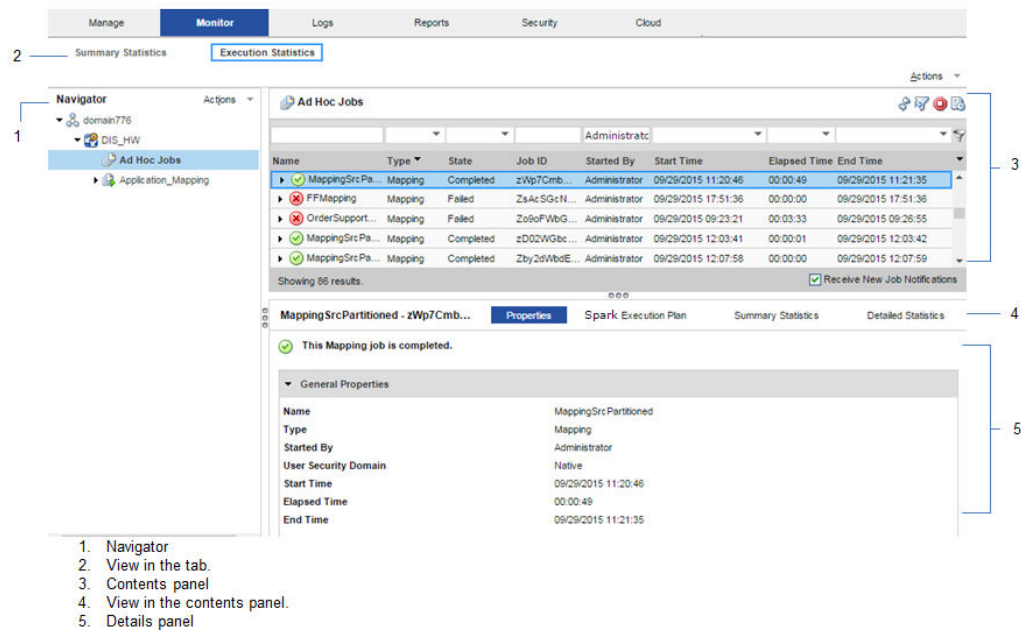
```
infagrid.orchestrator.svc.sunset.time=[HOURS]
```

You can also disable sunset by setting the property to 0 or a negative value. If you disable sunset, the Blaze Grid Manager never shuts down.

Spark Engine Monitoring

You can monitor statistics and view log events for a Spark engine mapping job in the Monitor tab of the Administrator tool. You can also monitor mapping jobs for the Spark engine in the YARN web user interface.

The following image shows the Monitor tab in the Administrator tool:



The Monitor tab has the following views:

Summary Statistics

Use the **Summary Statistics** view to view graphical summaries of object states and distribution across the Data Integration Services. You can also view graphs of the memory and CPU that the Data Integration Services used to run the objects.

Execution Statistics

Use the **Execution Statistics** view to monitor properties, run-time statistics, and run-time reports. In the Navigator, you can expand a Data Integration Service to monitor **Ad Hoc Jobs** or expand an application to monitor deployed mapping jobs or workflows

When you select **Ad Hoc Jobs**, deployed mapping jobs, or workflows from an application in the Navigator of the **Execution Statistics** view, a list of jobs appears in the contents panel. The contents panel displays jobs that are in the queued, running, completed, failed, aborted, and cancelled state. The Data Integration Service submits jobs in the queued state to the cluster when enough resources are available.

The contents panel groups related jobs based on the job type. You can expand a job type to view the related jobs under it.

Access the following views in the **Execution Statistics** view:

Properties

The **Properties** view shows the general properties about the selected job such as name, job type, user who started the job, and start time of the job.

Spark Execution Plan

When you view the Spark execution plan for a mapping, the Data Integration Service translates the mapping to a Scala program and an optional set of commands. The execution plan shows the commands and the Scala program code.

Summary Statistics

The **Summary Statistics** view appears in the details panel when you select a mapping job in the contents panel. The **Summary Statistics** view displays the following throughput statistics for the job:

- **Pre Job Task.** The name of each job task that reads source data and stages the row data to a temporary table before the Spark job runs. You can also view the bytes and average bytes processed for each second.
- **Source.** The name of the source.
- **Target.** The name of the target.
- **Rows.** For source, the number of rows read by the Spark application. For target, the sum of rows written to the target and rejected rows.
- **Post Job Task.** The name of each job task that writes target data from the staged tables. You can also view the bytes and average bytes processed for each second.

When a mapping contains a Union transformation with multiple upstream sources, the sources appear in a comma separated list in a single row under Sources.

Note: In a Hive mapping with an Update Strategy transformation containing the DD_UPDATE condition, the target contains only the temporary tables after the Spark job runs. The result of the mapping job statistics appears in the post job task and indicates twice the number of records updated.

The following image shows the **Summary Statistics** view in the details panel for a mapping run on the Spark engine:

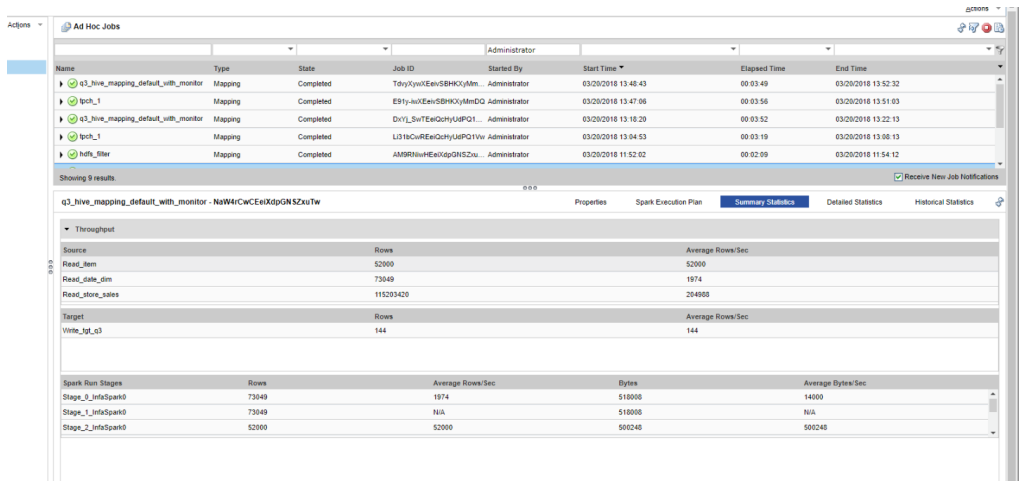
Name	Type	State	Job ID	Started By	Start Time	Elapsed Time	End Time
SparkOracle	Mapping	Completed	nD4SeqRREifBiamJun5w	admin	08/20/2018 13:48:26	00:03:52	08/20/2018 13:52:19
Post_UserDefinedJob_Task_1	UserDefinedTask	Completed	nD4SeqRREifBiamJun5w	admin	08/20/2018 13:52:04	00:00:14	08/20/2018 13:52:19
InfaSpark0	Spark Application	Completed	nD4SeqRREifBiamJun5w	admin	08/20/2018 13:51:05	00:00:58	08/20/2018 13:52:04
Pre_UserDefinedJob_Task_1	UserDefinedTask	Completed	nD4SeqRREifBiamJun5w	admin	08/20/2018 13:50:25	00:00:38	08/20/2018 13:51:04
Pre_UserDefinedJob_Task_2	UserDefinedTask	Completed	nD4SeqRREifBiamJun5w	admin	08/20/2018 13:50:25	00:00:37	08/20/2018 13:51:03

Section	Item	Rows	Average Rows/Sec	Bytes	Average Bytes/Sec
Pre Job Task	Pre_UserDefinedJob_Task_1	16384	0	96591	2
	Pre_UserDefinedJob_Task_2	0	0	247	0
Source	Read_ORACLE_TEST	16384	4096		
	Read_ORACLE_TEST1	0		N/A	
Spark Run Stages	Stage_0_InfaSpark0	16384	4096	487288	121822
	Stage_1_InfaSpark0	0	0	N/A	N/A
Target	Write_ORACLE_TEST	0		N/A	
Post Job Task	Post_UserDefinedJob_Task_1	0	0	N/A	N/A

You can also view the Spark run stages information in the details pane of the Summary Statistics view on the Execution Statistics Monitor tab. It appears as a list after the sources and before the targets.

The **Spark Run Stages** displays the absolute counts and throughput of rows and bytes related to the Spark application stage statistics. Rows refer to the number of rows that the stage writes, and bytes refer to the bytes broadcasted in the stage.

The following image displays the Spark Run Stages:



For example, the Spark Run Stages column contains the Spark application staged information starting with stage_<ID>. In the example, Stage_0 shows the statistics related to the Spark run stage with ID=0 in the Spark application.

Consider when the Spark engine reads source data that includes a self-join with verbose data enabled. In this scenario, the optimized mapping from the Spark application does not contain any information on the second instance of the same source in the Spark engine logs.

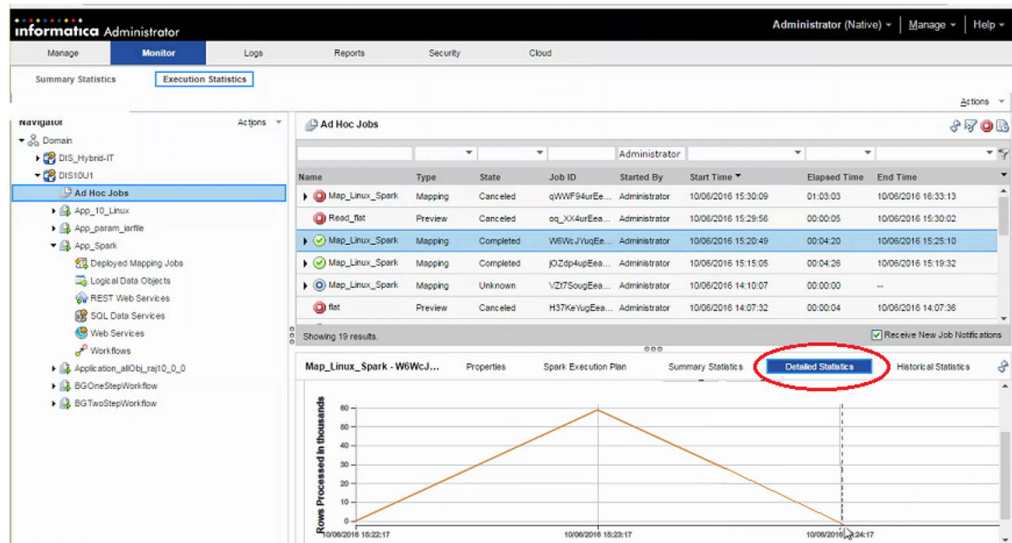
Consider when you read data from the temporary table and the Hive query of the customized data object leads to the shuffling of the data. In this scenario, the filtered source statistics appear instead of reading from the temporary source table in the Spark engine log.

When you run a mapping with Spark monitoring enabled, performance varies based on the mapping complexity. It can take up to three times longer than usual processing time with monitoring enabled. By default, monitoring is disabled.

Detailed Statistics

The **Detailed Statistics** view appears in the details panel when you select a mapping job in the contents panel. The **Detailed Statistics** view displays a graph of the row count for the job run.

The following image shows the **Detailed Statistics** view in the details panel for a mapping run on the Spark engine:



Viewing Hive Tasks

When you have a Hive source with a transactional table, you can view the Hive task associated with the Spark job.

When you run a mapping on Spark that launches Hive tasks, you can view the Hive query statistics in the session log and in the Administrator tool for monitoring along with the Spark application. For example, you can monitor information related to the Update Strategy transformation and SQL authorization associated to the mapping on Spark.

You can view the Summary Statistics for a Hive task in the Administrator tool. The Spark statistics continue to appear. When the Spark engine launches a Hive task, you can see Source Load Summary and Target Load summary including Spark data frame with Hive task statistics. Otherwise, when you have only a Spark task, the Source Load Summary and Target Load Summary do not appear in the session log.

Under Target Load Summary, all Hive instances will be prefixed with 'Hive_Target_'. You can see same instance name in the Administrator tool.

You can view the Tez job statistics in the Administrator tool when reading and writing to Hive tables that the Spark engine launches in any of the following scenarios:

- You have resources present in the Amazon buckets.
- You have transactional Hive tables.
- You have table columns secured with fine grained authorization.

Incorrect statistics appears for all the Hive sources and targets indicating zero rows for average rows for each second, bytes, average bytes for each second, and rejected rows. You can see that only processed rows contain correct values, and the remaining columns will contain either 0 or N/A.

Spark Engine Logs

The Spark engine logs appear in the LDTM log. The LDTM logs the results of the Spark engine execution plan run for the mapping. You can view the LDTM log from the Developer tool or the Monitoring tool for a mapping job.

The log for the Spark engine shows the step to translate the mapping to an internal format, steps to optimize the mapping, steps to render the mapping to Spark code, and the steps to submit the code to the Spark

executor. The logs also show the Scala code that the Logical Data Translation Generator creates from the mapping logic.

When you run Sqoop mappings on the Spark engine, the Data Integration Service prints the Sqoop log events in the mapping log.

Viewing Spark Logs

You can view logs for a Spark mapping from the YARN web user interface.

1. In the YARN web user interface, click an application ID to view.
2. Click the application **Details**.
3. Click the **Logs** URL in the application details to view the logs for the application instance.

The log events appear in another browser window.

Troubleshooting Spark Engine Monitoring

Do I need to configure a port for Spark Engine Monitoring?

Spark engine monitoring requires the cluster nodes to communicate with the Data Integration Service over a socket. The Data Integration Service picks the socket port randomly from the port range configured for the domain. The network administrators must ensure that the port range is accessible from the cluster nodes to the Data Integration Service. If the administrators cannot provide a port range access, you can configure the Data Integration Service to use a fixed port with the `SparkMonitoringPort` custom property. The network administrator must ensure that the configured port is accessible from the cluster nodes to the Data Integration Service.

Recovered jobs show 0 elapsed time in monitoring statistics

When a job is recovered, the Monitoring tab shows the same start and end time for the job, and the elapsed time = 0. While this statistic is not the actual elapsed time, it enables you to identify jobs that were recovered. For a more accurate view of the elapsed time for the job, view the Spark job logs on the cluster or the session logs on the Data Integration Service.

I enabled big data recovery, but recovered jobs show missing or incorrect statistics in the monitoring tab

CHAPTER 9

Hierarchical Data Processing

This chapter includes the following topics:

- [Overview of Hierarchical Data Processing, 155](#)
- [How to Develop a Mapping to Process Hierarchical Data, 156](#)
- [Complex Data Types, 158](#)
- [Complex Ports, 163](#)
- [Complex Data Type Definitions, 165](#)
- [Type Configuration, 170](#)
- [Complex Operators, 174](#)
- [Complex Functions, 176](#)
- [Hierarchical Data Preview, 177](#)

Overview of Hierarchical Data Processing

The Spark and Databricks Spark engines can process mappings that use complex data types, such as array, struct, and map. With complex data types, the run-time engine directly reads, processes, and writes hierarchical data in complex files.

The Spark and Databricks Spark engines can process hierarchical data in Avro, JSON, and Parquet complex files. They use complex data types to represent the native data types for hierarchical data in complex files. For example, a hierarchical data of type record in an Avro file is represented as a struct data type by the run-time engine.

Hierarchical Data Processing Scenarios

You can develop mappings for the following hierarchical data processing scenarios:

- To generate and modify hierarchical data.
- To transform relational data to hierarchical data.
- To transform hierarchical data to relational data.
- To convert data from one complex file format to another. For example, read hierarchical data from an Avro source and write to a JSON target.

Hierarchical Data Processing Configuration

You create a connection to access complex files and create a data object to represent data in the complex file. Then, configure the data object read and write operations to project columns as complex data types. To

read and write hierarchical data in complex files, you create a mapping, add a Read transformation based on the read operation, and add a Write transformation based on the write operation. To process hierarchical data, configure the following objects and transformation properties in a mapping:

- **Complex ports.** To pass hierarchical data in a mapping, create complex ports. You create complex ports by assigning complex data types to ports.
- **Complex data type definitions.** To process hierarchical data of type struct, create or import complex data type definitions that represent the schema of struct data.
- **Type configuration.** To define the properties of a complex port, specify or change the type configuration.
- **Complex operators and functions.** To generate or modify hierarchical data, create expressions using complex operators and functions.

You can also use hierarchical conversion wizards to simplify some of the mapping development tasks.

How to Develop a Mapping to Process Hierarchical Data

Develop a mapping with complex ports, operators, and functions to process hierarchical data on the Spark or Databricks Spark engine.

Note: The tasks and the order in which you perform the tasks to develop the mapping depend on the mapping scenario.

The following list outlines the high-level tasks to develop and run a mapping to read, write, and process hierarchical data in complex files.

Create a connection.

Create a connection to access data in complex files based on the file storage.

Create a data object.

1. Create a data object to represent the complex files as sources or targets.
2. Configure the data object properties.
3. In the read and write operations, enable the column file properties to project columns in the complex files as complex data types.

Create a mapping and add mapping objects.

1. Create a mapping, and add Read and Write transformations.
To read from and write to a complex file, add Read and Write transformations based on the data object read and write operations.
To write to an Avro or Parquet file, you can also create a complex file target from an existing transformation in the mapping.
2. Based on the mapping logic, add other transformations that are supported on the run-time engine.

Generate struct data.

Based on the mapping scenario, use one of the hierarchical conversion wizards to generate struct data. You can also perform the following steps manually:

Create or import complex data type definitions for struct ports.

1. Create or import complex data type definitions that represent the schema of the struct data. The complex data type definitions are stored in the type definition library, which is a Model repository object. The default name of the type definition library is `m_Type_Definition_Library`.
2. If a mapping uses one or more mapplets, rename the type definition libraries in the mapping and the mapplets to ensure that the names are unique.

Create and configure struct ports in transformations.

1. Create ports in transformations and assign struct data type.
2. Specify the type configuration for the struct ports.
You must reference a complex data type definition for the struct port.
3. Create expressions with complex functions to generate struct data.

Modify struct data.

You can convert struct data to relational or hierarchical data. If the struct data contains elements of primitive data types, you can extract the elements as relational data. If the struct data contains elements of complex data types, you can extract the elements as hierarchical data. Based on the mapping scenario, use one of the hierarchical conversion wizards to modify struct data. You can also perform the following steps manually:

1. Create output ports with port properties that match the element of the struct data that you want to extract.
2. Create expressions with complex operators or complex functions to modify the struct data.

Generate array data.

1. Create ports in transformations and assign array data type.
2. Specify the type configuration for the array ports.
3. Create expressions with complex functions to generate array data.

Modify array data.

You can convert array data to relational or hierarchical data. If the array data contains elements of primitive data types, you can extract the elements as relational data. If the array data contains elements of complex data types, you can extract the elements as hierarchical data. Based on the mapping scenario, use one of the hierarchical conversion wizards to modify array data. You can also perform the following steps manually:

1. Create output ports with port properties that match the element of the array data that you want to extract.
2. Create expressions with complex operators or complex functions to modify the array data.

Generate map data.

1. Create ports in transformations and assign map data type.
2. Specify the type configuration for the map ports.
3. Create expressions with complex functions to generate map data.

Modify map data.

You can convert map data to relational or hierarchical data. If the map data contains elements of primitive data types, you can extract the elements as relational data. If the map data contains elements of complex data types, you can extract the elements as hierarchical data. Based on the mapping

scenario, use one of the hierarchical conversion wizards to modify map data. You can also perform the following steps manually:

1. Create output ports with port properties that match the element of the map data that you want to extract.
2. Create expressions with complex operators or complex functions to modify the map data.

Configure the transformations.

Link the ports and configure the transformation properties based on the mapping logic.

Configure the mapping properties.

Configure the following mapping run-time properties:

- To run on the Spark engine, choose the Spark validation environment and Hadoop as the execution environment.
- To run on the Databricks Spark engine, choose the Databricks Spark validation environment and Hadoop as the execution environment.

Validate and run the mapping.

1. Validate the mapping to fix any errors.
2. Optionally, view the engine execution plan to debug the logic.
3. Run the mapping.

Complex Data Types

A complex data type is a transformation data type that represents multiple data values in a single column position. The data values are called elements. Elements in a complex data type can be of primitive or complex data types. Use complex data types to process hierarchical data in mappings that run on the Spark or Databricks Spark engine.

Transformation data types include the following data types:

Primitive data type

A transformation data type that represents a single data value in a single column position. Data types such as decimal, integer, and string are primitive data types. You can assign primitive data types to ports in any transformation.

Complex data type

A transformation data type that represents multiple data values in a single column position. Data types such as array, map, and struct are complex data types. You can assign complex data types to ports in some transformations.

Nested data type

A complex data type that contains elements of complex data types. Complex data types such as an array of structs or a struct with an array of other structs are nested data types.

The following table lists the complex data types:

Complex Data Type	Description
array	An array is an ordered collection of elements. The elements can be of a primitive or complex data type. All elements in the array must be of the same data type.
map	A map is an unordered collection of key-value pair elements. The key must be of a primitive data type. The value can be of a primitive or complex data type.
struct	A struct is a collection of elements of different data types. The elements can be of primitive or complex data types. Struct has a schema that defines the structure of the data.

The following image shows primitive, complex, and nested data types assigned to ports in a transformation:

	Name	Type	Type Configuration
1	emp_name	string	N/A
2	emp_sal	decimal	N/A
3	emp_phone	array	string []
4	emp_id_dept	map	< integer, string >
5	emp_address	struct	(typedef_adrs)
6	emp_bonus	array	(typedef_perf) []

1. Primitive data types
2. Complex data types
3. Nested data type

The ports emp_name and emp_sal are of primitive data types. The ports emp_phone, emp_id_dept, and emp_address are of complex data types. The port emp_bonus is of a nested data type. The array contains elements of type struct.

Array Data Type

An array data type represents an ordered collection of elements. To pass, generate, or process array data, assign array data type to ports.

An array is a zero-based indexed list. An array index indicates the position of the array element. For example, the array index 0 indicates the first element in an array. The transformation language includes operators to access array elements and functions to generate and process array data.

An array can be one-dimensional or multidimensional. A one-dimensional array is a linear array. A multidimensional array is an array of arrays. Array transformation data types can have up to five dimensions.

Format

```
array <data_type> []
```

The following table describes the arguments for this data type:

Argument	Description
array	Name of the array column or port.
data_type	Data type of the elements in an array. The elements can be primitive data types or complex data types. All elements in the array must be of the same data type.
[]	Dimension of the array represented as subscript. A single subscript [] represents a one-dimensional array. Two subscripts [] [] represent a two-dimensional array. Elements in each dimension are of the same data type.

The elements of an array do not have names. The number of elements in an array can be different for each row.

Array Examples

One-dimensional array

The following array column represents a one-dimensional array of string elements that contains customer phone numbers:

```
custphone string[]
```

The following example shows data values for the custphone column:

custphone

```
[205-128-6478,722-515-2889]
```

```
[107-081-0961,718-051-8116]
```

```
[107-031-0961,NULL]
```

Two-dimensional array

The following array column represents a two-dimensional array of string elements that contains customer work and personal email addresses.

```
email_work_pers string[][]
```

The following example shows data values for the email_work_pers column:

email_work_pers

```
[john_baer@xyz.com,jbaer@xyz.com][john.baer@fgh.com,jbaer@ijk.com]
```

```
[bobbi_apperley@xyz.com,bapperl@xyz.com][apperlbob@fgh.com,bobbi@ijk.com]
```

```
[linda_bender@xyz.com,lbender@xyz.com][l.bender@fgh.com,NULL]
```

Map Data Type

A map data type represents an unordered collection of key-value pair elements. A map element is a key and value pair that maps one thing to another. To pass, generate, or process map data, assign map data type to ports.

The key must be of a primitive data type. The value can be of a primitive or complex data type. A map data type with values of a complex data type is a nested map. A nested map can contain up to three levels of nesting of map data type.

The transformation language includes subscript operator to access map elements. It also includes functions to generate and process map data.

Format

```
map <primitive_type -> data_type>
```

The following table describes the arguments for this data type:

Argument	Description
map	Name of the map column or port.
primitive_type	Data type of the key in a map element. The key must be of a primitive data type.
data_type	Data type of the value in a map element. The value can be of a primitive or complex data type.

Map Example

The following map column represents map data with an integer key and a string value to map customer ids with customer names:

```
custid_name <integer -> string>
```

The following example shows data values for the custid_name column:

```
custid_name
```

```
<26745 -> 'John Baer'>
```

```
<56743 -> 'Bobbi Apperley'>
```

```
<32879 -> 'Linda Bender'>
```

Struct Data Type

A struct data type represents a collection of elements of different data types. A struct data type has an associated schema that defines the structure of the data. To pass, generate, or process struct data, assign struct data type to ports.

The schema for the struct data type determines the element names and the number of elements in the struct data. The schema also determines the order of elements in the struct data. Informatica uses complex data type definitions to represent the schema of struct data.

The transformation language includes operators to access struct elements. It also includes functions to generate and process struct data and to modify the schema of the data.

Format

```
struct {element_name1:value1 [, element_name2:value2, ...]}
```

Schema for the struct is of the following format:

```
schema {element_name1:data_type1 [, element_name2:data_type2, ...]}
```

The following table describes the arguments for this data type:

Argument	Description
struct	Name of the struct column or port.
schema	A definition of the structure of data. Schema is a name-type pair that determines the name and data type of the struct elements.
element_name	Name of the struct element.
value	Value of the struct element.
data_type	Data type of the element value. The element values can be of a primitive or complex data type. Each element in the struct can have a different data type.

Struct Example

The following schema is for struct data to store customer addresses:

```
address  
{st_number:integer,st_name:string,city:string,state:string,zip:string}
```

The following example shows struct data values for the `cust_address` column:

cust_address

```
{st_number:154,st_name:Addison Ave,city:Redwood City,state:CA,zip:94065}  
  
{st_number:204,st_name:Ellis St,city:Mountain View,state:CA,zip:94043}  
  
{st_number:357,st_name:First St,city:Sunnyvale,state:CA,zip:94085}
```

Rules and Guidelines for Complex Data Types

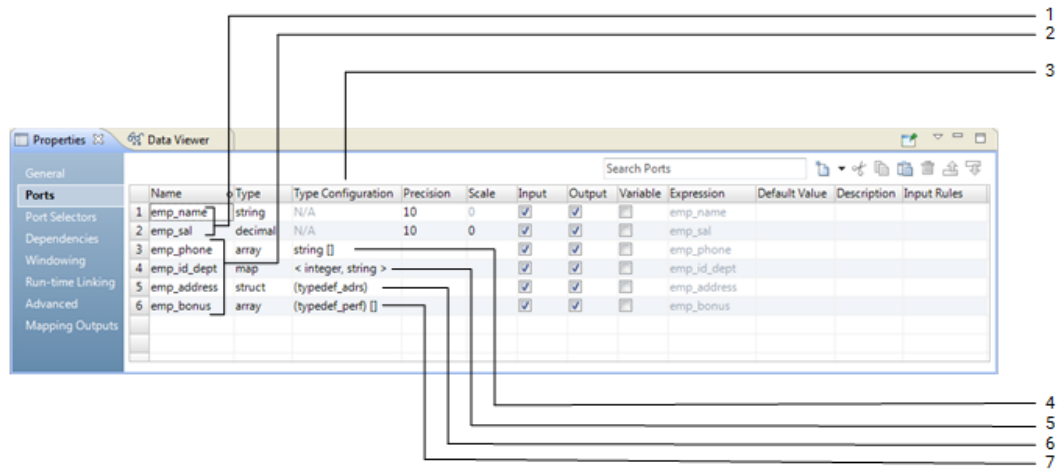
Consider the following rules and guidelines when you work with complex data types:

- A nested data type can contain up to 10 levels of nesting.
- A nested map can contain up to three levels of nesting of map data types.
- An array data type cannot directly contain an element of type array. Use multidimensional arrays to create a nested array. For example, an array with two dimensions is an array of arrays.
- A multidimensional array can contain up to five levels of nesting. The array dimension determines the levels of nesting.
- Each array in a multidimensional array must have elements of the same data type.

Complex Ports

A complex port is a port that is assigned a complex data type. Based on the complex data type, you must specify the complex port properties. Use complex ports in transformations to pass or process hierarchical data in a mapping.

The following image shows the complex ports and complex port properties on the Ports tab for a transformation:



1. Port
2. Complex port
3. Type configuration
4. Type configuration for an array port
5. Type configuration for a map port
6. Type configuration for a struct port
7. Type configuration for a port of nested data type

Based on the data type, a transformation can include the following ports and port properties:

Port

A port of a primitive data type that you can create in any transformation.

Complex port

A port of a complex or nested data type that you can create in some transformations. Array, map, and struct ports are complex ports. Based on the complex data type, you specify the complex port properties in the type configuration column.

Type configuration

A set of properties that you specify for the complex port. The type configuration determines the data type of the complex data type elements or the schema of the data. You specify the data type of the elements for array and map ports. You specify a complex data type definition for the struct port.

Type configuration for an array port

Properties that determine the data type of the array elements. In the image, the array port emp_phone is a one-dimensional array with an ordered collection of string elements. An array with string elements is also called an array of strings.

Type configuration for a map port

Properties that determine the data type of the key-value pair of the map elements. In the image, the map port `emp_id_dept` is an unordered collection of key-value pairs of type integer and string.

Type configuration for a struct port

Properties that determine the schema of the data. To represent the schema, you create or import a complex data type definition. In the image, the struct port `emp_address` references a complex data type definition `typedef_adrs`.

Type configuration for a port of nested data type

Properties that determine the nested data type. In the image, the array port `emp_bonus` is a one-dimensional array with an ordered collection of struct elements. The struct elements reference a complex data type definition `typedef_bonus`. An array with struct elements is also called an array of structs.

Complex Ports in Transformations

You can create complex ports in some transformations that are supported on the Spark and Databricks Spark engines. Read and Write transformations can represent ports that pass hierarchical data as complex data types.

You can create complex ports in the following transformations:

- Aggregator
- Expression
- Filter
- Java
- Joiner
- Lookup
- Normalizer
- Router
- Sorter
- Union

Note: The Databricks Spark engine does not support the Java transformation.

The Read and Write transformations can read and write hierarchical data in complex files. To read and write hierarchical data, the Read and Write transformations must meet the following requirements:

- The transformation must be based on a complex file data object.
- The data object read and write operations must project columns as complex data types.

Rules and Guidelines for Complex Ports

Consider the following rules and guidelines when you work with complex ports:

- Aggregator transformation. You cannot define a group by value as a complex port.
- Filter transformation. You cannot use the operators `>`, `<`, `>=`, and `<=` in a filter condition to compare data in complex ports.
- Joiner transformation. You cannot use the operators `>`, `<`, `>=`, and `<=` in a join condition to compare data in complex ports.

- Lookup transformation. You cannot use a complex port in a lookup condition.
- Rank transformation. You cannot define a group by or rank value as a complex port.
- Router transformation. You cannot use the operators `>`, `<`, `>=`, and `<=` in a group filter condition to compare data in complex ports.
- Sorter transformation. You cannot define a sort key value as a complex port.
- You can use complex operators to specify an element of a complex port that is of a primitive data type. For example, an array port "emp_names" contains string elements. You can define a group by value as `emp_names[0]`, which is of type string.

Creating a Complex Port

Create complex ports in transformations to pass or process hierarchical data in mappings that run on the Spark or Databricks Spark engine.

1. Select the transformation in the mapping.
2. Create a port.
3. In the **Type** column for the port, select a complex data type.
The complex data type for the port appears in the Type column.

After you create a complex port, specify the type configuration for the complex port.

Complex Data Type Definitions

A complex data type definition represents the schema of the data. Use complex data type definitions to define the schema of the data for a struct port. If the complex port is of type struct or contains elements of type struct, you must specify the type configuration to reference a complex data type definition.

You can create or import complex data type definitions. You import complex data type definitions from complex files. Complex data type definitions are stored in the type definition library, which is a Model repository object.

When you create a mapping or a mapplet, the Developer tool creates a type definition library with the name `Type_Definition_Library`. You can view and rename the type definition library and the complex data type definitions in the Outline view of a mapping. The complex data type definitions appear on the type definition library tab of the mapping editor. The tab name is based on the name of the type definition library. You create or import complex data type definitions on the type definition library tab. When a mapping uses one or more mapplets, rename the type definition libraries in the mapping and the mapplets to ensure that the names are unique.

Type definition library and complex data type definitions are specific to a mapping or mapplet. A mapping or a mapplet can have one or more complex data type definitions. You can reference the same complex data type definition in one or more complex ports in the mapping or mapplet.

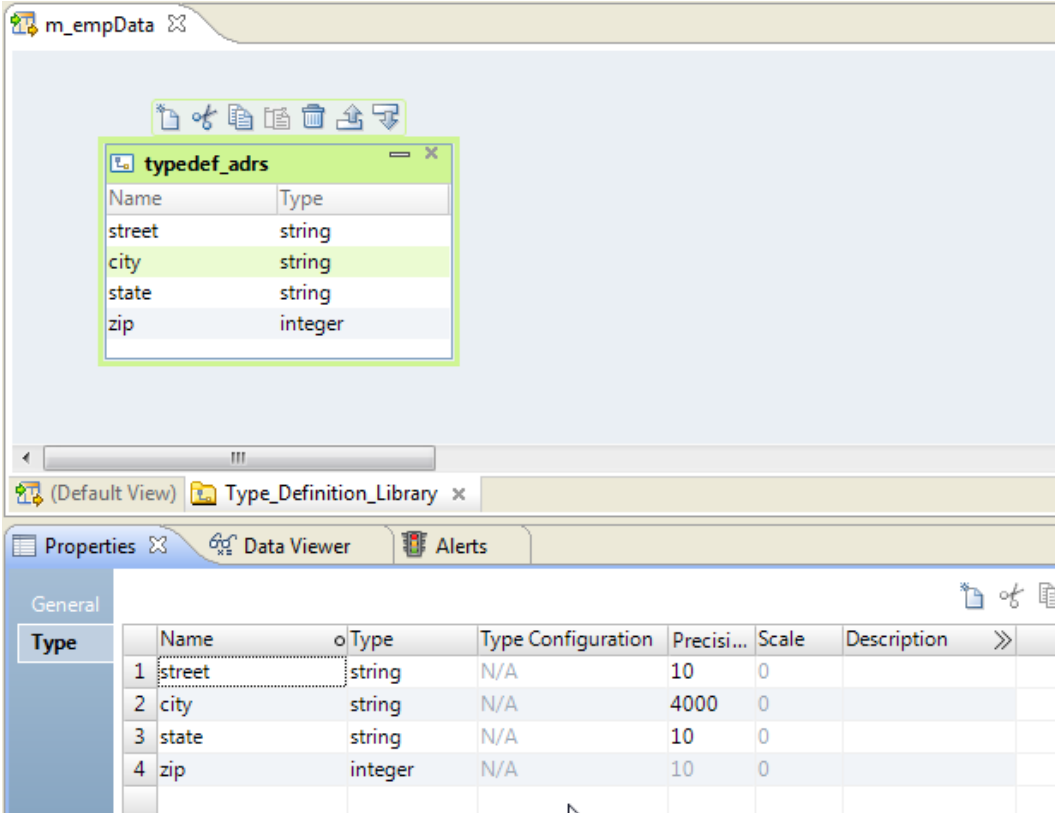
Complex Data Type Definition Example

The complex ports `work_address` and `home_address` of type struct have the following schema:

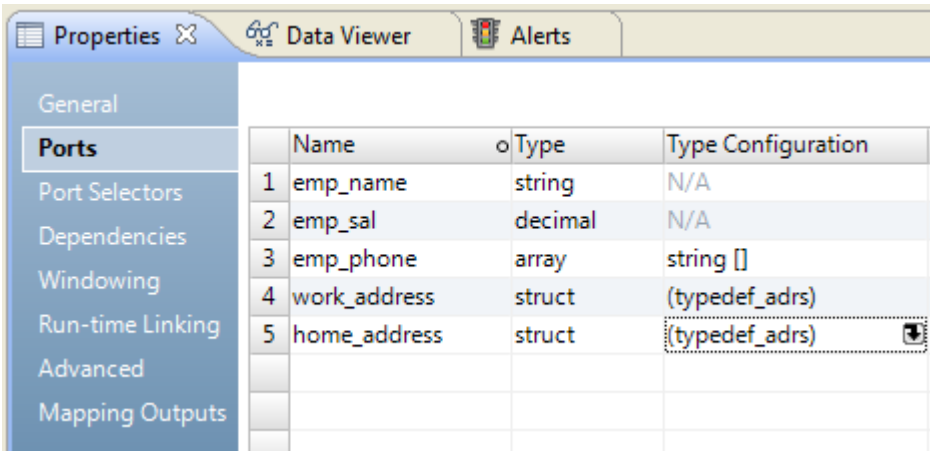
```
{street:string, city:string, state:string, zip:integer}
```

In the mapping, you create a complex data type definition `typedef_adrs` that defines the schema of the data. Then, you specify the type configuration of the struct ports `work_address` and `home_address` to use the `typedef_adrs` complex data type definition.

The following image shows a complex data type definition `typedef_adrs`:



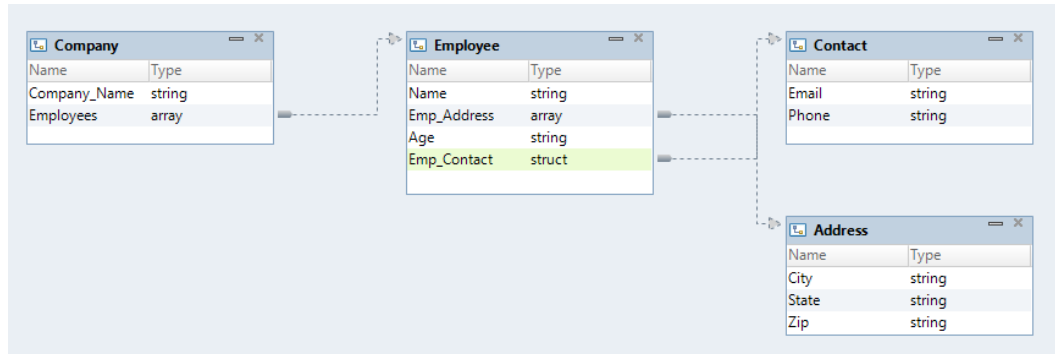
The following image shows two struct ports `work_address` and `home_address` that reference the complex data type definition `typedef_adrs`:



Nested Data Type Definitions

Elements of a complex data type definition can reference one or more complex data type definitions in the type definition library. Such complex data type definitions are called nested data type definitions.

The following image shows a nested data type definition Company on the type definition library tab:



The nested data type definition Company references the following complex data type definitions:

- In the complex data type definition Company, the array element Employees references the complex data type definition Employee.
- In the complex data type definition Employee, the Emp_Address element references the complex data type definition Address.
- In the complex data type definition Employee, the Emp_Contact element references the complex data type definition Contact.

Note: In a recursive data type definition, one of the complex data type definitions at any level is the same as any of its parents. You cannot reference a recursive data type definition to a struct port or a struct element in a complex port.

Rules and Guidelines for Complex Data Type Definitions

Consider the following rules and guidelines when you work with complex data type definitions:

- A struct port or a struct element in a complex port must reference a complex data type definition.
- You cannot reference a complex data type definition in one mapping to a complex port in another mapping.
- You cannot reference a recursive data type definition to a struct port or a struct element in a complex port.
- Changes to the elements in a complex data type definition are automatically propagated to complex ports that reference the complex data type definition.
- If you change the name of the type definition library or the complex data type definition, you must update the name in expressions that use complex functions such as STRUCT_AS, RESPEC, and CAST.

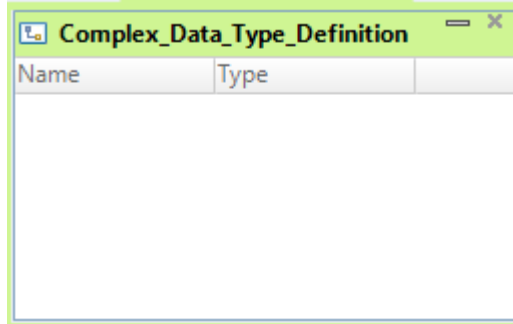
Creating a Complex Data Type Definition

A complex data type definition represents the schema of the data of type struct. You can create a complex data type definition for a mapping or a mapplet on the Type Definition Library tab of the mapping editor.

1. In the Developer tool, open the mapping or mapplet where you want to import a complex data type definition.

2. In the **Outline** view, select the **Type Definition Library** to view the **Type_Definition_Library** tab in the mapping editor.
3. Right-click an empty area of the editor and click **New Complex Data Type Definition**.

An empty complex data type definition appears on the Type Definition Library tab of the mapping editor.



4. Optionally, change the name of the complex data type definition.
5. Click the **New** button to add elements to the data type definition with a name and a type.
The data type of the element can be of primitive or complex data type.

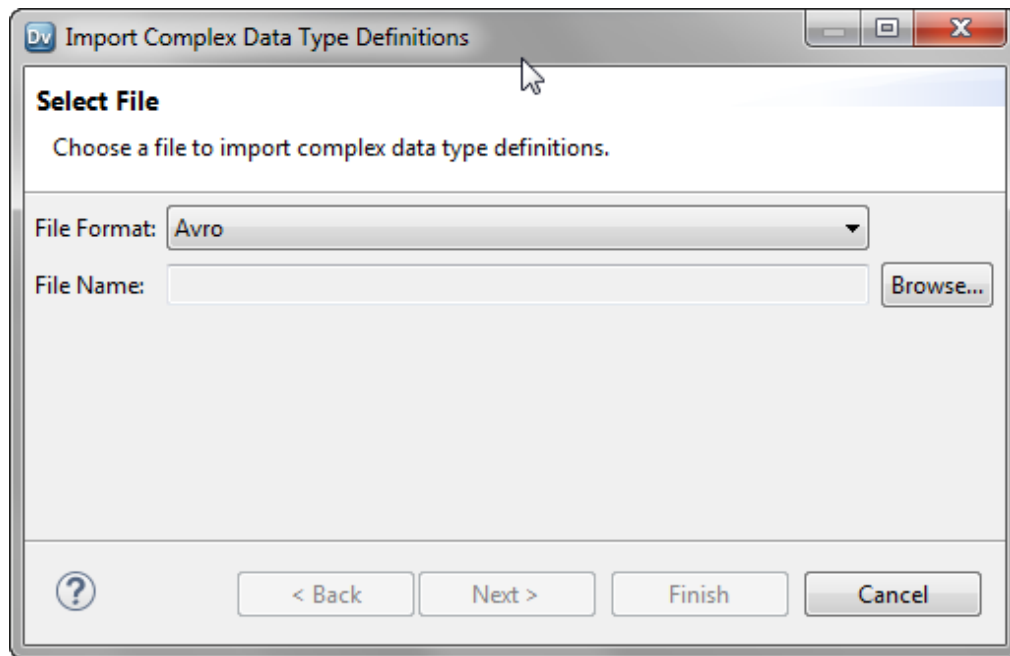
Importing a Complex Data Type Definition

A complex data type definition represents the schema of the data of type struct. You can import the schema of the hierarchical data in the complex file to the type definition library as a complex data type definition.

Import complex data type definitions from an Avro, JSON, or Parquet schema file. For example, you can import an Avro schema from an .avsc file as a complex data type definition.

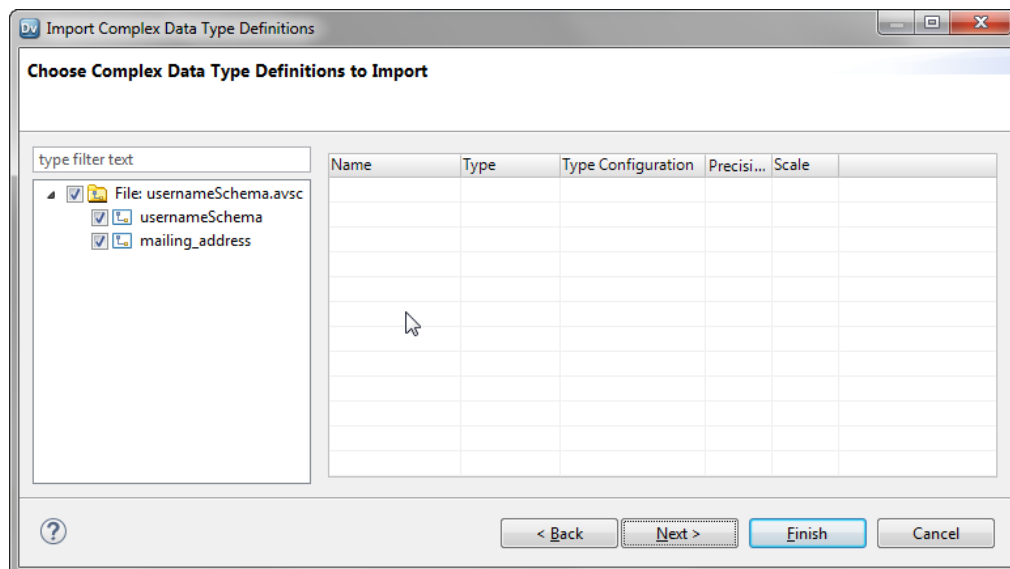
1. In the Developer tool, open the mapping or mapplet where you want to import a complex data type definition.
2. In the **Outline** view, select the **Type Definition Library** to view the **Type_Definition_Library** tab in the mapping editor.
3. Right-click an empty area of the editor and click **Import Complex Data Type Definitions**.

The **Import Complex Data Type Definitions** dialog box appears.



4. Select a complex file format from the **File Format** list.
5. Click **Browse** to select the complex file schema from which you want to import the complex data type definition.
6. Navigate to the complex file schema and click **Open**.
7. Click **Next**.

The **Choose Complex Data Type Definitions to Import** page appears.



8. Select one or more schemas from the list to import.
9. Click **Next**.

The **Complex Data Type Definitions to Import** page appears.

10. Review the complex data type definitions that you want to import and click **Finish**.

The complex data type definition appears in the **Type Definition Library** tab of the mapping or the mapplet.

Type Configuration

A type configuration is a set of complex port properties that specify the data type of the complex data type elements or the schema of the data. After you create a complex port, you specify or change the type configuration for the complex port on the Ports tab.

The type configuration for a complex port depends on the complex data type assigned to the port. You must specify a complex data type definition for a struct port or a struct element in a complex port. Array and map ports have a default type configuration that you can change.

The following table lists the type configuration properties for complex ports:

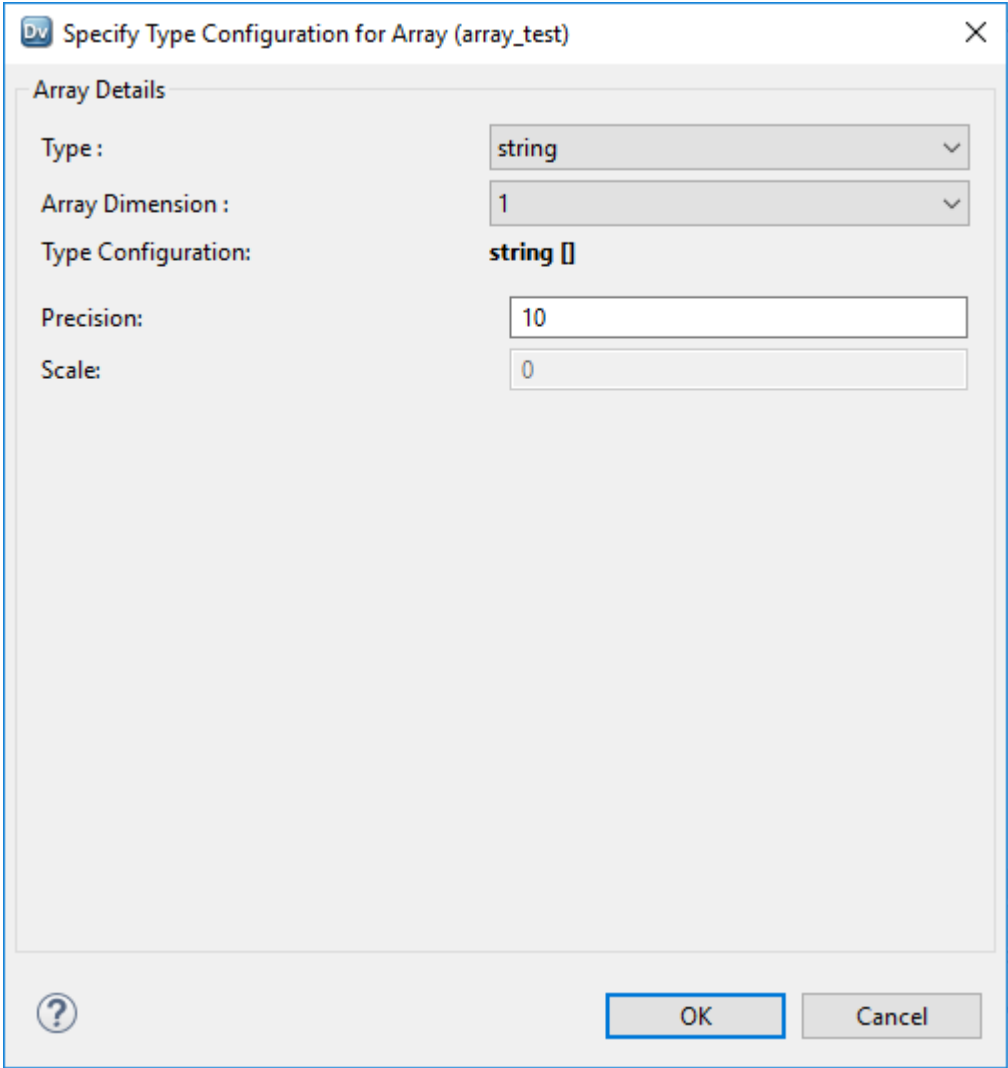
Complex Port	Type Configuration Properties
array	- Type. Data type of the elements of an array. Default is string. - Array Dimension. Number of levels for an array. Default is 1. - Precision. Maximum number of significant digits for numeric data types, or maximum number of characters for string data types. - Scale. Maximum number of digits after the decimal point of numeric values.
map	- Type. Data type of the key-value pair. Default is string for key and value. - Precision. Maximum number of significant digits for numeric data types, or maximum number of characters for string data types. - Scale. Maximum number of digits after the decimal point of numeric values.
struct	- Complex data type definition. The schema of the data that you created or imported as complex data type definition in the type definition library.

Changing the Type Configuration for an Array Port

The default type configuration for an array is `string[]`, which is a one-dimensional array of strings. You can change the type configuration for the array port. Specify the data type of the array elements and the array dimension.

1. In the transformation, select the array port for which you want to specify the type configuration.
2. In the **Type Configuration** column, click the **Open** button.

The **Specify Type Configuration for Array** dialog box appears.

The image shows a dialog box titled "Specify Type Configuration for Array (array_test)". It has a close button (X) in the top right corner. The dialog is divided into two main sections. The top section, labeled "Array Details", contains five configuration options: "Type:" with a dropdown menu showing "string"; "Array Dimension:" with a dropdown menu showing "1"; "Type Configuration:" displaying "string []"; "Precision:" with a text input field containing "10"; and "Scale:" with a text input field containing "0". The bottom section of the dialog is a light gray area containing a help icon (question mark in a circle) on the left and two buttons, "OK" and "Cancel", on the right.

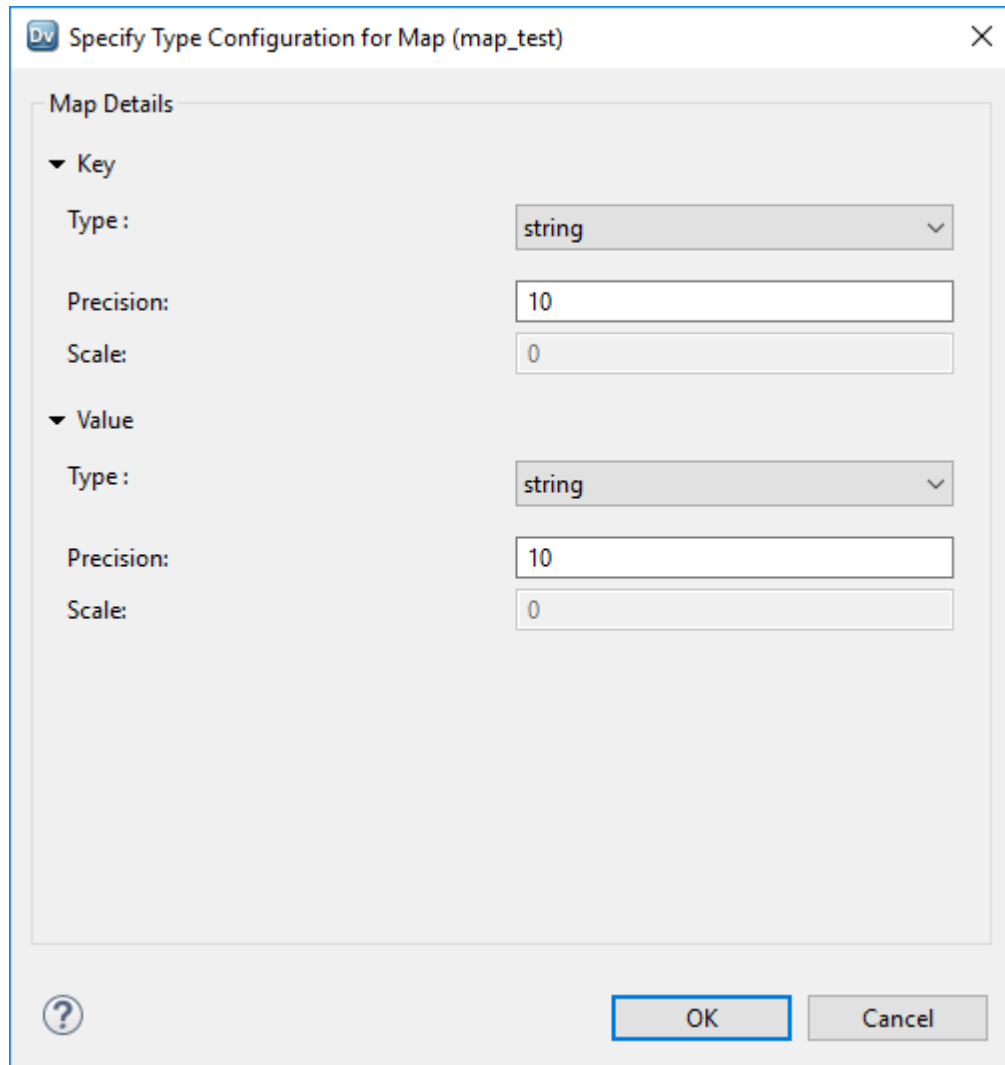
3. In the **Type** drop-down menu, change the data type of the elements in the array.
For example, if the elements in the array are of type integer, select **integer**.
The type configuration value appears as **integer []**.
4. Optionally, in the **Array Dimension** drop-down menu, select a number to create a nested array.
For example, array dimension 3 creates a nested array of 3 levels. The type configuration value appears as **integer [] [] []**.
5. Optionally, configure the following properties:
 - **Precision.** Maximum number of significant digits for numeric data types, or maximum number of characters for string data types.
 - **Scale.** Maximum number of digits after the decimal point of numeric values.
6. Click **OK**.
On the **Ports** tab, the specified type configuration appears in the **Type Configuration** column of the array port.

Changing the Type Configuration for a Map Port

The default type configuration for a map is `<string, string>`. You can change the type configuration for a map port. Specify the data type of the key-value pair.

1. In the transformation, select the map port for which you want to set the type configuration.
2. In the **Type Configuration** column, click the **Open** button.

The **Type Configuration** dialog box appears:



The dialog box titled "Specify Type Configuration for Map (map_test)" contains a "Map Details" section. Under "Key", there are fields for "Type:" (a dropdown menu set to "string"), "Precision:" (a text box with "10"), and "Scale:" (a text box with "0"). Under "Value", there are similar fields for "Type:" (dropdown set to "string"), "Precision:" (text box with "10"), and "Scale:" (text box with "0"). At the bottom right are "OK" and "Cancel" buttons, and at the bottom left is a help icon (a question mark in a circle).

3. Specify the properties for the key and value of the map data type.
 - a. In the **Type** drop-down menu, select the data type.
 - b. Optionally, in the **Precision** box, specify the maximum number of significant digits for numeric data types, or maximum number of characters for string data types.
 - c. Optionally, in the **Scale** box, specify the maximum number of digits after the decimal point of numeric values.
4. Click **OK**.

On the **Ports** tab, the specified type configuration appears in the **Type Configuration** column of the map port.

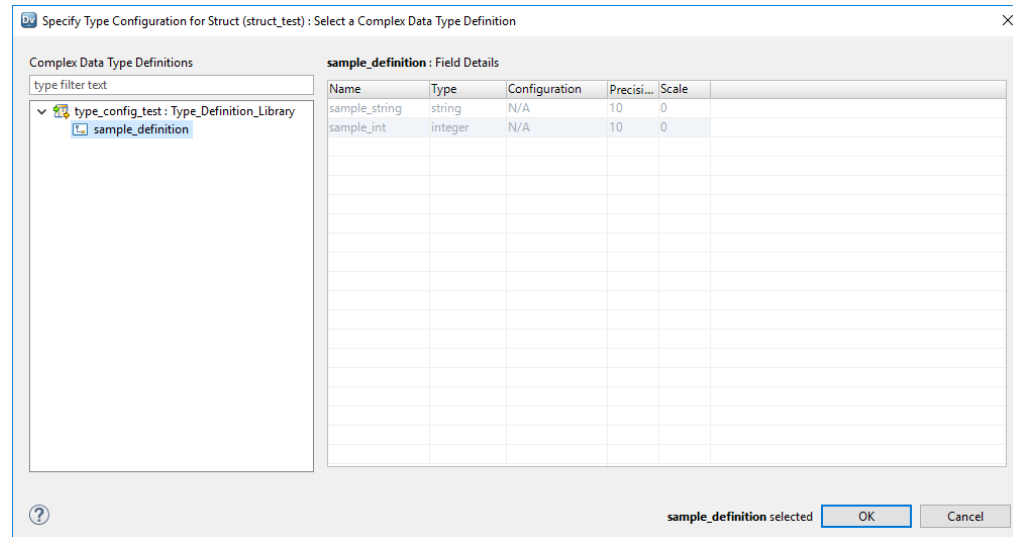
Specifying the Type Configuration for a Struct Port

The type configuration for a struct port requires a complex data type definition in the type definition library. Specify the complex data type definition that represents the schema of the struct data.

Before you specify the type configuration for the struct port, create or import a complex data type definition. You can also reference an existing complex data type definition in the type definition library.

1. In the transformation, select the struct port for which you want to specify the type configuration.
2. In the **Type Configuration** column, click the **Open** button.

The **Type Configuration** dialog box appears:



3. Select a complex data type definition from the list of definitions in the type definition library for the mapping or the mapplet.
4. Click **OK**.

On the **Ports** tab, the specified type configuration appears in the **Type Configuration** column of the struct port.

Complex Operators

A complex operator is a type of operator to access elements in a complex data type. The transformation language includes complex operators for complex data types. Use complex operators to access or extract elements in hierarchical data.

The following table lists the complex operators:

Complex Operators	Description
[]	Subscript operator. Use to access array and map elements. Syntax: <code>array[index]</code> - <code>array</code> is the array from which you want to access an element. - <code>index</code> is the position of an element within an array. For example, use <code>[0]</code> to access the first element in an one-dimensional array. Syntax: <code>map[key]</code> - <code>map</code> is the map from which you want to access the value corresponding to a key. - <code>key</code> is the map key for which you want to retrieve the map value.
.	Dot operator. Use to access struct elements. Syntax: <code>struct.element_name</code> - <code>struct</code> is the struct from which you want to access an element. - <code>element_name</code> is the name of the struct element.

Use complex operators in expressions to convert hierarchical data to relational data.

For more information about the operator syntax, return values, and examples, see the *Informatica Developer Transformation Language Reference*.

Complex Operator Examples

- Convert array data in a JSON file to relational data.

A JSON file has a `quarterly_sales` column that is of the array type with integer elements. You develop a mapping to transform the hierarchical data to relational data. The array port `quarterly_sales` contains four elements. Use subscript operators in expressions to extract the four elements into the four integer ports, `q1_sales`, `q2_sales`, `q3_sales`, and `q4_sales`.

The expressions for the integer output port are as follows:

```
quarterly_sales[0]
quarterly_sales[1]
quarterly_sales[2]
quarterly_sales[3]
```

- Convert struct data in a Parquet file to relational data.

A Parquet file has an `address` column that contains customer address as a struct type. You develop a mapping to transform hierarchical data to relational data. The struct port `address` contains three elements `city`, `state`, and `zip`, of type string. Use dot operators in expressions to extract the three struct elements into the three string ports, `city`, `state`, and `zip`.

The expressions for the string output port are as follows:

```
address.city
address.state
address.zip
```

Extracting an Array Element Using a Subscript Operator

Use the subscript operator in expressions to extract one or more elements of an array.

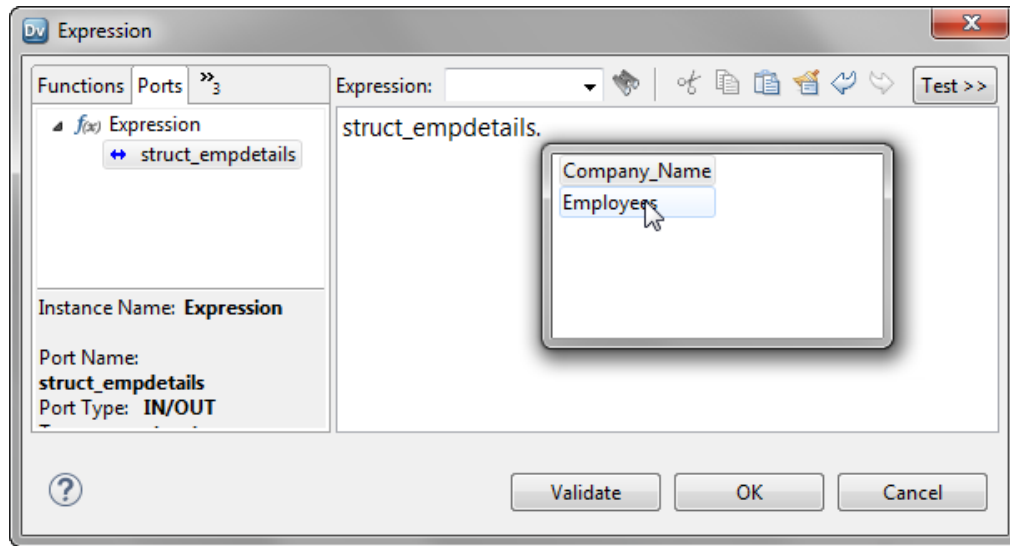
1. In the transformation, create an output port of the same data type as the return value of the array element that you want to extract.
To extract more than one element in an array, create an array output port. The data type of the elements in the type configuration for the array port must match the data type of the return value of the array elements.
2. In the **Expression** column for the output port, click the **Open** button.
The **Expression Editor** opens.
3. Delete the existing expression in the editor.
4. Click the **Ports** tab and select an array port.
5. Enter one or more index values within the subscript operator.
To extract array elements in a multidimensional array, enter index values within each subscript operator.
6. Click **Validate** to validate the expression.
7. Click **OK** to exit the **Validate Expression** dialog box.
8. Click **OK** to exit the **Expression Editor**.

Extracting a Struct Element Using the Dot Operator

Use the dot operator in expressions to extract elements of a struct.

1. In the transformation, create an output port of the same data type as the return value of the struct element that you want to extract.
2. In the **Expression** column for the output port, click the **Open** button.
The Expression Editor opens.
3. Delete the existing expression in the editor.
4. Click the **Ports** tab and select the struct port.
5. Enter a dot after the port name.

The Developer tool displays a list of elements in the struct.



6. Select the element that you want to extract.
7. Click **Validate** to validate the expression.
8. Click **OK** to exit the Validate Expression dialog box.
9. Click **OK** to exit the Expression Editor.

Complex Functions

A complex function is a type of pre-defined function in which the value of the input or the return type is of a complex data type. The transformation language includes complex functions for complex data types. Use complex functions to generate and process hierarchical data.

The following table describes the complex functions for array data type:

Complex Function	Description
ARRAY (<i>element1</i> [<i>element2</i>] ...)	Generates an array with the specified elements. The data type of the argument determines the data type of the array.
COLLECT_LIST (<i>value</i>)	Returns an array with elements in the specified port. The data type of the argument determines the data type of the array. COLLECT_LIST is an aggregate function.
CONCAT_ARRAY ('', <i>array_of_strings</i>)	Concatenates string elements in an array based on a separator that you specify and returns a string.
SIZE (<i>array</i>)	Returns the size of the array.

The following table describes the complex functions for map data type:

Complex Function	Description
MAP (<i>key1</i> , <i>value1</i> [, <i>key2</i> , <i>value2</i>] ...)	Generates a map with elements based on the specified key-value pair. The data type of arguments determines the data type of the map elements.
MAP_FROM_ARRAYS (<i>key_array</i> , <i>value_array</i>)	Generates a map from the specified key and value arrays. The data type of arguments determines the data type of the map elements.
MAP_KEYS (<i>map</i>)	Returns an array of keys for the specified map.
MAP_VALUES (<i>map</i>)	Returns an array of values for the specified map.
COLLECT_MAP (<i>value</i>)	Returns a map with elements based on the specified arguments. The data type of the argument determines the data type of the map. COLLECT_MAP is an aggregate function.
SIZE (<i>map</i>)	Returns the size of the map.

The following table describes the complex functions for struct data type:

Complex Function	Description
STRUCT_AS (<i>type_definition</i> , <i>struct</i>)	Generates a struct with a schema based on the specified complex data type definition and the values you pass as argument.
STRUCT (<i>element_name1</i> , <i>value1</i>)	Generates a struct with the specified names and values for the elements. The data type of the value is based on the specified value argument.
RESPEC (<i>type_definition</i> , <i>struct</i>)	Renames each element of the given struct value based on the names of the elements in the specified complex data type definition.
CAST (<i>type_definition</i> , <i>struct</i>)	Changes the data type and the name of each element for the given struct value based on the corresponding data type and name of the element in the specified complex data type definition.

For more information about the function syntax, return value, and examples, see the *Informatica Developer Transformation Language Reference*.

Hierarchical Data Preview

You can preview hierarchical data within a mapping in the Developer tool. You can choose sources and transformations as preview points in a mapping. Previewing hierarchical data helps to design and debug big data mappings.

When you preview data with hierarchical types, the Data Integration Service passes the request to the Spark engine. When you view the results, you can view the complex data types, the schema, and values. You can also navigate through the structure with breadcrumbs. You can export the results to a JSON file.

Note: Data preview on hierarchical data is available for technical preview. Technical preview functionality is supported but is unwarranted and is not production-ready. Informatica recommends that you use in non-production environments only.

Note: Data preview is supported for the Spark engine only. You cannot preview hierarchical data in a mapping configured for the Databricks Spark engine.

Hierarchical Data Preview Process

When you preview data, the Data Integration Service determines whether to run the job in the native environment or on the Spark engine.

If the preview point or any upstream transformation in the mapping contains a hierarchical data type, the Data Integration Service pushes the job to the Spark engine. Otherwise, the Data Integration Service runs the data preview. The preview point is the object in a mapping that you choose to view data for.

The Data Integration Service uses the following process when it pushes a data preview job to the Spark engine:

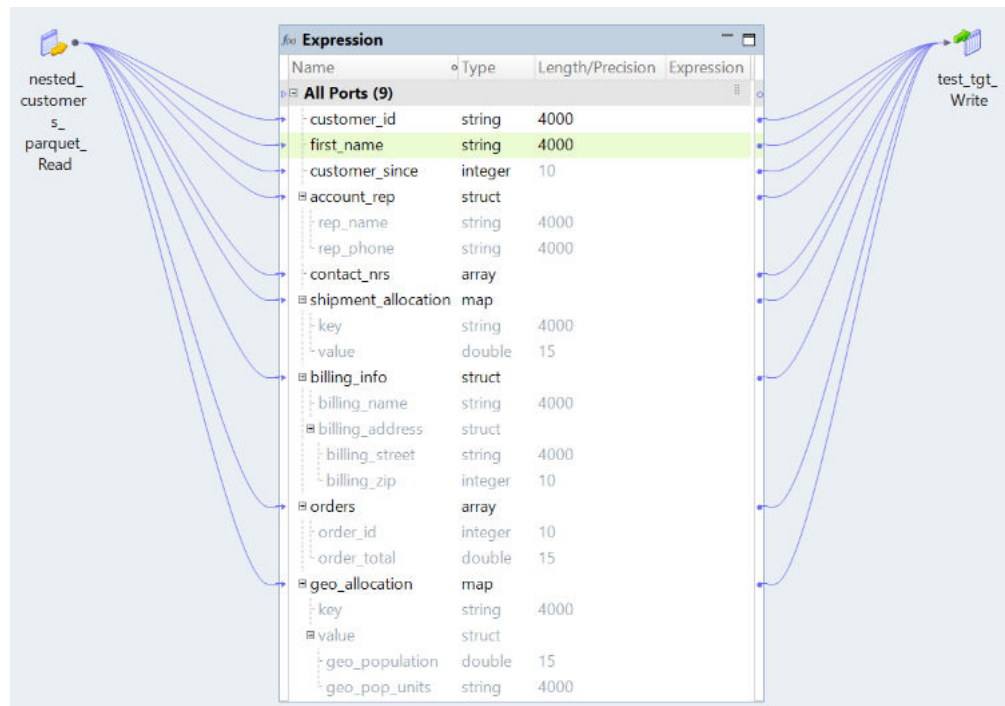
1. The Data Integration Service generates a mapping that includes a target based on the preview point.
2. It passes the mapping and the preview request to the Spark engine.
3. The Spark engine runs the mapping and stages the data based on the configured staging directories.
4. The Data Integration Service passes the staged data to the Developer tool and then deletes the staged data.
5. The results of the preview appear in the data viewer of the Developer tool.

Note: When you run data preview, the Data Integration Service validates the validation environments you have selected in the **Run-time** view.

Data Preview Interface

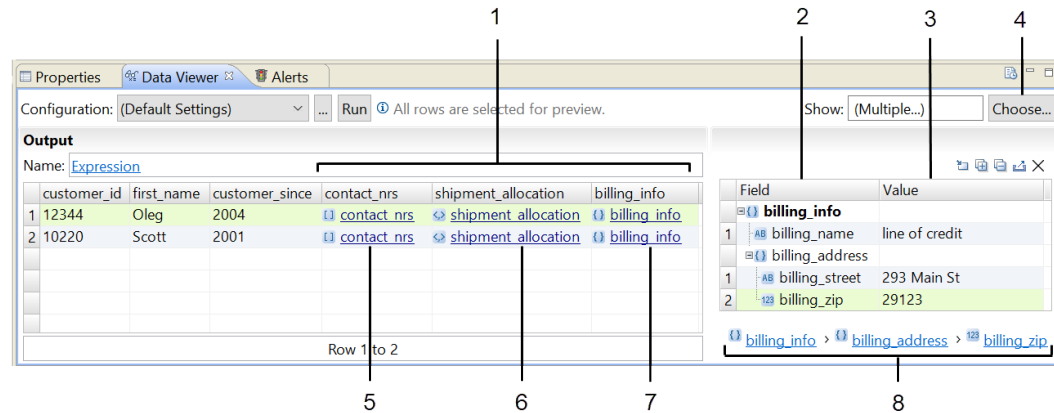
The data preview interface consists of Data Viewer view and hierarchical type panel. Based on the preview point you can preview hierarchical data.

The following image shows a mapping which has complex data types:



To view the hierarchical data after expressions have been evaluated, you use the Expression transformation as a preview point.

The following image shows the results of the data preview:



1. Hierarchy keys
2. Schema
3. Value
4. Port filter
5. Array data type
6. Map data type
7. Struct data type
8. Breadcrumbs

Data Viewer

When you preview the hierarchical data, you view ports in the preview point represented as columns in the Data Viewer. To identify the hierarchical types, the following punctuation icons are used:

- [] is used to indicate array data type.
- { } is used to indicate struct data type.
- <> is used to indicate map data type.

You can also view tooltips to view the details regarding the type of hierarchical data. If you need to limit the number of ports that appear, you can choose the ports you want to view in the output.

Highlight a field to see the details in the hierarchical type panel. When you right click any hierarchy field, you can export the data in that field in JSON format.

Exporting Data

You can export the data that appears in the **Data Viewer** view to a JSON file. Export data when you want to create a local copy of the data.

1. In the **Data Viewer** view, right-click the results and select **Export Data**.
2. Enter a file name.
3. Select the location where you want to save the file.
4. Click **OK**.

Hierarchical Type Panel

When you click a hierarchy field, the details of the hierarchy field appears in the hierarchical type panel.

In the hierarchical type panel, you can see the schema and values for the hierarchical data selected. You can navigate through the structure with the breadcrumbs. If the structure is very deep, you can expand and collapse the fields. When you expand the structure, you can view the data within the hierarchy. For example, when you view billing information in the hierarchical type panel, you can see the values for the billing name, street and ZIP Code.

Previewing Hierarchical Data

When you run data preview in the Data Viewer view, you can perform tasks, such as viewing the structure and data of a hierarchy field and filtering ports to view. You can also export the results to a JSON file.

Before you run data preview, add the complex file object to the mapping and verify that Spark is configured as the mapping run-time environment.

1. To perform hierarchical data preview, open a mapping in the Developer tool, and select a preview point.
2. To run the data preview, click Run in the Data Viewer view.

The output of the data preview can be viewed in the Data Viewer view.

3. To filter ports in the output, choose the ports you want to view.
4. To view the structure and the data in the field, select a hierarchy field.
The details of the hierarchy field appears on the hierarchical type panel.

5. To export the results to a JSON file, right-click the hierarchy field and select **Export Data**.

Rules and Guidelines for Hierarchical Data Preview

Consider the following rules and guidelines when you work with hierarchical data preview:

- Dynamic mappings with hierarchical data types are not supported.
- Date fields will have an additional timestamp component (00:00:00) in the data viewer.
- You can view the logs in the following location: `<Install location>/logs/<node name>/services/DataIntegrationService/disLogs/ms`. You cannot view the logs from the Developer tool.
- You cannot preview hierarchical data on an Expression transformation configured for windowing.
- You cannot preview hierarchical data in transformations that contain multiple output groups.

Normalizer

Router

Union

Update Strategy

CHAPTER 10

Hierarchical Data Processing Configuration

This chapter includes the following topics:

- [Hierarchical Data Conversion, 181](#)
- [Convert Relational or Hierarchical Data to Struct Data, 182](#)
- [Convert Relational or Hierarchical Data to Nested Struct Data, 184](#)
- [Extract Elements from Hierarchical Data, 192](#)
- [Flatten Hierarchical Data, 194](#)

Hierarchical Data Conversion

In the Developer tool, use **Hierarchical Conversion** wizards that simplify the tasks for developing a mapping to process hierarchical data on the Spark engine. These wizards add one or more transformations to the mapping and create output ports. The expressions in the output ports use complex functions and operators to convert relational data to hierarchical data or hierarchical data to relational data.

The following table lists the hierarchical conversion wizards and when to use these wizards:

Hierarchical Conversion Wizard	When to Use
Create Struct Port	To convert relational and hierarchical data to struct data. For example, you want to convert data in one or more columns of any data type to hierarchical data column of type struct.
Create Nested Complex Port	To convert relational and hierarchical data in two tables to a nested struct data. You can select one or more columns of any data type. For example, you want to convert data in columns from two tables to a hierarchical data column of type struct with a nested schema. The wizard creates a struct with an array of structs.

Hierarchical Conversion Wizard	When to Use
Extract from Complex Port	To convert hierarchical data to relational data or modify the hierarchical data. For example, you want to perform the following conversions: - Convert one or more elements of primitive data types in a hierarchical data to relational data columns. - Convert one or more elements of complex data types in a hierarchical data to hierarchical data columns.
Flatten Complex Port	To convert hierarchical data to relational data. For example, you want to convert one or more elements of a hierarchical data to relational data columns.

Convert Relational or Hierarchical Data to Struct Data

You can convert relational or hierarchical data in one or more columns to hierarchical data of type struct. Use the **Create Struct Port** wizard in the Developer tool to perform the conversion.

For example, a relational table contains three columns city, state, and zip. You create a mapping to convert the data in the three columns to one hierarchical column. Select the three ports in the transformation and use the Create Struct Port wizard to generate struct data with the selected ports as its elements.

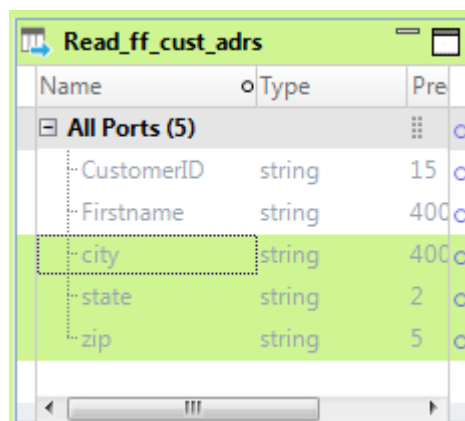
The wizard performs the following tasks:

- Creates a complex data type definition based on the ports that you select.
- Adds an Expression transformation to the mapping.
- Creates a struct output port to represent the struct data.
- Creates an expression that uses the STRUCT_AS function to generate struct data.

Creating a Struct Port

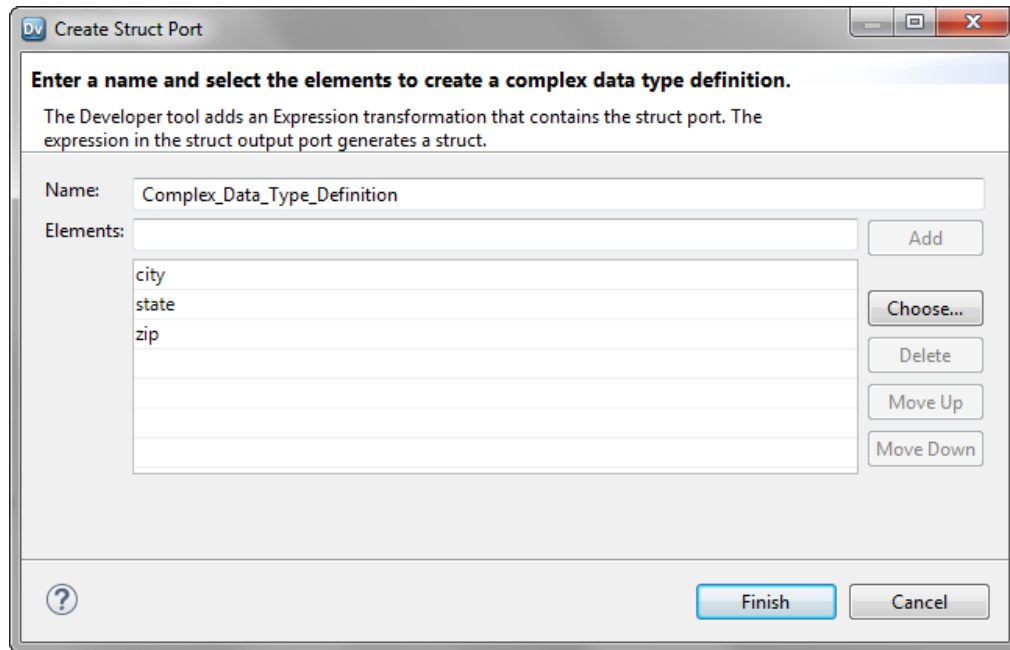
Use the **Create Struct Port** wizard to convert data that passes through one or more ports to struct data.

1. In the transformation, select one or more ports that you want to convert as elements of the struct data.
The ports you select also determine the elements of the complex data type definition.



2. Right-click the selected ports, and select **Hierarchical Conversions > Create Struct Port**.

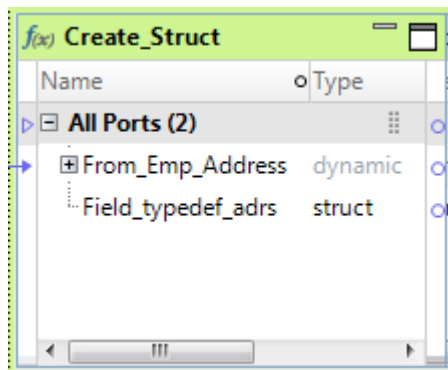
The **Create Struct Port** wizard appears with the list of ports that you selected.



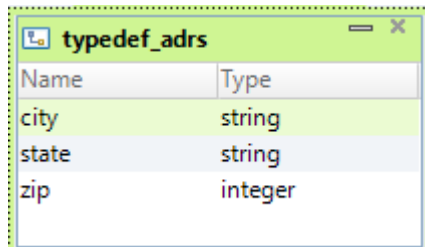
3. Optionally, in the **Name** box, change the name of the complex data type definition.
For example, typedef_address.
4. Optionally, click **Choose** to select other ports in the transformation.
5. Click **Finish**.

You can see the following changes in the mapping:

- The mapping contains a new Expression transformation `Create_Struct` with a struct output port and a dynamic port with ports from the upstream transformation.



- The type definition library contains the new complex data type definition.



- The struct output port references the complex data type definition.
- The struct output port contains an expression with the STRUCT_AS function. For example,

```
STRUCT_AS(:Type.Type_Definition_Library.typedef_address,city,state,zip)
```

Convert Relational or Hierarchical Data to Nested Struct Data

You can convert relational or hierarchical data in one or more columns in two transformations to nested struct data. Use the **Create Nested Complex Port** wizard in the Developer tool to perform the conversion.

The wizard converts relational data from two tables to struct data with a nested data type definition.

For example, a relational table contains employee bank account details. Another table contains the bank details. You create a mapping to convert data in the two tables into a hierarchical format that the payment system can process. Select the ports in the two transformations and use the Create Nested Complex Port wizard. The wizard generates struct data that contains an array element for each bank. The array contains a struct element for the employee bank account details.

Relational input

The Emp_Details table contains the following columns: employee ID, name, address, bank id, and bank account number.

The Bank_Details table contains the following columns: bank id, bank name, SWIFT code.

Struct output

```
banks{
  bank_swift:integer
  [bank_name{account_number:integer,employee_id:integer,name:string,address:string}]
}
```

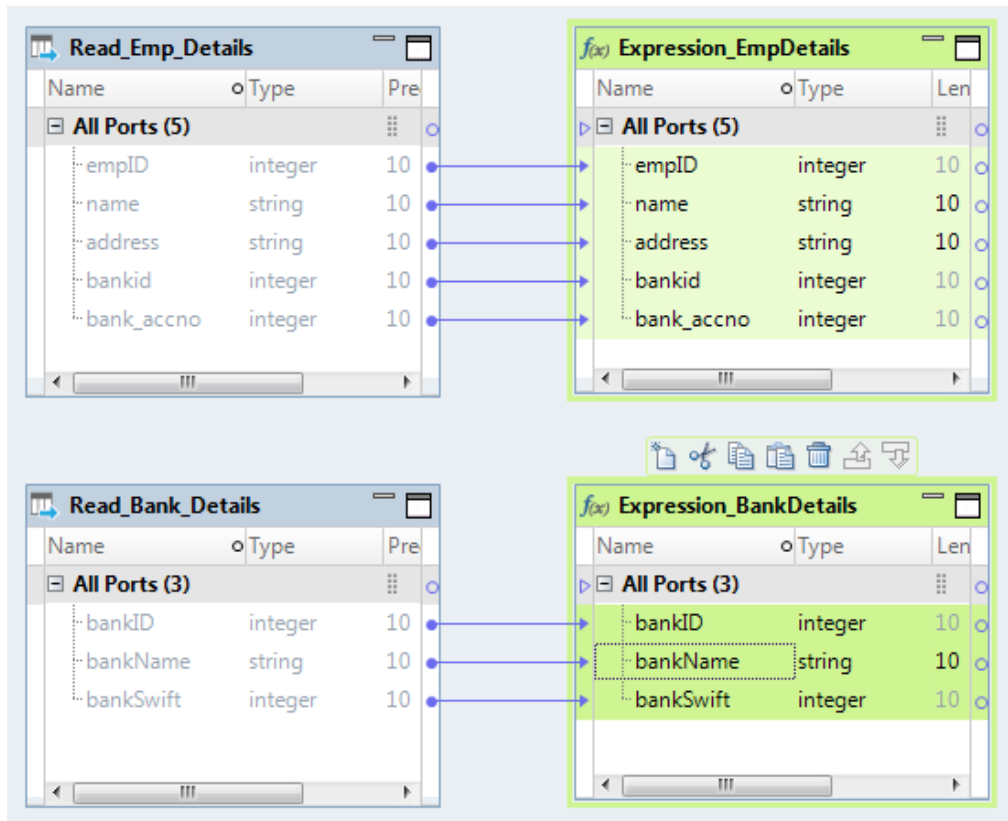
The wizard performs the following tasks:

- Adds a Joiner transformation to join source data from the parent and child tables based on the join keys. The transformation that contains ports for the parent complex data type definition is the parent transformation, and the other transformation is the child transformation.
- Creates a child complex data type definition and adds an Expression transformation that contains the child struct port. The expression in the struct port generates a struct based on the child complex data type definition.
- Adds an Aggregator transformation that contains an array port. The expression in the array port generates an array with the child struct as its element, and groups the struct values based on the group by port.
- Creates a parent complex data type definition and adds an Expression transformation that contains the nested complex port. The expression in the nested complex port generates a struct with an array of structs as one of its elements.

Creating A Nested Complex Port

Use the **Create Nested Complex Port** wizard to convert data that passes through one or more ports in two transformations to a struct data that references a nested data type definition. You specify the transformations from which the wizard creates the parent and child complex data type definitions.

1. Open the mapping or maplet in which you want to create a nested complex port.
2. Press Ctrl and select the ports in the two transformations that you want to convert as elements of the nested complex port.



3. Right-click the selected ports, and select **Hierarchical Conversions > Create Nested Complex Port**.

The **Create Nested Complex Port** wizard appears.

Specify the transformations from which to create parent and child complex data type definitions.
The Developer tool adds a Joiner transformation to join source data from the parent and child transformations.

Choose the transformation that contains the ports for the parent complex data type definition.

☒ Expression_EmpDetails
☐ Expression_BankDetails

Select a port from each transformation to use as keys to join the tables.

Parent Key:

Child Key:

? < Back Next > Finish Cancel

4. Choose the transformation that contains ports for the parent complex data type definition and select the join keys to join the tables.

Specify the transformations from which to create parent and child complex data type definitions.
The Developer tool adds a Joiner transformation to join source data from the parent and child transformations.

Choose the transformation that contains the ports for the parent complex data type definition.

☐ Expression_EmpDetails
☒ Expression_BankDetails

Select a port from each transformation to use as keys to join the tables.

Parent Key:

Child Key:

? < Back Next > Finish Cancel

5. Click **Next**.

The wizard page to create child complex data type definition appears.

Create Nested Complex Port

Enter a name and select the elements to create the child complex data type definition.

The Developer tool adds an Expression transformation that contains the child struct port. The expression in the output port generates a struct.

Name:

Elements:

empID	
name	
address	
bankid	
bank_accno	

Buttons: Add, Choose..., Delete, Move Up, Move Down

Navigation: ? < Back Next > Finish Cancel

6. Optionally, change the name of the child complex data type definition and make any changes to the elements.

Create Nested Complex Port

Enter a name and select the elements to create the child complex data type definition.

The Developer tool adds an Expression transformation that contains the child struct port. The expression in the output port generates a struct.

Name:

Elements:

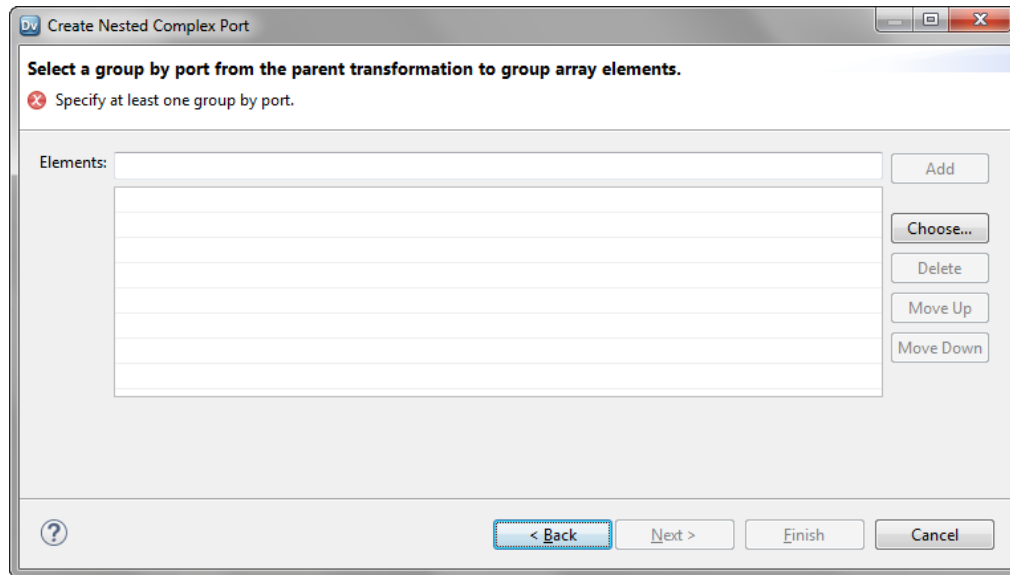
empID	
name	
address	
bank_accno	

Buttons: Add, Choose..., Delete, Move Up, Move Down

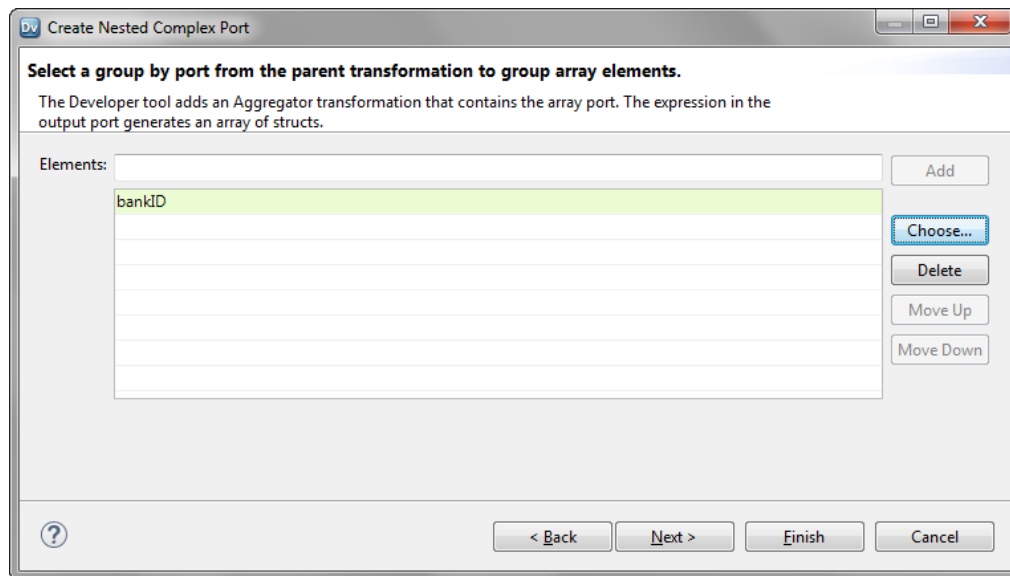
Navigation: ? < Back Next > Finish Cancel

7. Click **Next**.

The wizard page to select group by port from the parent transformation appears.



8. Click **Choose**.
The **Ports** dialog box appears.
9. Select a group by port from the parent transformation and click **OK**.
The selected port appears on the wizard page.



10. Click **Next**.

The final wizard page to create parent complex data type definition and nested complex port appears.

Enter a name and select the elements to create the parent complex data type definition.

The Developer tool adds an Expression transformation that contains the nested complex port, which is a struct port with an array of structs. The expression in the output port generates a struct.

Name:

Elements:

bankID	
bankName	
bankSwift	

Buttons: Add, Choose..., Delete, Move Up, Move Down

Bottom buttons: < Back, Next >, Finish, Cancel

11. Optionally, change the name of the parent complex data type definition and make any changes to the elements.

You cannot delete the port that you selected as the group by port.

Enter a name and select the elements to create the parent complex data type definition.

The Developer tool adds an Expression transformation that contains the nested complex port, which is a struct port with an array of structs. The expression in the output port generates a struct.

Name:

Elements:

bankID	
bankName	
bankSwift	

Buttons: Add, Choose..., Delete, Move Up, Move Down

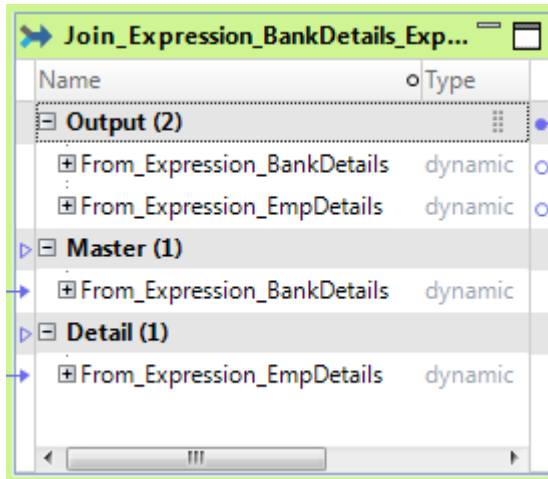
Bottom buttons: < Back, Next >, Finish, Cancel

Cannot delete the element because it is used as the group by port.

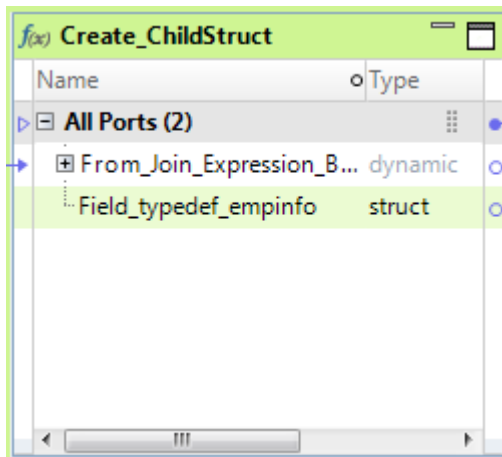
12. Click **Finish**.

You can see the following changes in the mapping:

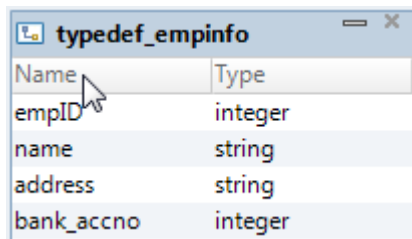
- The mapping contains a new Joiner transformation `Join_<transformation1>_<transformation1>` that joins the two tables.



- The mapping contains a new Expression transformation `Create_ChildStruct` with a struct output port and a dynamic port with ports from the upstream transformation.



- The type definition library contains the new child complex data type definition. The struct output port references the child complex data type definition.



The output port contains an expression with the `STRUCT_AS` function:

```
STRUCT_AS(:Type.Type_Definition_Library.<child_typedef>,<comma_delimited_child_elements>)
```

- The mapping contains a new Aggregator transformation `Agg_ArrayOfStruct` with the group by port to group the child structs into an array.

Name	Type
All Ports (2)	
From_Create_ChildStruct	dynamic
Field_typedef_empinfo_List	array

- The mapping contains a new Expression transformation `Final_NestedComplexPort` with a struct output port and a dynamic port with ports from the upstream transformation.

Name	Type
All Ports (2)	
From_Agg_ArrayOfStruct	dynamic
Field_typedef_bank	struct

- The type definition library contains the new parent complex data type definition, which is a nested data type definition. The struct output port references the parent complex data type definition.

Name	Type
bankID1	integer
bankName	string
bankSwift	integer
Field_typedef_empinfo_List	array

Name	Type
empID	integer
name	string
address	string
bank_accno	integer

The output port contains an expression with the `STRUCT_AS` function:

```
STRUCT_AS(:Type.Type_Definition_Library.<parent_typedef>,<comma_delimited_struct_elements>)
```

Extract Elements from Hierarchical Data

You can extract elements of a primitive or complex data type from hierarchical data. Use the **Extract from Complex Port** wizard in the Developer tool to perform the conversion.

Based on the data type of the elements in the complex port that you select, the wizard performs the following conversions:

- If the elements are of primitive data types, the wizard converts the element to a relational data.
- If the elements are of complex data types, the wizard converts the element to a hierarchical data.

The wizard adds an Expression transformation that contains one or more extracted ports. The number of elements that you select to extract determine the number of output ports in the transformation. The data type of the element determines the data type of the output port. The expression in the output ports use complex operators to extract elements from the complex port.

The following table describes the wizard conversions based on the data type of the complex port that you select:

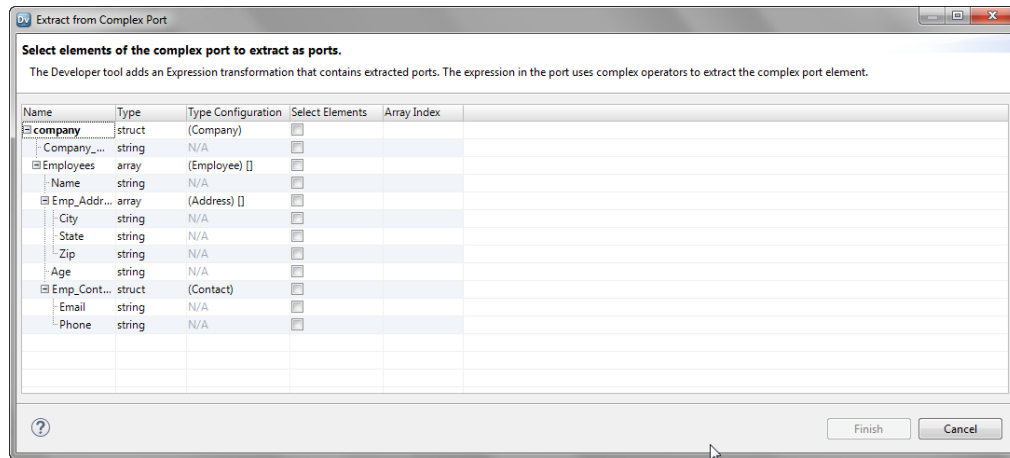
Complex Port Data Type	Wizard Conversion
array	Relational data. If you specify an array index, the wizard extracts an element in the array. Default is 0. The wizard extracts the first element of the array. Hierarchical data. If you do not specify an array index, the wizard extracts the entire array.
struct	Relational data. The wizard extracts the element in the struct that you selected.
array of structs	Relational data. If you select an element in the struct, the wizard extracts the element. The wizard requires an array index value. Default is 0. Hierarchical data. If you select the array and specify an array index, the wizard extracts the struct element in the array. If you do not specify an array index, the wizard extracts the entire array.
array of maps	Relational data. If you select the key or value element in the array and specify an array index, the wizard extracts the element. Hierarchical data. If you select the array and specify an array index, the wizard extracts the map element in the array. If you do not specify an array index, the wizard extracts the entire array.
struct of structs	Relational data. If you select an element in the parent or the child struct, the wizard extracts the element. Hierarchical data. If you select the parent struct or child struct, the wizard extracts that struct.
struct of maps	Hierarchical data. If you select the map element, the wizard extracts the map data. If you select the struct, the wizard extracts the struct data in another port.

Extracting Elements from a Complex Port

Use the **Extract from Complex Port** wizard to extract elements from hierarchical data.

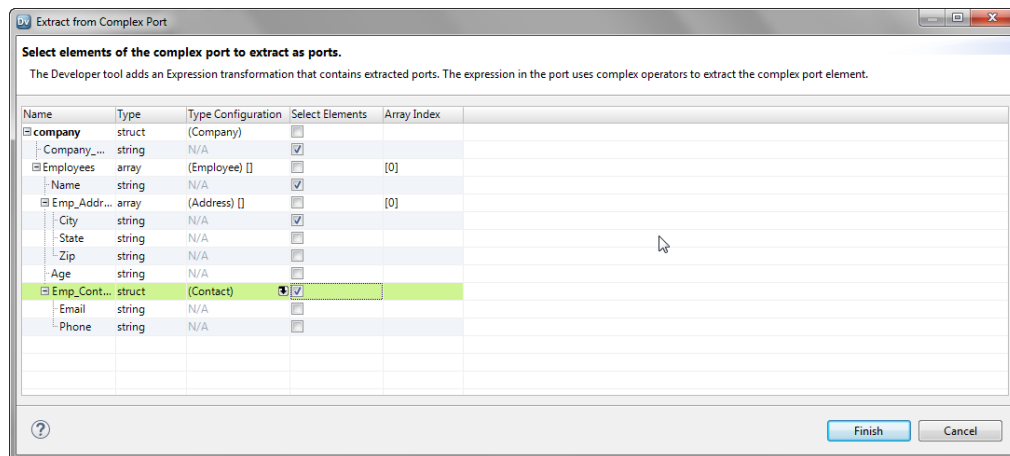
1. In the transformation, select a complex port from which you want to extract elements.
2. Right-click the selected ports, and select **Hierarchical Conversions > Extract from Complex Port**.

The **Extract from Complex Port** wizard appears with the list of elements in the complex port.



3. In the **Select Elements** column, select the check box for each element that you want to extract.
4. To extract an array element, specify an array index in the **Array Index** column.

If you delete the array index, the wizard extracts the entire array. To extract an element in the struct from a complex port for an array of structs, you must specify an array index value.

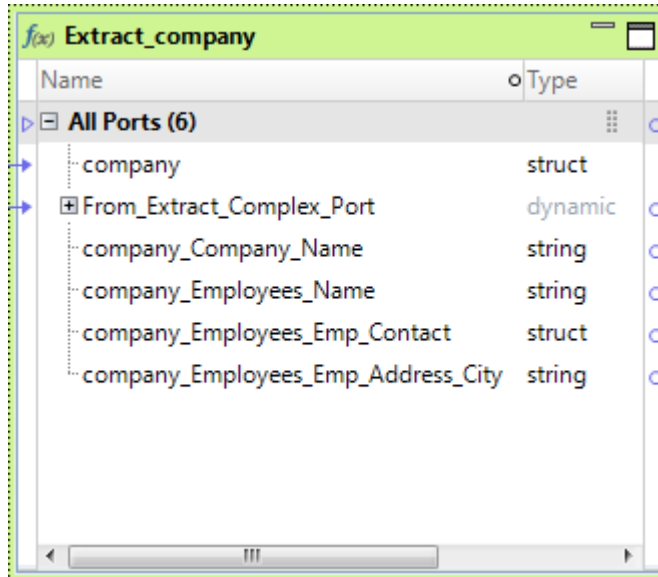


5. Click **Finish**.

You can see the following changes in the mapping:

- The mapping contains a new Expression transformation `Extract_<complex_port_name>` with the following ports:
 - The complex port from which you want to extract elements as the input port.
 - One or more output ports for the extracted elements. The number of elements that you selected to extract determines the number of output ports in the transformation.

- A dynamic port with ports from the upstream transformation.



- The output ports contain expressions that use complex operators to extract elements. The following image shows the expressions for the output ports in the Expression column on the Ports tab:

Name	Type	Type Configuration	Precisi...	Scale	Input	Output	Variable	Expression
1 company	struct	(Company)			☒	☐		company
2 From_Extract_Complex_Port	dynamic	N/A			☒	☒		From_Extract_Complex_Port
3 company_Company_Name	string	N/A	10	0	☐	☒		company.Company_Name
4 company_Employees_Name	string	N/A	10	0	☐	☒		company.Employees[0].Name
5 company_Employees_Emp_Contact	struct	(Contact)			☐	☒		company.Employees[0].Emp_Contact
6 company_Employees_Emp_Address_City	string	N/A	10	0	☐	☒		company.Employees[0].Emp_Address[0].City

Flatten Hierarchical Data

You can flatten elements of hierarchical data into relational data. Use the **Flatten Complex Port** wizard in the Developer tool to perform the conversion.

The wizard converts hierarchical data to relational data. When you have hierarchical data with nested data type, you can select specific elements or all elements of complex data type to flatten.

Based on the data type of the complex port, the wizard performs the following tasks:

struct

- Adds an Expression transformation with flattened output ports. The expression for the output ports uses the dot operator to extract elements in the struct.
- Adds a final Expression transformation that contains a dynamic port with all ports from the upstream transformation including the flattened struct ports.

array

Adds a Normalizer transformation with flattened output ports. The wizard flattens the array field in the Normalizer view.

Adds a final Expression transformation that contains a dynamic port with all ports from the upstream transformation including the flattened array ports.

nested data type

Adds one or more Expression and Normalizer transformations with flattened output ports. If you select a child element of a nested complex port, the wizard flattens both the parent and child elements.

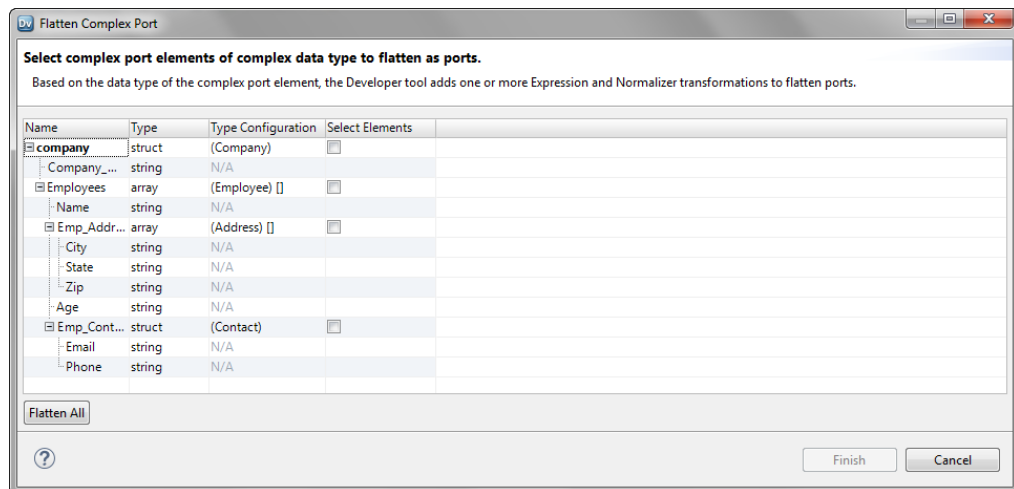
Adds a final Expression transformation that contains a dynamic port with all ports from the upstream transformation including the flattened ports.

Flattening a Complex Port

Use the **Flatten Complex Port** wizard to convert hierarchical data to relational data.

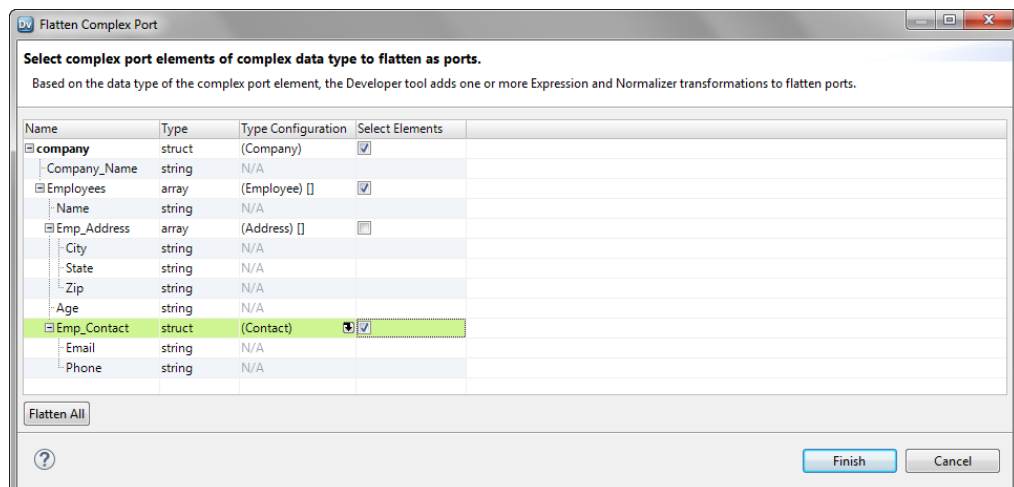
1. In the transformation, select a complex port that you want to flatten.
2. Right-click the selected ports, and select **Hierarchical Conversions > Flatten Complex Port**.

The **Flatten Complex Port** wizard appears with the list of elements in the complex port.



3. In the **Select Elements** column, select the check box for each element of a struct or an array data type that you want to extract.

If you select a child element of a nested complex port, the wizard automatically selects the parent element to flatten.

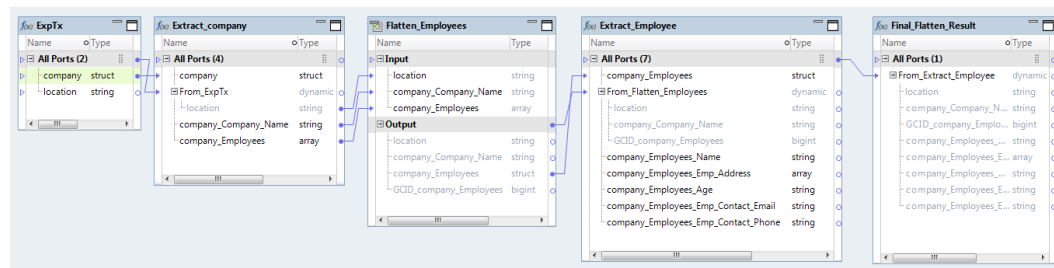


4. Optionally, click **Flatten All** to flatten all elements of struct or data type in the complex port. The wizard selects all elements of struct or data type in the complex port to flatten.
5. Click **Finish**.

You can see the following changes in the mapping:

- For each struct element that you selected to flatten, the mapping contains a new Expression transformation `Extract_<element_of_struct_type>` with the following ports:
 - The complex port from which you want to extract elements.
 - One or more output ports for the extracted elements.
 - The output ports contain expressions that use complex operators to extract elements.
 - A dynamic port with ports from the upstream transformation
- For each array element that you selected to flatten, the mapping contains a new Normalizer transformation `Flatten_<element_of_array_type>` with an input group and an output group. The output group contains the flattened normalized fields for the array element.
- A final Expression transformation `Final_Flatten_Result` with a dynamic port that contains the flattened ports for the complex port and other ports in the upstream transformation.

The following image shows an example of mapping changes:



CHAPTER 11

Hierarchical Data Processing with Schema Changes

This chapter includes the following topics:

- [Overview of Hierarchical Data Processing with Schema Changes, 197](#)
- [How to Develop a Dynamic Mapping to Process Schema Changes in Hierarchical Data, 198](#)
- [Dynamic Complex Ports, 199](#)
- [Input Rules for a Dynamic Complex Port, 201](#)
- [Port Selectors for Dynamic Complex Ports, 204](#)
- [Dynamic Expressions, 204](#)
- [Complex Operators, 206](#)
- [Complex Functions, 206](#)
- [Rules and Guidelines for Dynamic Complex Ports, 206](#)
- [Optimized Mappings, 206](#)

Overview of Hierarchical Data Processing with Schema Changes

You can use dynamic complex ports to manage schema changes to hierarchical data. You use dynamic ports to receive new and changed hierarchical data columns from a dynamic complex file source. Then, to manage schema or metadata changes of a complex port, you use dynamic complex ports.

To process hierarchical data on the Spark engine, you use complex ports in transformations. To handle metadata changes at run time, you create dynamic mappings with dynamic ports. A dynamic port can propagate hierarchical data through a generated port of a complex data type. However, if the schema of the hierarchical data change at run time, you must use dynamic complex ports to pass hierarchical data in the mapping.

Complex ports require you to specify the type configuration based on the data type of the elements in an array or a map or the schema of a struct. With dynamic complex ports, you do not have to specify the type configuration for a hierarchical column. For example, you do not have to specify a complex data type definition for a dynamic struct.

With dynamic complex ports, you can perform the following tasks in a dynamic mapping:

- Receive new or changed elements of a complex port if the schema of the complex port changes at run time.
- Modify hierarchical data or filter elements of a hierarchical data with input rules for a dynamic complex port.
- Extract elements of a hierarchical data that might have schema changes at run time using complex operators.
- Generate arrays and structs with dynamic expressions that use complex functions.

For more information about dynamic mappings, see the *Informatica Developer Mapping Guide*.

Note: You cannot process hierarchical data with schema changes on the Databricks Spark engine.

How to Develop a Dynamic Mapping to Process Schema Changes in Hierarchical Data

To process hierarchical data with schema changes on the Spark engine, develop a dynamic mapping with dynamic complex sources and targets, dynamic ports, complex ports, dynamic complex ports, complex operators and functions, port selectors, and dynamic expressions.

Note: The tasks and the order in which you perform the tasks to develop a dynamic mapping depend on the mapping scenario.

The following list outlines the high-level tasks to develop and run a mapping to read, write, and process hierarchical data that can change at run time.

Create dynamic complex sources and targets.

If a complex file source changes, you can configure the Read transformation to accommodate changes. For example, you can configure the Read transformation to use a different data source or to update the data object based on the data source. If a complex file target changes, you can configure the Write transformation to accommodate target changes. For example, you can configure the Write transformation to generate columns based on the mapping flow.

You can configure a Read or a Write transformation to get columns from a complex file data object at run time.

Create dynamic ports and define input rules.

When you configure transformations in a mapping, you create dynamic ports to receive new or changed columns based on the data flow. You configure input rules to determine the columns that a dynamic port receives and to rename or reorder the generated ports.

Create dynamic complex ports and define input rules.

Create a dynamic complex port to handle changes to the schema of a complex port. Configure input rules to determine the schema or element types of a dynamic complex port.

Create port selectors and dynamic expressions for dynamic complex output ports.

Create a dynamic expression by using dynamic ports or port selectors in the expressions. When you include a dynamic port, the expression generates a complex data type with elements based on each port that the dynamic port generates. When you include a port selector, the expression generates a complex data type with elements based on each port determined by the selection rule.

Create parameters.

You can use parameters to change values at run time. Use parameters to change values such as sources, targets, connections, and rules within the mapping.

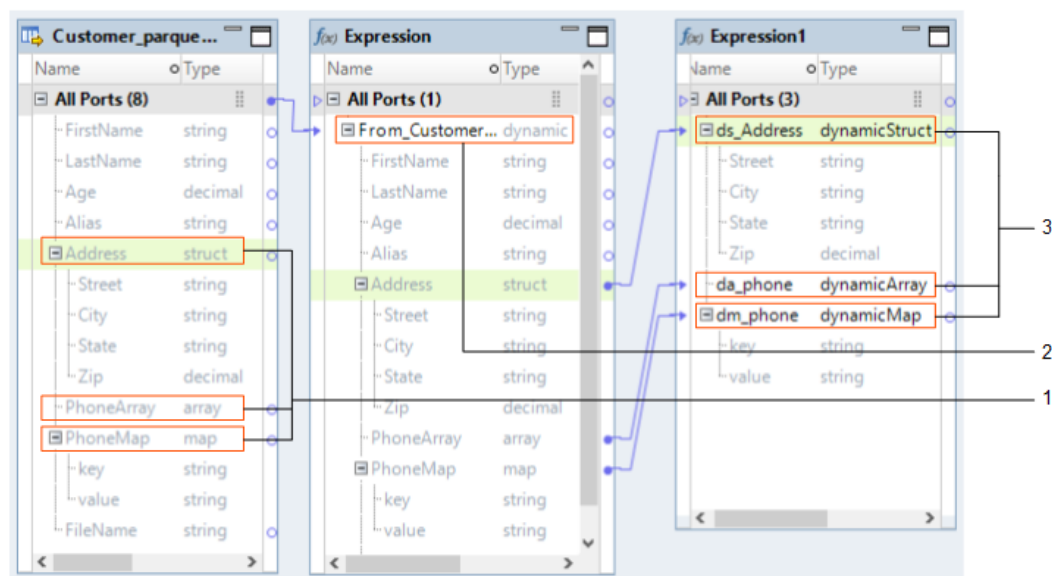
Create run-time links.

Transformations might change in such a way that you cannot create direct links when you design the mapping. If you cannot create links at design time, you can configure run-time links.

Dynamic Complex Ports

A dynamic complex port receives one or more elements of a complex port or a hierarchical data column based on input rules. You can use dynamic complex ports such as dynamic array, dynamic map, and dynamic struct in some transformations on the Spark engine.

The following image shows the ports that you can use in a mapping to process hierarchical data with schema changes:



1. Complex port
2. Dynamic port
3. Dynamic complex port

Complex port

A port of a complex data type in a transformation that can pass or process hierarchical data. Array, map, and struct ports are complex ports. Based on the complex data type, you specify the complex port properties as the type configuration.

Dynamic port

A port in a transformation that can receive one or more columns from an upstream transformation. Dynamic ports can receive new or changed columns as generated ports based on the metadata that passes through the mapping. The input rules for a dynamic port determine the generated ports.

Dynamic complex port

A port in a transformation that can pass or process hierarchical data with new or changed elements based on the schema changes of the hierarchical data. The input rules for a dynamic complex port determine the data type of the hierarchical data elements or the schema of the hierarchical data.

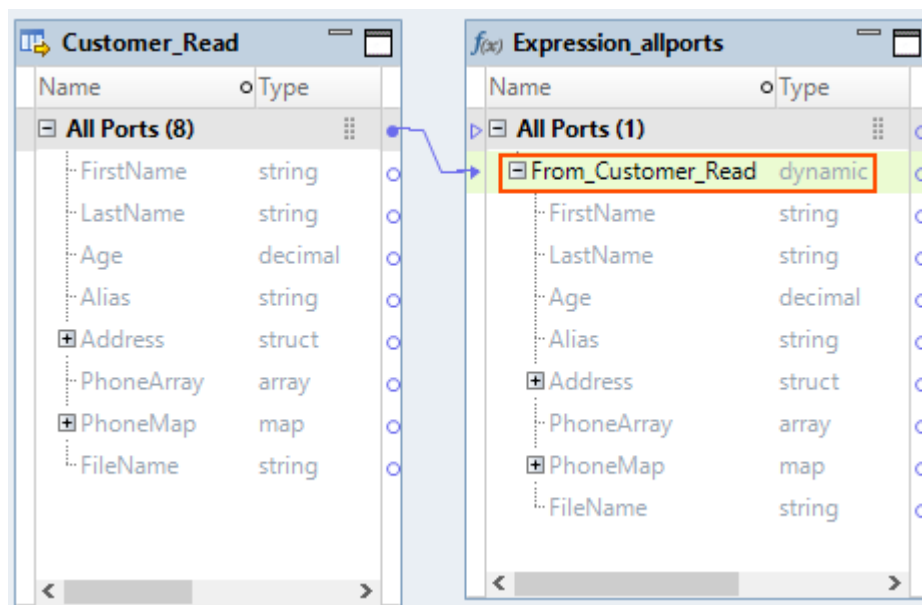
Dynamic Ports and Dynamic Complex Ports

Dynamic complex ports are different from dynamic ports. A dynamic complex port represents one hierarchical column of a complex data type, whereas a dynamic port represents one or more columns of any data type.

You create dynamic ports to accommodate changes to columns at run time based on the metadata that passes through the mapping. You create dynamic complex ports to accommodate changes to the schema of the hierarchical column at run time.

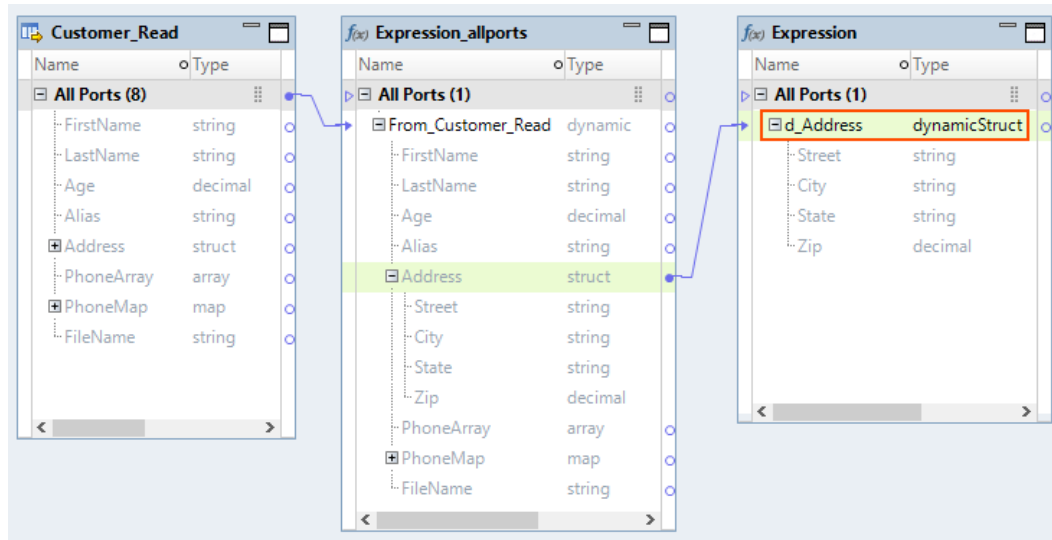
You cannot use dynamic complex ports in a mapplet, whereas you can use dynamic ports in a mapplet. For example, you can add a dynamic port to transformations in a mapplet. But you cannot use dynamic ports in the input and output instances of a mapplet.

The following image shows a dynamic port `From_Customer_Read`:



The dynamic port `From_Customer_Read` in the `Expression_allports` transformation includes generated ports of type string, decimal, struct, array, and map. The generated ports can change at run time based on the columns in the upstream transformation.

The following image shows a dynamic complex port d_Address:



The dynamic struct port d_Address in the Expression transformation includes elements of type string and decimal. The elements of a dynamic complex port can change at run time based on the schema changes to the hierarchical column in the upstream transformation.

Dynamic Complex Ports in Transformations

You can create dynamic complex ports in some transformations on the Spark engine.

The following transformations support dynamic complex ports:

- Aggregator
- Expression
- Filter
- Joiner
- Lookup
- Normalizer
- Rank
- Router
- Sorter

Input Rules for a Dynamic Complex Port

Input rules for a dynamic complex port are one or more rules that determine the data type of the hierarchical data elements or the schema of the hierarchical data. Define input rules to modify hierarchical data that is represented as a complex port in an upstream mapping.

Based on the dynamic complex port, you define the input rules in the corresponding **Input Rules** dialog box as follows:

dynamic array

Select one or more data types for the array elements. The default input rule includes elements of all data types.

dynamic map

Select one or more data types for the key and value of map elements. The default input rules for key and map include elements of all data types.

dynamic struct

Specify which elements to include or exclude in a struct based on the element names or data type. Include all elements is the default input rule.

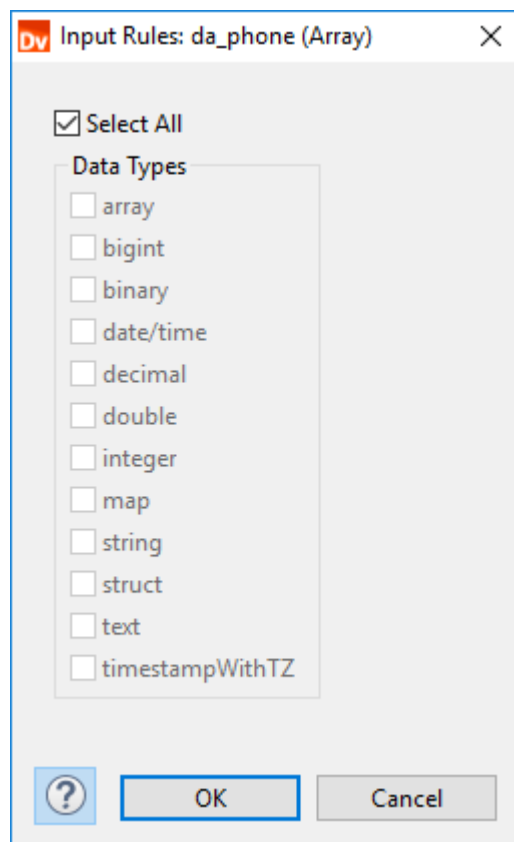
Input Rule for a Dynamic Array

An input rule for a dynamic array determines the data type of the array elements. You can select one or more data types. The default input rule includes elements of all data types.

For example, you develop a dynamic mapping to process two different complex file sources. You expect that an array in one of the complex files is of string elements and an array in the other complex file is of integer elements. Then, you can select string and integer data types as input rules for the dynamic array port.

If you are not certain about the data type of the array elements, you can select all data types.

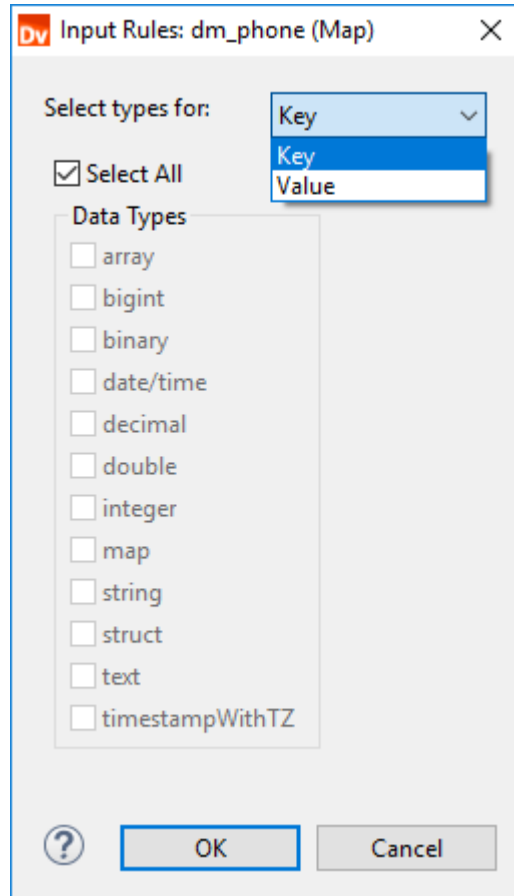
The following image shows the Dynamic Array Input Rule dialog box for a dynamic array port:



Input Rules for a Dynamic Map

An input rule for a dynamic map determines the data type of the key and value pair of a map. You can select one or more data types for the key and value. The default input rules for key and map include elements of all data types.

The following image shows the Input Rules dialog box for a dynamic map port:



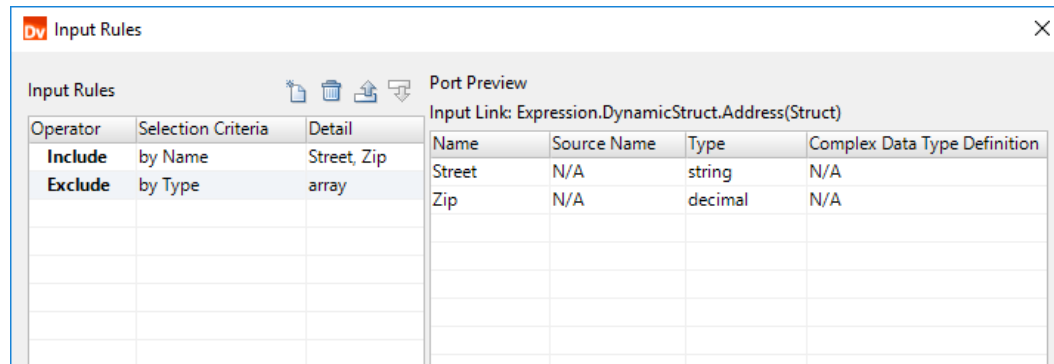
Input Rules for a Dynamic Struct

An input rule for a dynamic struct determines the schema of the struct. You can define the input rules based on the name or type of the struct element in an upstream transformation.

Specify which elements to include or exclude in a struct based on the element names or data type. You can define multiple rules. The Data Integration Service applies rules in the order in which they appear in the input rules list. Include all ports is the default input rule. Create at least one include input rule for a dynamic struct port.

Note: You cannot use input rules to rename elements of a struct port.

The following image shows the Input Rules dialog box for a dynamic struct port:



Port Selectors for Dynamic Complex Ports

You can create port selectors for dynamic complex ports to use in a dynamic expression or a join condition.

You can use port selectors in the following transformations for complex dynamic ports:

- Expression transformation. You can reference a port selector in STRUCT and ARRAY functions as a dynamic expression to construct a dynamic array or a dynamic struct.
- Joiner transformation. You can reference port selectors in STRUCT and ARRAY functions in a join condition.

For more information about port selectors, see the *Informatica Developer Mapping Guide*.

Dynamic Expressions

You can configure a dynamic expression to construct a dynamic complex port. Use STRUCT or ARRAY functions in dynamic expressions to construct a dynamic struct or a dynamic array. You can reference a port selector or a dynamic port in a dynamic expression for a dynamic complex output port.

A dynamic expression for an output port of a primitive data type runs the expression against each port to generate multiple output ports. A dynamic expression for a dynamic complex output port runs the expression to construct a dynamic complex port and returns a single output port of a complex data type.

You must specify a base port for the dynamic expression. Base port setting determines how the elements of the output complex port must look like. For example, elements of a port selector determine the elements of a dynamic output array. Base port for a struct determines the types of the elements of the output struct port.

Note: When you configure a dynamic expression, the Developer tool does not validate whether you have specified a base port.

Example - Dynamic Expression to Construct a Dynamic Struct

An Expression transformation has a dynamic port from which you want to construct a struct.

An Expression transformation has the following generated input ports:

Name	String
Street	String
City	String
State	String
Zip	Integer

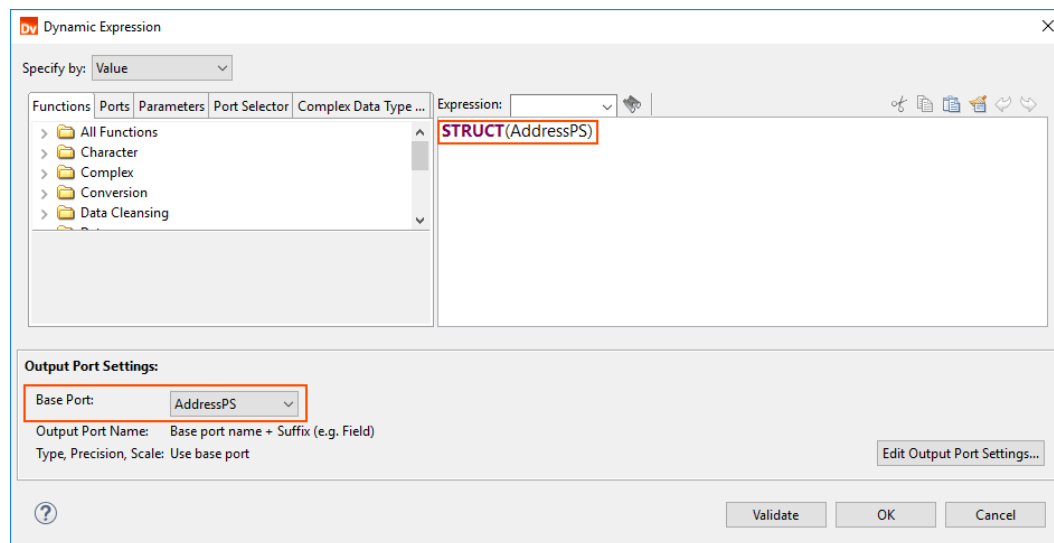
The transformation contains a dynamic complex output port called d_structAddress. You want to construct a struct and return the result to d_structAddress.

You create a port selector AddressPS and define the selection rules to include the following ports by name:

Street	String
City	String
State	String
Zip	Integer

Then, you create a dynamic expression with the STRUCT function.

The following image shows a dynamic expression that references a port selector AddressPS:



In the **Output Port Settings** area, specify the port selector AddressPS as the base port.

The dynamic expression creates the following struct with the elements from the port selector:

```
d_structAddress{
  Street    String
  City      String
  State     String
  Zip       Integer
}
```

Complex Operators

You can use complex operators in expressions to access or extract elements of a dynamic complex port.

For example, you can use the dot and subscript operators as follows to access elements of a dynamic complex port:

```
dynamicStruct.element  
dynamicArray[0]  
dynamicMap[key]
```

For more information about complex operators, see [“Complex Operators” on page 174](#).

Complex Functions

You can use ARRAY and STRUCT functions in dynamic expressions to construct dynamic arrays and dynamic structs. You use a dynamic port or a port selector as the complex function argument.

When you construct a dynamic struct, the number of elements of the struct changes based on the number of ports in the dynamic port or the port selector.

When you construct a dynamic array, the size of the array changes based on the number of ports in the dynamic port or the port selector.

For more information about complex functions, see [“Complex Functions” on page 176](#).

Rules and Guidelines for Dynamic Complex Ports

Consider the following rules and guidelines when you work with dynamic complex ports:

- You cannot use dynamic complex ports in a maplet.
- Filter transformation. You cannot use a filter condition to compare dynamic complex ports. However, you can use complex operators to extract elements of a dynamic complex port and use them in a filter condition. For example, you can use the following filter condition:

```
dynamicStruct.field == 100
```

Optimized Mappings

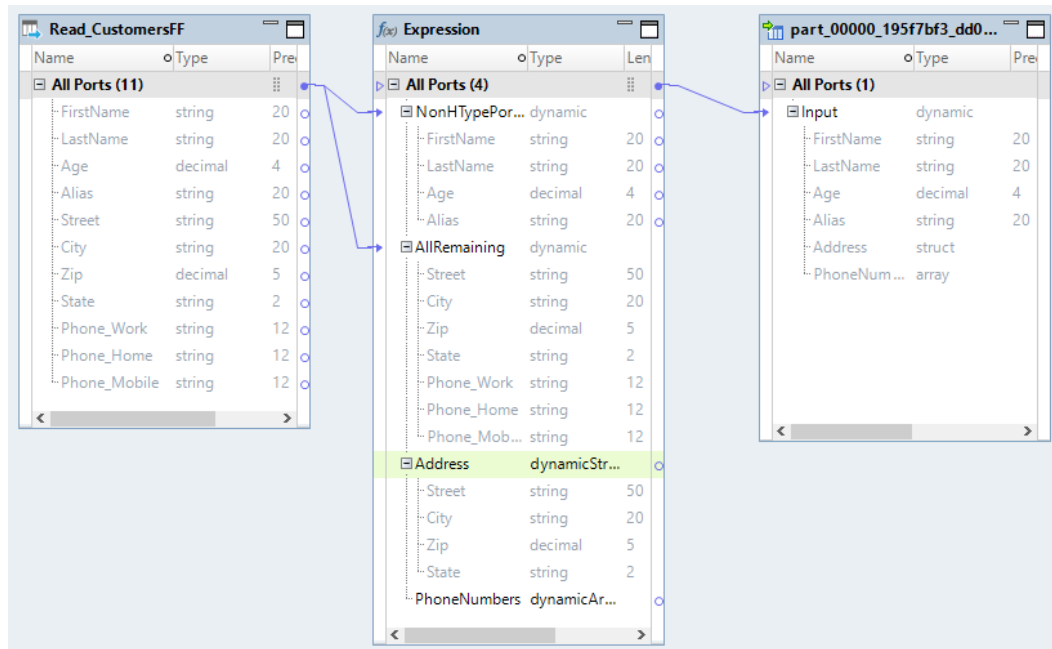
When you develop a dynamic mapping, you might want to review the dynamic mapping and fix any issues before you run the mapping. To review the mapping, you can view an optimized mapping. An optimized mapping is a compiled version of the dynamic mapping.

To view the optimized mapping, right-click an empty area in the mapping editor and click **Show Optimized Mapping**.

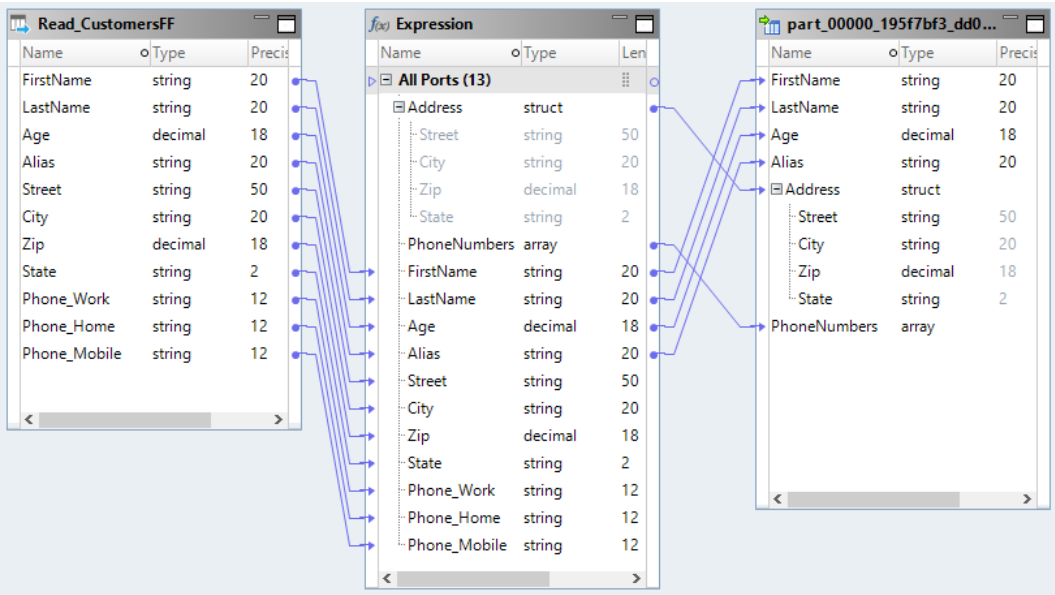
The Data Integration Service generates the optimized mapping. The Data Integration Service compiles the mapping to complete the following tasks:

- Expand dynamic ports and convert generated ports to static ports.
- Convert dynamic complex ports to complex ports.
- Resolve dynamic expressions.
- Resolve parameter values to default values.
- Link static ports.
- Resolve run-time links to connect the ports.

The following image shows a dynamic mapping with dynamic ports, dynamic complex ports, and a mapping flow link:



The following image shows the optimized mapping in which the generated ports are converted to static ports, dynamic complex ports are converted to complex ports, and the port links are resolved:



CHAPTER 12

Intelligent Structure Models

This chapter includes the following topics:

- [Overview of Intelligent Structure Models, 209](#)
- [Intelligent Structure Discovery Process, 210](#)
- [Use Case, 210](#)
- [Using an Intelligent Structure Model in a Mapping, 211](#)
- [Rules and Guidelines for Intelligent Structure Models, 211](#)
- [Before You Begin, 212](#)
- [How to Develop and Run a Mapping to Process Data with an Intelligent Structure Model , 212](#)
- [Creating an Informatica Intelligent Cloud Services Account, 215](#)
- [Creating an Intelligent Structure Model, 215](#)
- [Exporting an Intelligent Structure Model, 216](#)
- [Checking for Data Loss, 216](#)

Overview of Intelligent Structure Models

You can use CLAIRE™ Intelligent Structure Discovery to parse semi-structured or structured data in mappings that run on the Spark engine.

Long, complex files with little or no structure can be difficult to understand much less parse. CLAIRE Intelligent Structure Discovery can automatically discover the structure in unstructured data.

CLAIRE uses machine learning algorithms to decipher data in semi-structured or unstructured data files and create a model of the underlying structure of the data. You can generate an Intelligent structure model, a model of the pattern, repetitions, relationships, and types of fields of data discovered in a file, in Informatica Intelligent Cloud Services.

To use the model, you export it from Data Integration, and then can associate it with a data object in a Big Data Management mapping. You can run the mapping on the Spark engine to process the data. The mapping uses the Intelligent structure model to extract and parse data from input files based on the structure expressed in the model.

Intelligent Structure Discovery Process

You can create a CLAIRE™ Intelligent structure model in Intelligent Structure Discovery. Intelligent Structure Discovery is a service in Data Integration.

When you provide a sample file, Intelligent Structure Discovery determines the underlying structure of the information and creates a model of the structure. After you create an Intelligent structure model you can view, edit, and refine it. For example, you can select to exclude or combine structure elements. You can normalize repeating groups.

When you finish refining the model, you can export it and then associate it with a data object in a Big Data Management mapping.

The following image shows the process by which Intelligent Structure Discovery deciphers the underlying patterns of data and creates a model of the data patterns:



You can create models for semi-structured data from Microsoft Excel, Microsoft Word tables, PDF forms, and CSV files, or unstructured text files. You can also create models for structured data such as XML and JSON files.

You can quickly model data for files whose structure is very difficult, time consuming, and costly to find, such as log files, clickstreams, customer web access, error text files, or other internet, sensor, or device data that does not follow industry standards.

Use Case

You work in an operations group for an insurance company. Your team wants to process web logs to identify operations and security issues.

Your back-end system collects data on system access and security activities in various branches of the company. When the data is collected, the data is stored in the corporate data center and hosted in Amazon S3 storage.

Your team wants to understand the types of operations issues that have caused the most errors and system downtime in the past few weeks. You want to store data afterwards for auditing purposes.

Before your data analysts can begin working with the data, you need to parse the data in Amazon S3 input buckets and produce actionable data. But you cannot spend the time and resources required to sift through the data to create models of analysis. You might have to develop numerous mappings and parameter sets to parse the data to make sure that actionable data is created from the web logs.

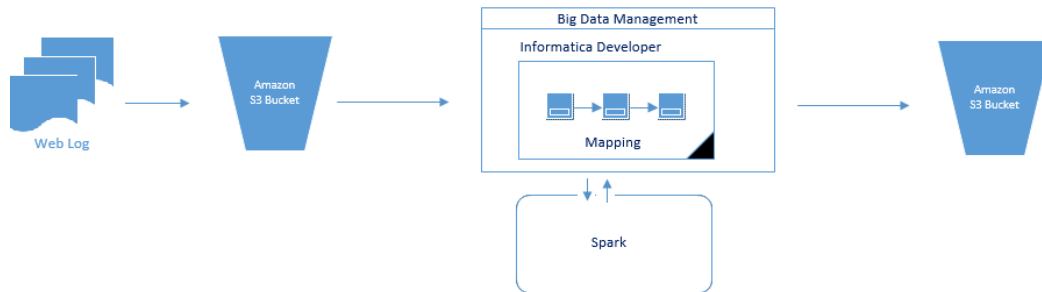
Instead of manually creating individual transformations, your team can use automatically generated Intelligent structure models to determine the relevant data sets. You create an Intelligent structure model in

Intelligent Structure Discovery, an application in Data Integration that uses machine learning algorithms to decipher data in structured or unstructured data files and discover the underlying structure of the data.

Intelligent Structure Discovery creates an Intelligent structure model that represents the input file data structure. You create a mapping with a data object that uses the intelligent structure model to output actionable data sets.

After the mapping fetches data from Amazon S3 input buckets, the mapping processes the data with an Intelligent structure model to prepare the data, and can then write the data to Amazon S3 output buckets.

The following image shows the process to fetch the data from Amazon S3 input buckets, parse and prepare the data, and then write the data to Amazon S3 output buckets. Analysts can then use the data to handle security issues and improve operations management.



Using an Intelligent Structure Model in a Mapping

To use an Intelligent structure model in a mapping, you associate it with a data object. The data object can then be used to process files that are similar to the one used to create the model.

You can add a Read transformation based on the data object to a mapping. If you want to process the data any further, such as checking data quality, or structuring the data into relational format, you add relevant downstream transformations to the mapping.

When the mapping runs, the Read transformation reads one or more input files and parses the data into fields, arrays, and structs.

Depending on the model and input, the data object output might contain primitive data types, complex data types, or nested data types. For more information about working with these data types in transformations, see [“Complex Data Types” on page 158](#).

Rules and Guidelines for Intelligent Structure Models

Consider the following rules and guidelines when you work with an Intelligent structure model:

- The model that you select for a data object should match the structure of the expected input files for the data object as much as possible. If an input file does not match the model, or partially matches the model, there might be a large amount of unidentified data and data loss. Therefore, when you create the model, it is important to choose a file that represents the expected input files.
- When you create an Intelligent structure model in Intelligent Structure Discovery, do not use duplicate names for different elements.

- Ensure that the Intelligent structure model is valid before you add it to the Column Projection properties of the data object properties.
- You can only use a Read transformation with an Intelligent structure model in a mapping. Do not create or use a Write transformation with an Intelligent structure model, as the mapping will fail.
- When you create the Intelligent structure model, select the Data Integration Version that corresponds to your current Big Data Management version.
- A data object with an Intelligent structure model parse PDF forms, Microsoft Word tables, and XML files whose size is less than the supported Hadoop split size of 256 MB.
- An Intelligent structure model can parse the data within PDF form fields but not data outside the fields.
- An Intelligent structure model can parse data within Microsoft Word tables. Other data is unidentified.

Before You Begin

Before you create intelligent structure models to use in Big Data Management, you need to meet certain requirements.

Make sure that you meet the following requirements:

- Ensure that you have an active Informatica Intelligent Cloud Services account.
- In Informatica Intelligent Cloud Services, configure a separate organization to create and export models to use in a mapping in the Developer tool.
- Verify that your Informatica Intelligent Cloud Services administrator has set the correct user-defined roles to create and export intelligent structure models.
- The version of the license of the Informatica Intelligent Cloud Services that you use to create intelligent structure models must correspond to your current Big Data Management version.

For more information about registration and roles, see the [Informatica Intelligent Cloud Services Administrator help](#).

How to Develop and Run a Mapping to Process Data with an Intelligent Structure Model

You can create a mapping with a data object that incorporates an intelligent structure model to parse data. You run the mapping on the Spark engine to process the data.

Note: The tasks and the order in which you perform the tasks to develop the mapping depend on the mapping scenario.

The following list outlines the high-level tasks to develop and run a mapping to read and process data in files of any type that an intelligent structure model can process, and then write the data to a target.

In Data Integration, a service in Informatica Intelligent Cloud Services, create an intelligent structure model.

Create an intelligent structure model using a representative file. Export the model to an .amodel file. After you save the file locally, you can copy it to the relevant file storage system.

For more information, see [Intelligent structure models](#) in the Data Integration help.

In Big Data Management, create a connection.

Create a connection to access data in files that are stored in the relevant system. You can create the following types of connections that will work with the data objects that can incorporate an intelligent structure:

- Hadoop Distributed File System
- Amazon S3
- Microsoft Azure Blob
- Microsoft Azure Data Lake Store

Create a data object with an intelligent structure model to read and parse source data.

1. Create a data object with an intelligent structure model to represent the files stored as sources. You can create the following types of data objects with an intelligent structure model:
 - Complex file
 - Amazon S3
 - Microsoft Azure Blob
 - Microsoft Azure Data Lake Store
2. Configure the data object properties.
3. In the data object read operation, configure columns to project hierarchical data as a complex data type. Enable the Project Column as Complex Data Type property in the Column Projection properties.

Note: You must have the relevant Informatica Intelligent Cloud Services license to import an intelligent structure model as a data object. For more information, see [“Before You Begin” on page 212](#).

Create a data object to write data to a target.

1. Create a data object to write the data to target storage.
2. Configure the data object properties.

Warning: Do not associate an intelligent structure model with a data object write operation. If you use a write operation that is associated with an intelligent structure model in a mapping, the mapping is not valid.

Create a mapping and add mapping objects.

1. Create a mapping.
2. Add a Read transformation based on the data object with the intelligent structure model.
3. Based on the mapping logic, add other transformations that are supported on the Spark engine. Link the ports and configure the transformation properties based on the mapping logic.
4. Add a Write transformation based on the data object that passes the data to the target storage or output. Link the ports and configure the transformation properties based on the mapping logic.

Configure the mapping to run on the Spark engine.

Configure the following mapping run-time properties:

1. Select Hadoop as the validation environment and Spark as the engine.
2. Select Hadoop as the execution environment and select a Hadoop connection.

Validate and run the mapping on the Spark engine.

1. Validate the mapping and fix any errors.
2. Optionally, view the Spark engine execution plan to debug the logic.

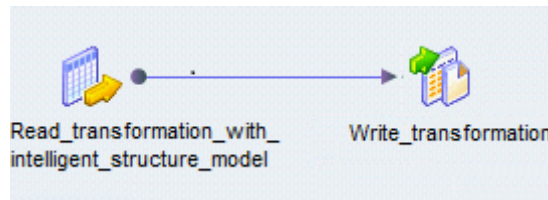
3. Run the mapping.

Mapping Example

Your organization needs to analyze purchase order details such as customer ID, item codes, and item quantity. The purchase order details are stored in Microsoft Excel spreadsheets in HDFS. The data must be changed into text files for storage. In addition, private customer data must be removed. Create a mapping that reads all the purchase records from the file in HDFS using a data object with an intelligent structure. The mapping must parse the data and write it to a storage target.

You can use the extracted data for auditing.

The following figure shows the example mapping:



You can use the following objects in the HDFS mapping:

HDFS input

The input object, `Read_transformation_with_intelligent_structure_model`, is a Read transformation that processed a Microsoft Excel file stored in HDFS and creates field output.

Amazon S3 output

The output object, `Write_transformation`, is a Write transformation that represents an Amazon S3 bucket.

When you run the mapping, the Data Integration Service reads the file in a binary stream and passes it to the Read transformation. The Read transformation extracts the relevant data in accordance to the intelligent structure and passes data to the Write transformation. The Write transformation writes the data to the storage target.

You can configure the mapping to run in a Hadoop run-time environment.

To configure the mapping, perform the following tasks:

1. Create an HDFS connection to read files from the Hadoop cluster.
2. Create a complex file data object read operation. Specify the following parameters:
 - The intelligent structure as the resource in the data object. The intelligent structure was configured so that it does not pass sensitive data.
 - The HDFS file location.
 - The input file folder location in the read data object operation.
3. Drag and drop the complex file data object read operation into a mapping.
4. Create an Amazon S3 connection.
5. Create a Write transformation for the Amazon S3 data object and add it to the mapping.

Creating an Informatica Intelligent Cloud Services Account

Create an Informatica Intelligent Cloud Services account if you do not already have one.

1. To create an account, access the Informatica Intelligent Cloud Services login page at the following link:
<https://dm-us.informaticacloud.com/identity-service/home>.
2. On the Informatica Intelligent Cloud Services login page, click **Don't have an account?**.
3. In the setup page, enter your contact details and email address. You can select to use your email address as a username.

Note: You will receive a registration email with access details and a link to confirm your account. You must confirm the account within 48 hours or register again.

4. In the registration email, click the account confirmation access link.
The Informatica Intelligent Cloud Services password setup window appears.
5. On the password setup window, define your Informatica Intelligent Cloud Services password, and then click **Log In**.

The Informatica Intelligent Cloud Services service picker appears. You can now access Data Integration.

Creating an Intelligent Structure Model

Create an Intelligent structure model in Data Integration. It is recommended to use a simplified sample file to generate the model. The sample file should have a representative data structure that is similar to the files that you want to parse.

For more information about creating and refining an Intelligent structure model, refer to the *Informatica Intelligent Cloud Services Mappings*.

1. In the Informatica Intelligent Cloud Services service picker, select **Data Integration**.
2. Click **New > Component > Intelligent Structure**.
3. In the **Intelligent Structures** page, click **New**.
4. In the **New Intelligent Structure** page, enter a name and description. You must provide a name for the Intelligent structure model.
5. To create an Intelligent structure model, browse for a sample file and click **Discover Structure**.
Intelligent Structure Discovery creates and displays an Intelligent structure model of the data.
6. To refine the Intelligent structure model and customize the structure of the output data, you can select a data node and select to combine, exclude, flatten, or collapse the node.
7. To save the Intelligent structure model, click **OK**.

Exporting an Intelligent Structure Model

After you create or edit a model, you can export the model to your local drive.

1. In Data Integration, select the **Explore** page.
2. On the **Explore** page, navigate to the project and folder with the Intelligent structure model.
The **Explore** page displays all the assets in the folder.
3. Click to select the row that contains the relevant Intelligent structure model. In the **Actions** menu, select **Edit**. In the **Intelligent Structure Details** panel, click **Edit**.
The Intelligent structure model appears in a separate page.
4. In the **Intelligent Structure** page, find the icon menu in the upper right corner of the Visual Model tab, and click the **Export model** icon.
Intelligent Structure Discovery downloads the `.amodel` file to your local drive.

Checking for Data Loss

To verify if a mapping with an Intelligent structure model lost data when it processed an input file, check the Spark history log.

- To get the log, run the following command on the Spark engine host machine:

```
yarn logs -applicationId<application ID>
```

CHAPTER 13

Stateful Computing

This chapter includes the following topics:

- [Overview of Stateful Computing, 217](#)
- [Windowing Configuration, 218](#)
- [Window Functions, 221](#)
- [Windowing Examples, 226](#)

Overview of Stateful Computing

You can use window functions in Expression transformations to perform stateful computations on the Spark engine.

A stateful computation is a function that takes some state and returns a value along with a new state.

You can use a window function to perform stateful computations. A window function takes a small subset of a larger data set for processing and analysis.

Window functions operate on a group of rows and calculate a return value for every input row. This characteristic of window functions makes it easy to express data processing tasks that are difficult to express concisely without window functions.

Use window functions to perform the following tasks:

- Retrieve data from upstream or downstream rows.
- Calculate a cumulative sum based on a group of rows.
- Calculate a cumulative average based on a group of rows.

Before you define a window function in an Expression transformation, you must describe the window by configuring the windowing properties. Windowing properties include a frame specification, partition keys, and order keys. The frame specification states which rows are included in the overall calculation for the current row. The partition keys determine which rows are in the same partition. The order keys determine how rows in a partition are ordered.

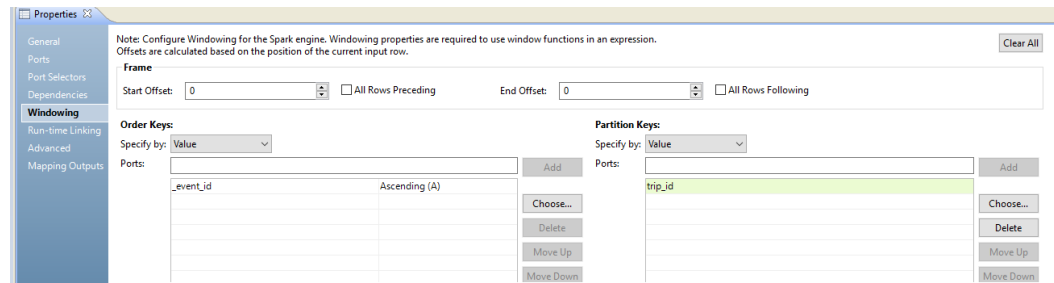
After you configure windowing properties, you define a window function in the Expression transformation. Spark supports the window functions LEAD and LAG. You can also use aggregate functions as window functions in an Expression transformation.

Windowing Configuration

When you include a window function in an Expression transformation, you configure the windowing properties associated with the function. Windowing properties define the partitioning, ordering, and frame boundaries associated with a particular input row.

Configure a transformation for windowing on the Windowing tab.

The following image shows the Windowing tab:



You configure the following groups of properties on the Windowing tab:

Frame

Defines the rows that are included in the frame for the current input row, based on physical offsets from the position of the current input row.

You configure a frame if you use an aggregate function as a window function. The window functions LEAD and LAG reference individual rows and ignore the frame specification.

Partition Keys

Separate the input rows into different partitions. If you do not define partition keys, all rows belong to a single partition.

Order Keys

Define how rows in a partition are ordered. The ports you choose determine the position of a row within a partition. The order key can be ascending or descending. If you do not define order keys, the rows have no particular order.

Frame

The frame determines which rows are included in the calculation for the current input row, based on their relative position to the current row.

If you use an aggregate function instead of LEAD or LAG, you must specify a window frame. LEAD and LAG reference individual row sand ignore the frame specification.

The start offset and end offset describe the number of rows that appear before and after the current input row. An offset of "0" represents the current input row. For example, a start offset of -3 and an end offset of 0 describes a frame including the current input row and the three rows before the current row.

The following image shows a frame with a start offset of -1 and an end offset of 1:

Type	Category	Revenue	
Action	Video game	1000	
Arcade	Video game	1000	← 1 PRECEDING
Sports	Video game	2000	
Adventure	Video game	3000	← 1 FOLLOWING
Strategy	Video game	4000	

Current input row →

For every input row, the function performs an aggregate operation on the rows inside the frame. If you configure an aggregate expression like SUM with the preceding frame, the expression calculates the sum of the values within the frame and returns a value of 6000 for the input row.

You can also specify a frame that does not include the current input row. For example, a start offset of 10 and an end offset of 15 describes a frame that includes six total rows, from the tenth to the fifteenth row after the current row.

Note: The start offset must be less than or equal to the end offset.

Offsets of **All Rows Preceding** and **All Rows Following** represent the first row of the partition and the last row of the partition. For example, if the start offset is All Rows Preceding and the end offset is -1, the frame includes one row before the current row and all rows before that.

The following figure illustrates a frame with a start offset of 0 and an end offset of All Rows Following:

Genre	Recordings	Revenue	
Jazz	233	5000	
Gospel	214	1000	
Country	145	2000	All Rows Following ↓
Ethnic	154	9000	
Pop	317	4000	
Rock	237	2100	
Classical	221	3200	
EDM	153	950	
Hip Hop	839	2300	
Punk	415	7650	

Current input row →

Partition and Order Keys

Configure partition and order keys to form groups of rows and define the order or sequence of rows within each partition.

Use the following keys to specify how to group and order the rows in a window:

Partition keys

Configure partition keys to define partition boundaries, rather than performing the calculation across all inputs. The window function operates across the rows that fall into the same partition as the current row.

You can specify the partition keys by value or parameter. Select **Value** to use port names. Choose **Parameter** to use a sort key list parameter. A sort key list parameter contains a list of ports to sort by. If you do not specify partition keys, all the data is included in the same partition.

Order keys

Use order keys to determine how rows in a partition are ordered. Order keys define the position of a particular row in a partition.

You can specify the order keys by value or parameter. Select **Value** to use port names. Choose **Parameter** to use a sort key list parameter. A sort key list parameter contains a list of ports to sort by. You must also choose to arrange the data in ascending or descending order. If you do not specify order keys, the rows in a partition are not arranged in any particular order.

Example

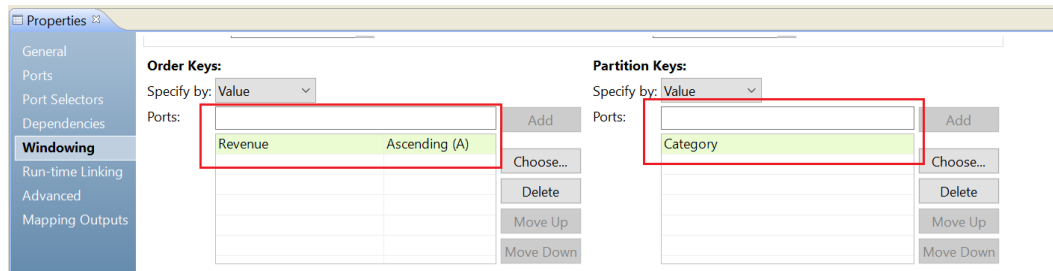
You are the owner of a coffee and tea shop. You want to calculate the best-selling and second best-selling coffee and tea products.

The following table lists the products, the corresponding product categories, and the revenue from each product:

Product	Category	Revenue
Espresso	Coffee	600
Black	Tea	550
Cappuccino	Coffee	500
Americano	Coffee	600
Oolong	Tea	250
Macchiato	Coffee	300
Green	Tea	450
White	Tea	650

You partition the data by category and order the data by descending revenue.

The following image shows the properties you configure on the Windowing tab:



The following table shows the data grouped into two partitions according to category. Within each partition, the revenue is organized in descending order:

Product	Category	Revenue
Espresso	Coffee	600
Americano	Coffee	600
Cappuccino	Coffee	500
Macchiato	Coffee	300
White	Tea	650
Black	Tea	550
Green	Tea	450
Oolong	Tea	250

Based on the partitioning and ordering specifications, you determine that the two best-selling coffees are espresso and Americano, and the two best-selling teas are white and black.

Rules and Guidelines for Windowing Configuration

Certain guidelines apply when you configure a transformation for windowing.

Consider the following rules and guidelines when you define windowing properties for a window function:

- When you configure a frame, the start offset must be less than or equal to the end offset. Otherwise, the frame is not valid.
- Configure a frame specification if you use an aggregate function as a window function. LEAD and LAG operate based on the offset value and ignore the frame specification.
- You cannot use complex ports as partition or order keys.
- Assign unique port names to partition and order keys to avoid run-time errors.
- The partition and order keys cannot use both a dynamic port and one or more generated ports of the same dynamic port. You must select either the dynamic port or the generated ports.

Window Functions

Window functions calculate a return value for every input row of a table, based on a group of rows.

A window function performs a calculation across a set of table rows that are related to the current row. You can also perform this type of calculation with an aggregate function. But unlike regular aggregate functions, a window function does not group rows into a single output row. The rows retain unique identities.

You can define the LEAD and LAG analytic window functions in an Expression transformation. LEAD and LAG give access to multiple rows within a table, without the need for a self-join.

LEAD

The LEAD function returns data from future rows.

LEAD uses the following syntax:

```
LEAD ( Column name, Offset, Default )
```

LEAD returns the value at an offset number of rows after the current row. Use the LEAD function to compare values in the current row with values in a following row. Use the following arguments with the LEAD function:

- Column name. The column name whose value from the subsequent row is to be returned.
- Offset. The number of rows following the current row from which the data is to be retrieved. For example, an offset of "1" accesses the next immediate row, and an offset of "3" accesses the third row after the current row.
- Default. The default value to be returned if the offset is outside the scope of the partition. If you do not specify a default, the default is NULL.

For more information about the LEAD function, see the *Informatica Transformation Language Reference*.

LAG

The LAG function returns data from preceding rows.

LAG uses the following syntax:

```
LAG ( Column name, Offset, Default )
```

LAG returns the value at an offset number of rows before the current row. Use the LAG function to compare values in the current row with values in a previous row. Use the following arguments with the LAG function:

- Column name. The column name whose value from the prior row is to be returned.
- Offset. The number of rows preceding the current row from which the data is to be retrieved. For example, an offset of "1" accesses the previous row, and an offset of "3" accesses the row that is three rows before the current row.
- Default. The default value to be returned if the offset is outside the scope of the partition. If you do not specify a default, the default is NULL.

For more information about the LAG function, see the *Informatica Transformation Language Reference*.

Aggregate Functions as Window Functions

In addition to LEAD and LAG, you can also use aggregate functions as window functions. When you use aggregate functions like SUM and AVG as window functions, you can perform running calculations that are similar to the stateful functions MOVINGSUM, MOVINGAVG, and CUME. Window functions are more flexible than stateful functions because you can set a specific end offset.

To use an aggregate function as a window function, you must define a frame in the windowing properties. You define a frame to limit the scope of the calculation. The aggregate function performs a calculation across the frame and produces a single value for each row.

Example

You are a lumber salesperson who sold different quantities of wood over the past two years. You want to calculate a running total of sales quantities.

The following table lists each sale ID, the date, and the quantity sold:

Sale_ID	Date	Quantity
30001	2016-08-02	10
10001	2016-12-24	10
10005	2016-12-24	30
40001	2017-01-09	40
10006	2017-01-18	10
20001	2017-02-12	20

A SUM function adds all the values and returns one output value. To get a running total for each row, you can define a frame for the function boundaries.

The following image shows the frame you specify on the Windowing tab:

Note: Configure Windowing for the Spark engine. Windowing properties are required to use window functions in an expression. Offsets are calculated based on the position of the current input row.

Frame

Start Offset: 0 ☒ All Rows Preceding End Offset: 0 ☐ All Rows Following

Order Keys:

Specify by: Value

Ports:

Order Key	Order
Date	Ascending (A)

Partition Keys:

Specify by: Value

Ports:

Partition Key

You configure the following windowing properties:

- Start offset: All Rows Preceding
- End offset: 0
- Order Key: Date Ascending
- Partition Key: Not specified

You define the following aggregate function:

```
SUM (Quantity)
```

SUM adds the quantity in the current row to the quantities in all the rows preceding the current row. The function returns a running total for each row.

The following table lists a running sum for each date:

Sale_ID	Date	Quantity	Total
30001	2016-08-02	10	10
10001	2016-12-24	10	20
10005	2016-12-24	30	50
40001	2017-01-09	40	90

Sale_ID	Date	Quantity	Total
10006	2017-01-18	10	100
20001	2017-02-12	20	120

Aggregate Offsets

An aggregate function performs a calculation on a set of values inside a partition. If the frame offsets are outside the partition, the aggregate function ignores the frame.

If the offsets of a frame are not within the partition or table, the aggregate function calculates within the partition. The function does not return NULL or a default value.

For example, you partition a table by seller ID and you order by quantity. You set the start offset to -3 and the end offset to 4.

The following image shows the partition and frame for the current input row:



The frame includes eight total rows, but the calculation remains within the partition. If you define an AVG function with this frame, the function calculates the average of the quantities inside the partition and returns 18.75.

Nested Aggregate Functions

A nested aggregate function in a window function performs a calculation separately for each partition. A nested aggregate function behaves differently in a window function and an Aggregator transformation.

A nested aggregate function in an Aggregator transformation performs a calculation globally across all rows. A nested aggregate function in an Expression transformation performs a separate calculation for each partition.

For example, you configure the following nested aggregate function in an Aggregator transformation and an Expression transformation:

```
MEDIAN ( COUNT ( P1 ) )
```

You define P2 as the group by port in the Aggregator transformation.

The function takes the median of the count of the following set of data:

P1	P2
10	1
7	1
12	1
11	2
13	2
8	2
10	2

RETURN VALUE: 3.5

When you include nested aggregate functions in an Expression transformation and configure the transformation for windowing, the function performs the calculation separately for each partition.

You partition the data by P2 and specify a frame of all rows preceding and all rows following. The window function performs the following calculations:

1. COUNT (P1) produces one value for every row. COUNT returns the number of rows in the partition that have non-null values.
2. MEDIAN of that value produces the median of a window of values generated by COUNT.

The window function produces the following outputs:

P1	P2	Output
10	1	3
7	1	3
12	1	3
11	2	4
13	2	4
8	2	4
10	2	4

You can nest aggregate functions with multiple window functions. For example:

```
LAG ( LEAD( MAX( FIRST ( p1 ) ) )
```

Note: You can nest any number of the window functions LEAD and LAG, but you cannot nest more than one aggregate function within another aggregate function.

Rules and Guidelines for Window Functions

Certain guidelines apply when you use window functions on the Spark engine.

Consider the following rules and guidelines when you define window functions in a transformation:

- Specify a constant integer as the offset argument in a window function.
- Specify a default argument that is the same data type as the input value.
- You cannot specify a default argument that contains complex data type or a SYSTIMESTAMP argument.
- To use the LEAD and LAG window functions, you must configure partition and order keys in the windowing properties.
- To use an aggregate function as a window function in an Expression transformation, you must configure a frame specification in the windowing properties.

Windowing Examples

The examples in this section demonstrate how to use LEAD, LAG, and other aggregate functions as window functions in an Expression transformation.

Financial Plans Example

You are a banker with information about the financial plans of two of your customers. Each plan has an associated start date.

For each customer, you want to know the expiration date for the current plan based on the activation date of the next plan. The previous plan ends when a new plan starts, so the end date for the previous plan is the start date of the next plan minus one day.

The following table lists the customer codes, the associated plan codes, and the start date of each plan:

CustomerCode	PlanCode	StartDate
C1	00001	2014-10-01
C2	00002	2014-10-01
C2	00002	2014-11-01
C1	00004	2014-10-25
C1	00001	2014-09-01
C1	00003	2014-10-10

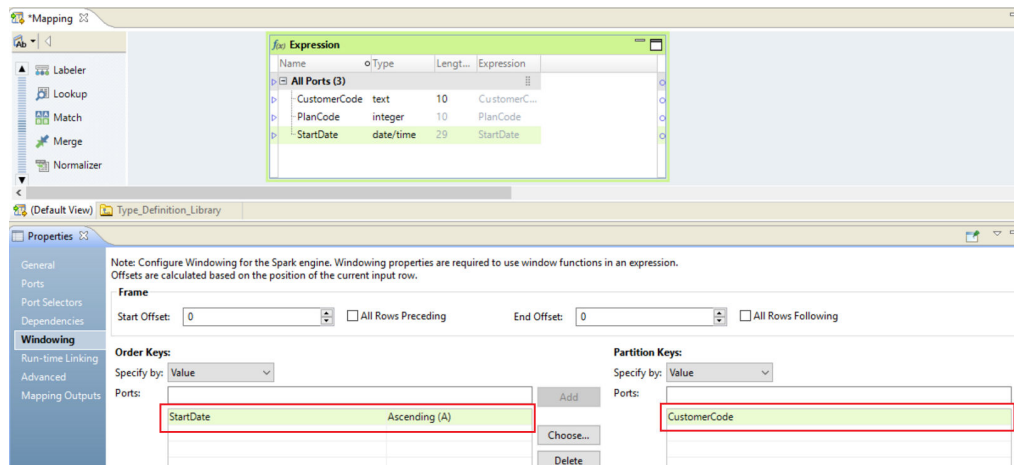
Define partition and order keys

You partition the data by customer code and order the data by ascending start date.

You configure the following windowing properties:

Property	Description
Order key	StartDate Ascending. Arranges the data chronologically by ascending start date.
Partition key	CustomerCode. Groups the rows according to customer code so that calculations are based on individual customers.
Frame	Not specified. Window functions access rows based on the offset argument and ignore the frame specification.

The following image shows the windowing properties you configure on the Windowing tab:



The following table lists the data grouped by customer code and ordered by start date:

CustomerCode	PlanCode	StartDate
C1	00001	2014-09-01
C1	00002	2014-10-01
C1	00003	2014-10-10
C1	00004	2014-10-25
C2	00001	2014-10-01
C2	00002	2014-11-01

The start dates for each customer are arranged in ascending order so that the dates are chronological.

Define a window function

You define a LEAD function to access the subsequent row for every input.

You define the following function on the Ports tab of the Expression transformation:

```
LEAD ( StartDate, 1, '01-Jan-2100' )
```

Where:

- `StartDate` indicates the target column that the function operates on.
- `1` is the offset. This value accesses the next immediate row.
- `01-Jan-2100` is the default value. The expression returns "01-Jan-2100" if the returned value is outside the bounds of the partition.

Define an `ADD_TO_DATE` function

You use an `ADD_TO_DATE` function to subtract one day from the date you accessed.

You define the following expression on the Ports tab of the Expression transformation:

```
ADD_TO_DATE ( LEAD ( StartDate, 1, '01-Jan-2100' ), 'DD', -1, )
```

By subtracting one day from the start date of the next plan, you find the end date of the current plan.

The following table lists the end dates of each plan:

CustomerCode	PlanCode	StartDate	EndDate
C1	00001	2014-09-01	2014-09-30
C1	00002	2014-10-01	2014-10-09
C1	00003	2014-10-10	2014-10-24
C1	00004	2014-10-25	2099-12-31*
C2	00001	2014-10-01	2014-10-31
C2	00002	2014-11-01	2099-12-31*

*The `LEAD` function returned the default value because these plans have not yet ended. The rows were outside the partition, so the `ADD_TO_DATE` function subtracted one day from 01-Jan-2100, returning 2099-12-31.

GPS Pings Example

Your organization receives GPS pings from vehicles that include trip and event IDs and a time stamp. You want to calculate the time difference between each ping and flag the row as skipped if the time difference with the previous row is less than 60 seconds.

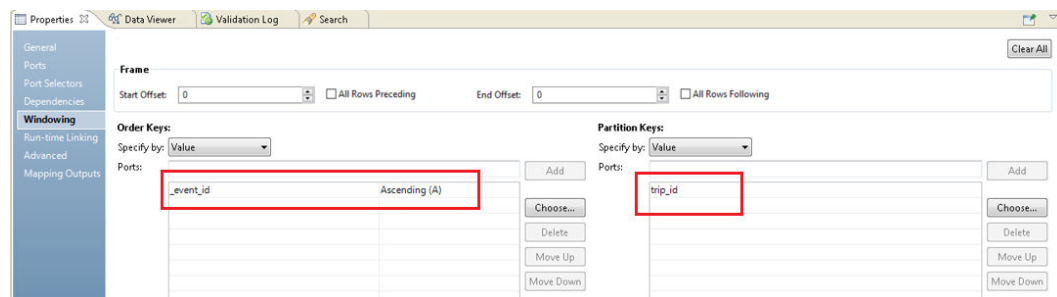
You order the events chronologically and partition the events by trip. You define a window function that accesses the event time from the previous row, and you use an `ADD_TO_DATE` function to calculate the time difference between the two events.

Windowing Properties

You define the following windowing properties on the Windowing tab:

Property	Description
Order key	<code>_event_id</code> Ascending. Arranges the data chronologically by ascending event ID.
Partition key	<code>trip_id</code> . Groups the rows according to trip ID so calculations are based on events from the same trip.
Frame	Not specified. Window functions access rows based on the offset argument and ignore the frame specification.

The following image shows the windowing properties you configure in the Expression transformation:



Window Function

You define the following LAG function to get the event time from the previous row:

```
LAG ( _event_time, 1, NULL )
```

Where:

- `_event_time` is the column name whose value from the previous row is to be returned.
- `1` is the offset. This value represents the row immediately before the current row.
- `NULL` is the default value. The function returns NULL if the return value is outside the bounds of the partition.

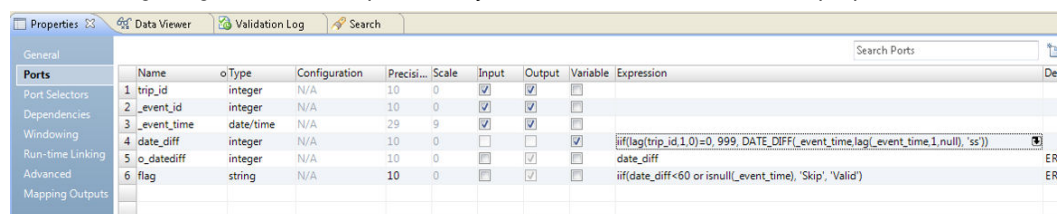
You define the following DATE_DIFF function to calculate the length of time between the two dates:

```
DATE_DIFF ( _event_time, LAG ( _event_time, 1, NULL ), 'ss' )
```

You flag the row as skipped if the DATE_DIFF is less than 60 seconds, or if the `_event_time` is NULL:

```
IF ( DATE_DIFF < 60 or ISNULL ( _event_time ), 'Skip', 'Valid' )
```

The following image shows the expressions you define in the transformation properties:



Output

The transformation produces the following outputs:

Trip ID	Event ID	Event Time	Time Difference	Flag
101	1	2017-05-03 12:00:00	NULL*	Skip
101	2	2017-05-03 12:00:34	34	Skip
101	3	2017-05-03 12:02:00	86	Valid
101	4	2017-05-03 12:02:23	23	Skip
102	1	2017-05-03 12:00:00	NULL*	Skip
102	2	2017-05-03 12:01:56	116	Valid
102	3	2017-05-03 12:02:00	4	Skip
102	4	2017-05-03 13:00:00	3480	Valid
103	1	2017-05-03 12:00:00	NULL*	Skip
103	2	2017-05-03 12:00:12	12	Skip
103	3	2017-05-03 12:01:12	60	Valid

*The rows preceding these rows are outside the bounds of the partition, so the transformation produced NULL values.

Aggregate Function as Window Function Example

You work for a human resources group and you want to compare each of your employees' salaries with the average salary in his or her department:

The following table lists the department names, the employee identification number, and the employee's salary:

Department	Employee	Salary
Development	11	5200
Development	7	4200
Development	9	4500
Development	8	6000
Development	10	5200
Personnel	5	3500
Personnel	2	3900
Sales	3	4800

Department	Employee	Salary
Sales	1	5000
Sales	4	4800

You set an unbounded frame to include all employees in the calculation, and you define an aggregate function to calculate the difference between each employee's salary and the average salary in the department.

Windowing Properties

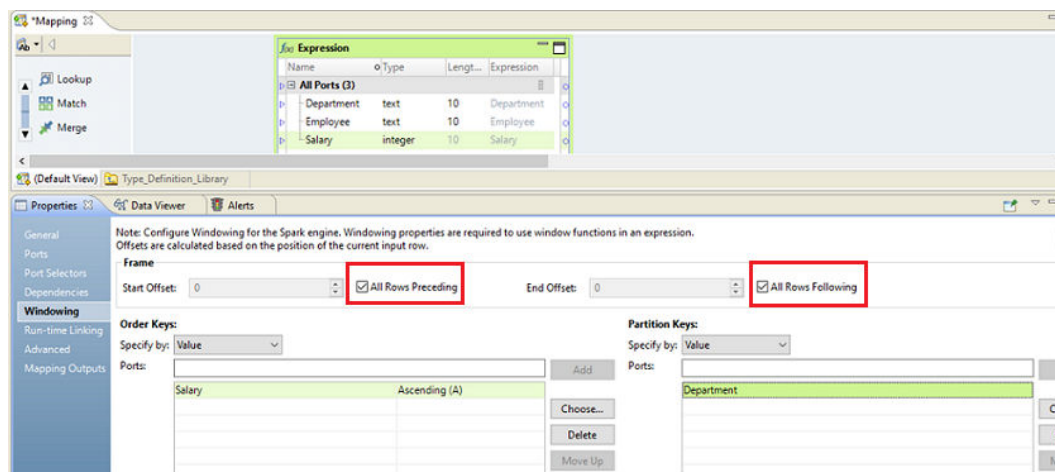
You define the following windowing properties on the Windowing tab:

Property	Description
Order key	Salary Ascending. Arranges the data by increasing salary.
Partition key	Department. Groups the rows according to department.
Start offset	All Rows Preceding
End offset	All Rows Following

With an unbounded frame, the aggregate function includes all partition rows in the calculation.

For example, suppose the current row is the third row. The third row is in the "Development" partition, so the frame includes the third row in addition to all rows before and after the third row in the "Development" partition.

The following image shows the windowing properties you configure in the Expression transformation:



Window Function

An aggregate function acts as a window function when you configure the transformation for windowing.

You define the following aggregate function to calculate the difference between each employee's salary and the average salary in his or her department:

```
Salary - AVG ( Salary ) = Salary_Diff
```

Output

The transformation produces the following salary differences:

Department	Employee	Salary	Salary_Diff
Development	11	5200	-820
Development	7	4200	-520
Development	9	4500	180
Development	8	6000	180
Development	10	5200	980
Personnel	5	3500	200
Personnel	2	3900	200
Sales	3	4800	-66
Sales	1	5000	-66
Sales	4	4800	134

You can identify which employees are making less or more than the average salary for his or her department. Based on this information, you can add other transformations to learn more about your data. For example, you might add a Rank transformation to produce a numerical rank for each employee within his or her department.

APPENDIX A

Connections

This appendix includes the following topics:

- [Connections, 234](#)
- [Cloud Provisioning Configuration, 234](#)
- [Amazon Redshift Connection Properties, 240](#)
- [Amazon S3 Connection Properties, 241](#)
- [Cassandra Connection Properties, 243](#)
- [Databricks Connection Properties, 244](#)
- [Google Analytics Connection Properties, 246](#)
- [Google BigQuery Connection Properties, 246](#)
- [Google Cloud Spanner Connection Properties, 247](#)
- [Google Cloud Storage Connection Properties, 248](#)
- [Hadoop Connection Properties, 249](#)
- [HDFS Connection Properties, 254](#)
- [HBase Connection Properties, 256](#)
- [HBase Connection Properties for MapR-DB, 256](#)
- [Hive Connection Properties, 257](#)
- [JDBC Connection Properties, 260](#)
- [Kafka Connection Properties, 267](#)
- [Microsoft Azure Blob Storage Connection Properties, 267](#)
- [Microsoft Azure Cosmos DB SQL API Connection Properties, 269](#)
- [Microsoft Azure Data Lake Store Connection Properties, 269](#)
- [Microsoft Azure SQL Data Warehouse Connection Properties, 270](#)
- [Snowflake Connection Properties, 272](#)
- [Creating a Connection to Access Sources or Targets, 273](#)
- [Creating a Hadoop Connection, 273](#)
- [Configuring Hadoop Connection Properties, 274](#)

Connections

Create a connection to access non-native environments, Hadoop and Databricks. If you access HBase, HDFS, or Hive sources or targets in the Hadoop environment, you must also create those connections. You can create the connections using the Developer tool, Administrator tool, and infacmd.

You can create the following types of connections:

Hadoop connection

Create a Hadoop connection to run mappings in the Hadoop environment.

HBase connection

Create an HBase connection to access HBase. The HBase connection is a NoSQL connection.

HDFS connection

Create an HDFS connection to read data from or write data to the HDFS file system on a Hadoop cluster.

Hive connection

Create a Hive connection to access Hive as a source or target. You can access Hive as a source if the mapping is enabled for the native or Hadoop environment. You can access Hive as a target if the mapping runs on the Blaze engine.

JDBC connection

Create a JDBC connection and configure Sqoop properties in the connection to import and export relational data through Sqoop.

Databricks connection

Create a Databricks connection to run mappings in the Databricks environment.

Note: For information about creating connections to other sources or targets such as social media web sites or Teradata, see the respective PowerExchange adapter user guide for information.

Cloud Provisioning Configuration

The cloud provisioning configuration establishes a relationship between the Create Cluster task and the cluster connection that the workflows use to run mapping tasks. The Create Cluster task must include a reference to the cloud provisioning configuration. In turn, the cloud provisioning configuration points to the cluster connection that you create for use by the cluster workflow.

The properties to populate depend on the Hadoop distribution you choose to build a cluster on. Choose one of the following connection types:

- AWS Cloud Provisioning. Connects to an Amazon EMR cluster on Amazon Web Services.
- Azure Cloud Provisioning. Connects to an HDInsight cluster on the Azure platform.
- Databricks Cloud Provisioning. Connects to a Databricks cluster on the Azure Databricks platform.

AWS Cloud Provisioning Configuration Properties

The properties in the AWS cloud provisioning configuration enable the Data Integration Service to contact and create resources on the AWS cloud platform.

General Properties

The following table describes cloud provisioning configuration general properties:

Property	Description
Name	Name of the cloud provisioning configuration.
ID	ID of the cloud provisioning configuration. Default: Same as the cloud provisioning configuration name.
Description.	Optional. Description of the cloud provisioning configuration.
AWS Access Key ID	Optional. ID of the AWS access key, which AWS uses to control REST or HTTP query protocol requests to AWS service APIs. If you do not specify a value, Informatica attempts to follow the Default Credential Provider Chain.
AWS Secret Access Key	Secret component of the AWS access key. Required if you specify the AWS Access Key ID.
Region	Region in which to create the cluster. This must be the region in which the VPC is running. Use AWS region values. For a list of acceptable values, see AWS documentation. Note: The region where you want to create the cluster can be different from the region in which the Informatica domain is installed.

Permissions

The following table describes cloud provisioning configuration permissions properties:

Property	Description
EMR Role	Name of the service role for the EMR cluster that you create. The role must have sufficient permissions to create a cluster, access S3 resources, and run jobs on the cluster. When the AWS administrator creates this role, they select the "EMR" role. This contains the default AmazonElasticMapReduceRole policy. You can edit the services in this policy.
EC2 Instance Profile	Name of the EC2 instance profile role that controls permissions on processes that run on the cluster. When the AWS administrator creates this role, they select the "EMR Role for EC2" role. This includes S3 access by default.
Auto Scaling Role	Required if you configure auto-scaling for the EMR cluster. This role is created when the AWS administrator configures auto-scaling on any cluster in the VPC. Default: When you leave this field blank, it is equivalent to setting the Auto Scaling role to "Proceed without role" when the AWS administrator creates a cluster in the AWS console.

EC2 Configuration

The following table describes cloud provisioning configuration EC2 configuration properties:

Property	Description
EC2 Key Pair	EC2 key pair to enable communication with the EMR cluster master node. Optional. This credential enables you to log into the cluster. Configure this property if you intend the cluster to be non-ephemeral.
EC2 Subnet	ID of the subnet on the VPC in which to create the cluster. Use the subnet ID of the EC2 instance where the cluster runs.
Master Security Group	Optional. ID of the security group for the cluster master node. Acts as a virtual firewall to control inbound and outbound traffic to cluster nodes. Security groups are created when the AWS administrator creates and configures a cluster in a VPC. In the AWS console, the property is equivalent to ElasticMapReduce-master. You can use existing security groups, or the AWS administrator might create dedicated security groups for the ephemeral cluster. If you do not specify a value, the cluster applies the default security group for the VPC.
Additional Master Security Groups	Optional. IDs of additional security groups to attach to the cluster master node. Use a comma-separated list of security group IDs.
Core and Task Security Group	Optional. ID of the security group for the cluster core and task nodes. When the AWS administrator creates and configures a cluster In the AWS console, the property is equivalent to the ElasticMapReduce-slave security group If you do not specify a value, the cluster applies the default security group for the VPC.
Additional Core and Task Security Groups	Optional. IDs of additional security groups to attach to cluster core and task nodes. Use a comma-separated list of security group IDs.
Service Access Security Group	EMR managed security group for service access. Required when you provision an EMR cluster in a private subnet.

Azure Cloud Provisioning Configuration Properties

The properties in the Azure cloud provisioning configuration enable the Data Integration Service to contact and create resources on the Azure cloud platform.

Authentication Details

The following table describes authentication properties to configure:

Property	Description
Name	Name of the cloud provisioning configuration.
ID	ID of the cloud provisioning configuration. Default: Same as the cloud provisioning configuration name.
Description	Optional. Description of the cloud provisioning configuration.

Property	Description
Subscription ID	ID of the Azure account to use in the cluster creation process.
Tenant ID	A GUID string associated with the Azure Active Directory.
Client ID	A GUID string that is the same as the Application ID associated with the Service Principal. The Service Principal must be assigned to a role that has permission to create resources in the subscription that you identified in the Subscription ID property.
Client Secret	An octet string that provides a key associated with the client ID.

Storage Account Details

Choose to configure access to one of the following storage types:

- Azure Data Lake Storage (ADLS). See [Azure documentation](#).
- An Azure Storage Account, known as general or blob storage. See [Azure documentation](#).

The following table describes the information you need to configure Azure Data Lake Storage (ADLS) with the HDInsight cluster:

Property	Description
Azure Data Lake Store Name	Name of the ADLS storage to access. The ADLS storage and the cluster to create must reside in the same region.
Data Lake Service Principal Client ID	A credential that enables programmatic access to ADLS storage. Enables the Informatica domain to communicate with ADLS and run commands and mappings on the HDInsight cluster. The service principal is an Azure user that meets the following requirements: - Permissions to access required directories in ADLS storage. - Certificate-based authentication for ADLS storage. - Key-based authentication for ADLS storage.
Data Lake Service Principal Certificate Contents	The Base64 encoded text of the public certificate used with the service principal. Leave this property blank when you create the cloud provisioning configuration. After you save the cloud provisioning configuration, log in to the VM where the Informatica domain is installed and run <code>infacmd ccps updateADLSCertificate</code> to populate this property.
Data Lake Service Principal Certificate Password	Private key for the service principal. This private key must be associated with the service principal certificate.
Data Lake Service Principal Client Secret	An octet string that provides a key associated with the service principal.
Data Lake Service Principal OAUTH Token Endpoint	Endpoint for OAUTH token based authentication.

The following table describes the information you need to configure Azure General Storage, also known as blob storage, with the HDInsight cluster:

Property	Description
Azure Storage Account Name	Name of the storage account to access. Get the value from the Storage Accounts node in the Azure web console. The storage and the cluster to create must reside in the same region.
Azure Storage Account Key	A key to authenticate access to the storage account. To get the value from the Azure web console, select the storage account, then Access Keys. The console displays the account keys.

Cluster Deployment Details

The following table describes the cluster deployment properties that you configure:

Property	Description
Resource Group	Resource group in which to create the cluster. A resource group is a logical set of Azure resources.
Virtual Network Resource Group	Optional. Resource group to which the virtual network belongs. If you do not specify a resource group, the Data Integration Service assumes that the virtual network is a member of the same resource group as the cluster.
Virtual Network	Name of the virtual network or vnet where you want to create the cluster. Specify a vnet that resides in the resource group that you specified in the Virtual Network Resource Group property. The vnet must be in the same region as the region in which to create the cluster.
Subnet Name	Subnet in which to create the cluster. The subnet must be a part of the vnet that you designated in the previous property. Each vnet can have one or more subnets. The Azure administrator can choose an existing subnet or create one for the cluster.

External Hive Metastore Details

You can specify the properties to enable the cluster to connect to a Hive metastore database that is external to the cluster.

You can use an external relational database like MySQL or Amazon RDS as the Hive metastore database. The external database must be on the same cloud platform as the cluster to create.

If you do not specify an existing external database in this dialog box, the cluster creates its own database on the cluster. This database is terminated when the cluster is terminated.

The following table describes the Hive metastore database properties that you configure:

Property	Description
Database Name	Name of the Hive metastore database.
Database Server Name	Server on which the database resides. Note: The database server name on the Azure web console commonly includes the suffix <code>database.windows.net</code> . For example: <code>server123xyz.database.windows.net</code> . You can specify the database server name without the suffix and Informatica will automatically append the suffix. For example, you can specify <code>server123xyz</code> .
Database User Name	User name of the account for the domain to use to access the database.
Database Password	Password for the user account.

Databricks Cloud Provisioning Configuration Properties

The properties in the Databricks cloud provisioning configuration enable the Data Integration Service to contact and create resources on the Databricks cloud platform.

The following table describes the Databricks cloud provisioning configuration properties:

Property	Description
Name	Name of the cloud provisioning configuration.
ID	The cluster ID of the Databricks cluster.
Description	Optional description of the cloud provisioning configuration.
Databricks domain	Domain name of the Databricks deployment.
Databricks token ID	The token ID created within Databricks required for authentication. Note: If the token has an expiration date, verify that you get a new token from the Databricks administrator before it expires.

Amazon Redshift Connection Properties

When you set up an Amazon Redshift connection, you must configure the connection properties.

The following table describes the Amazon Redshift connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters:~`!\$%^&*()-+={} \:;'"<, > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Amazon Redshift in the Database.

The **Details** tab contains the connection attributes of the Amazon Redshift connection. The following table describes the connection attributes:

Property	Description
Username	User name of the Amazon Redshift account.
Password	Password for the Amazon Redshift account.
Schema	Optional. Amazon Redshift schema name. Do not specify the schema name if you want to use multiple schema. The Data Object wizard displays all the user-defined schemas available for the Amazon Redshift objects. Default is public.
AWS Access Key ID	Amazon S3 bucket access key ID. Note: Required if you do not use AWS Identity and Access Management (IAM) authentication.
AWS Secret Access Key	Amazon S3 bucket secret access key ID. Note: Required if you do not use AWS Identity and Access Management (IAM) authentication.
Master Symmetric Key	Optional. Provide a 256-bit AES encryption key in the Base64 format when you enable client-side encryption. You can generate a key using a third-party tool. If you specify a value, ensure that you specify the encryption type as client side encryption in the advanced target properties.

Property	Description
Customer Master Key ID	Optional. Specify the customer master key ID or alias name generated by AWS Key Management Service (AWS KMS). You must generate the customer master key corresponding to the region where Amazon S3 bucket resides. You can specify any of the following values: Customer generated customer master key Enables client-side or server-side encryption. Default customer master key Enables client-side or server-side encryption. Only the administrator user of the account can use the default customer master key ID to enable client-side encryption. Note: You can use customer master key ID when you run a mapping in the native environment or on the Spark engine.
Cluster Node Type	Node type of the Amazon Redshift cluster. You can select the following options: - ds1.xlarge - ds1.8xlarge - dc1.large - dc1.8xlarge - ds2.xlarge - ds2.8xlarge For more information about nodes in the cluster, see the Amazon Redshift documentation.
Number of Nodes in Cluster	Number of nodes in the Amazon Redshift cluster. For more information about nodes in the cluster, see the Amazon Redshift documentation.
JDBC URL	Amazon Redshift connection URL.

Note: If you upgrade the mappings created in versions 10.1.1 Update 2 or earlier, you must select the relevant schema in the connection property. Else, the mappings fail when you run them on current version.

Amazon S3 Connection Properties

When you set up an Amazon S3 connection, you must configure the connection properties.

The following table describes the Amazon S3 connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] \ ; , " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.

Property	Description
Location	The domain where you want to create the connection.
Type	The Amazon S3 connection type.
Access Key	The access key ID for access to Amazon account resources. Note: Required if you do not use AWS Identity and Access Management (IAM) authentication.
Secret Key	The secret access key for access to Amazon account resources. The secret key is associated with the access key and uniquely identifies the account. Note: Required if you do not use AWS Identity and Access Management (IAM) authentication.
Folder Path	The complete path to Amazon S3 objects. The path must include the bucket name and any folder name. Do not use a slash at the end of the folder path. For example, <bucket name>/<my folder name>.
Master Symmetric Key	Optional. Provide a 256-bit AES encryption key in the Base64 format when you enable client-side encryption. You can generate a master symmetric key using a third-party tool.
Customer Master Key ID	Optional. Specify the customer master key ID or alias name generated by AWS Key Management Service (AWS KMS). You must generate the customer master key for the same region where Amazon S3 bucket reside. You can specify any of the following values: Customer generated customer master key Enables client-side or server-side encryption. Default customer master key Enables client-side or server-side encryption. Only the administrator user of the account can use the default customer master key ID to enable client-side encryption. Note: Applicable when you run a mapping in the native environment or on the Spark engine.
Region Name	Select the AWS region in which the bucket you want to access resides. Select one of the following regions: - Asia Pacific (Mumbai) - Asia Pacific (Seoul) - Asia Pacific (Singapore) - Asia Pacific (Sydney) - Asia Pacific (Tokyo) - AWS GovCloud (US) - Canada (Central) - China (Beijing) - China (Ningxia) - EU (Ireland) - EU (Frankfurt) - EU (London) - EU (Paris) - South America (Sao Paulo) - US East (Ohio) - US East (N. Virginia) - US West (N. California) - US West (Oregon) Default is US East (N. Virginia).

Cassandra Connection Properties

When you set up a Cassandra connection, you must configure the connection properties.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes the Cassandra connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. The ID must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Cassandra .
Host Name	Host name or IP address of the Cassandra server.
Port	Cassandra server port number. Default is 9042.
User Name	User name to access the Cassandra server.
Password	Password corresponding to the user name to access the Cassandra server.
Default Keyspace	Name of the Cassandra keyspace to use by default.
SQL Identifier Character	Type of character that the database uses to enclose delimited identifiers in SQL or CQL queries. The available characters depend on the database type. Select None if the database uses regular identifiers. When the Data Integration Service generates SQL or CQL queries, the service does not place delimited characters around any identifiers. Select a character if the database uses delimited identifiers. When the Data Integration Service generates SQL or CQL queries, the service encloses delimited identifiers within this character.
Additional Connection Properties	Enter one or more JDBC connection parameters in the following format: <param1>=<value>;<param2>=<value>;<param3>=<value> PowerExchange for Cassandra JDBC supports the following JDBC connection parameters: - BinaryColumnLength - DecimalColumnScale - EnableCaseSensitive - EnableNullInsert - EnablePaging - RowsPerPage - StringColumnLength - VTTableNameSeparator

Property	Description
SSL Mode	Not applicable for PowerExchange for Cassandra JDBC. Select disabled .
SSL Truststore Path	Not applicable for PowerExchange for Cassandra JDBC.
SSL Truststore Password	Not applicable for PowerExchange for Cassandra JDBC.
SSL Keystore Path	Not applicable for PowerExchange for Cassandra JDBC.
SSL Keystore Password	Not applicable for PowerExchange for Cassandra JDBC.

Databricks Connection Properties

Use the Databricks connection to run mappings on a Databricks cluster.

A Databricks connection is a cluster type connection. You can create and manage a Databricks connection in the Administrator tool or the Developer tool. You can use infacmd to create a Databricks connection. Configure properties in the Databricks connection to enable communication between the Data Integration Service and the Databricks cluster.

The following table describes the general connection properties for the Databricks connection:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.
Connection Type	Choose Databricks.
Cluster Configuration	Name of the cluster configuration associated with the Databricks environment. Required if you do not configure the cloud provisioning configuration.
Cloud Provisioning Configuration	Name of the cloud provisioning configuration associated with a Databricks cloud platform. Required if you do not configure the cluster configuration.

Property	Description
Staging Directory	The directory where the Databricks Spark engine stages run-time files. If you specify a directory that does not exist, the Data Integration Service creates it at run time. If you do not provide a directory path, the run-time staging files are written to <i>/<cluster staging directory>/DATABRICKS</i> .
Advanced Properties	List of advanced properties that are unique to the Databricks environment. You can configure run-time properties for the Databricks environment in the Data Integration Service and in the Databricks connection. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Databricks connection. The Data Integration Service processes property overrides based on the following priorities: 1. Databricks connection advanced properties 2. Data Integration Service custom properties Note: Informatica does not recommend changing these property values before you consult with third-party documentation, Informatica documentation, or Informatica Global Customer Support. If you change a value without knowledge of the property, you might experience performance degradation or other unexpected results.

Advanced Properties

Configure the following properties in the **Advanced Properties** of the Databricks configuration section:

infaspark.json.parser.mode

Specifies the parser how to handle corrupt JSON records. You can set the value to one of the following modes:

- **DROPMALFORMED.** The parser ignores all corrupted records. Default mode.
- **PERMISSIVE.** The parser accepts non-standard fields as nulls in corrupted records.
- **FAILFAST.** The parser generates an exception when it encounters a corrupted record and the Spark application goes down.

infaspark.json.parser.multiLine

Specifies whether the parser can read a multiline record in a JSON file. You can set the value to true or false. Default is false. Applies only to non-native distributions that use Spark version 2.2.x and above.

infaspark.flatfile.writer.nullValue

When the Databricks Spark engine writes to a target, it converts null values to empty strings (" "). For example, 12, AB,"",23p09udj.

The Databricks Spark engine can write the empty strings to string columns, but when it tries to write an empty string to a non-string column, the mapping fails with a type mismatch.

To allow the Databricks Spark engine to convert the empty strings back to null values and write to the target, configure the following advanced property in the Databricks Spark connection:

```
infaspark.flatfile.writer.nullValue=true
```

Google Analytics Connection Properties

When you set up a Google Analytics connection, you must configure the connection properties.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes the Google Analytics connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters:~`!\$%^&*()-+={}[]\;:'<,>.?/
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. The ID must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Google Analytics .
Service Account ID	Specifies the client_email value present in the JSON file that you download after you create a service account.
Service Account Key	Specifies the private_key value present in the JSON file that you download after you create a service account.
APIVersion	API that PowerExchange for Google Analytics uses to read from Google Analytics reports. Select Core Reporting API v3 . Note: PowerExchange for Google Analytics does not support Analytics Reporting API v4.

Google BigQuery Connection Properties

When you set up a Google BigQuery connection, you must configure the connection properties.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes the Google BigQuery connection properties:

Property	Description
Service Account ID	Specifies the client_email value present in the JSON file that you download after you create a service account in Google BigQuery.
Service Account Key	Specifies the private_key value present in the JSON file that you download after you create a service account in Google BigQuery.

Property	Description
Connection mode	The mode that you want to use to read data from or write data to Google BigQuery. Select one of the following connection modes: - Simple. Flattens each field within the Record data type field as a separate field in the mapping. - Hybrid. Displays all the top-level fields in the Google BigQuery table including Record data type fields. PowerExchange for Google BigQuery displays the top-level Record data type field as a single field of the String data type in the mapping. - Complex. Displays all the columns in the Google BigQuery table as a single field of the String data type in the mapping. Default is Simple.
Schema Definition File Path	Specifies a directory on the client machine where the Data Integration Service must create a JSON file with the sample schema of the Google BigQuery table. The JSON file name is the same as the Google BigQuery table name. Alternatively, you can specify a storage path in Google Cloud Storage where the Data Integration Service must create a JSON file with the sample schema of the Google BigQuery table. You can download the JSON file from the specified storage path in Google Cloud Storage to a local machine.
Project ID	Specifies the project_id value present in the JSON file that you download after you create a service account in Google BigQuery. If you have created multiple projects with the same service account, enter the ID of the project that contains the dataset that you want to connect to.
Storage Path	This property applies when you read or write large volumes of data. Path in Google Cloud Storage where the Data Integration Service creates a local stage file to store the data temporarily. You can either enter the bucket name or the bucket name and folder name. For example, enter <code>gs://<bucket_name></code> or <code>gs://<bucket_name>/<folder_name></code>

Google Cloud Spanner Connection Properties

When you set up a Google Cloud Spanner connection, you must configure the connection properties.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes the Google Cloud Spanner connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. The ID must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.

Property	Description
Location	The domain where you want to create the connection.
Type	The connection type. Select Google Cloud Spanner.
Project ID	Specifies the project_id value present in the JSON file that you download after you create a service account. If you have created multiple projects with the same service account, enter the ID of the project that contains the bucket that you want to connect to.
Service Account ID	Specifies the client_email value present in the JSON file that you download after you create a service account.
Service Account Key	Specifies the private_key value present in the JSON file that you download after you create a service account.
Instance ID	Name of the instance that you created in Google Cloud Spanner.

Google Cloud Storage Connection Properties

When you set up a Google Cloud Storage connection, you must configure the connection properties.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes the Google Cloud Storage connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { }] [\ ; : " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. The ID must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Google Cloud Storage .
Project ID	Specifies the project_id value present in the JSON file that you download after you create a service account. If you have created multiple projects with the same service account, enter the ID of the project that contains the bucket that you want to connect to.

Property	Description
Service Account ID	Specifies the client_email value present in the JSON file that you download after you create a service account.
Service Account Key	Specifies the private_key value present in the JSON file that you download after you create a service account.

Hadoop Connection Properties

Use the Hadoop connection to configure mappings to run on a Hadoop cluster. A Hadoop connection is a cluster type connection. You can create and manage a Hadoop connection in the Administrator tool or the Developer tool. You can use infacmd to create a Hadoop connection. Hadoop connection properties are case sensitive unless otherwise noted.

Hadoop Cluster Properties

Configure properties in the Hadoop connection to enable communication between the Data Integration Service and the Hadoop cluster.

The following table describes the general connection properties for the Hadoop connection:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [}] \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. Enter a string that you can use to identify the connection. The description cannot exceed 4,000 characters.
Cluster Configuration	The name of the cluster configuration associated with the Hadoop environment. Required if you do not configure the Cloud Provisioning Configuration.
Cloud Provisioning Configuration	Name of the cloud provisioning configuration associated with a cloud platform such as Amazon AWS or Microsoft Azure. Required if you do not configure the Cluster Configuration.

Property	Description
Cluster Environment Variables*	Environment variables that the Hadoop cluster uses. For example, the variable ORACLE_HOME represents the directory where the Oracle database client software is installed. You can configure run-time properties for the Hadoop environment in the Data Integration Service, the Hadoop connection, and in the mapping. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Hadoop connection or in the mapping. The Data Integration Service processes property overrides based on the following priorities: 1. Mapping custom properties set using infacmd ms runMapping with the -cp option 2. Mapping run-time properties for the Hadoop environment 3. Hadoop connection advanced properties for run-time engines 4. Hadoop connection advanced general properties, environment variables, and classpaths 5. Data Integration Service custom properties
Cluster Library Path*	The path for shared libraries on the cluster. The \$DEFAULT_CLUSTER_LIBRARY_PATH variable contains a list of default directories.
Cluster Classpath*	The classpath to access the Hadoop jar files and the required libraries. The \$DEFAULT_CLUSTER_CLASSPATH variable contains a list of paths to the default jar files and libraries. You can configure run-time properties for the Hadoop environment in the Data Integration Service, the Hadoop connection, and in the mapping. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Hadoop connection or in the mapping. The Data Integration Service processes property overrides based on the following priorities: 1. Mapping custom properties set using infacmd ms runMapping with the -cp option 2. Mapping run-time properties for the Hadoop environment 3. Hadoop connection advanced properties for run-time engines 4. Hadoop connection advanced general properties, environment variables, and classpaths 5. Data Integration Service custom properties
Cluster Executable Path*	The path for executable files on the cluster. The \$DEFAULT_CLUSTER_EXEC_PATH variable contains a list of paths to the default executable files.
* Informatica does not recommend changing these property values before you consult with third-party documentation, Informatica documentation, or Informatica Global Customer Support. If you change a value without knowledge of the property, you might experience performance degradation or other unexpected results.	

Common Properties

The following table describes the common connection properties that you configure for the Hadoop connection:

Property	Description
Impersonation User Name	Required if the Hadoop cluster uses Kerberos authentication. Hadoop impersonation user. The user name that the Data Integration Service impersonates to run mappings in the Hadoop environment. The Data Integration Service runs mappings based on the user that is configured. Refer the following order to determine which user the Data Integration Services uses to run mappings: 1. Operating system profile user. The mapping runs with the operating system profile user if the profile user is configured. If there is no operating system profile user, the mapping runs with the Hadoop impersonation user. 2. Hadoop impersonation user. The mapping runs with the Hadoop impersonation user if the operating system profile user is not configured. If the Hadoop impersonation user is not configured, the Data Integration Service runs mappings with the Data Integration Service user. 3. Informatica services user. The mapping runs with the operating user that starts the Informatica daemon if the operating system profile user and the Hadoop impersonation user are not configured.
Temporary Table Compression Codec	Hadoop compression library for a compression codec class name. Note: The Spark engine does not support compression settings for temporary tables. When you run mappings on the Spark engine, the Spark engine stores temporary tables in an uncompressed file format.
Codec Class Name	Codec class name that enables data compression and improves performance on temporary staging tables.
Hive Staging Database Name	Namespace for Hive staging tables. Use the name <code>default</code> for tables that do not have a specified database name. If you do not configure a namespace, the Data Integration Service uses the Hive database name in the Hive target connection to create staging tables. When you run a mapping in the native environment to write data to Hive, you must configure the Hive staging database name in the Hive connection. The Data Integration Service ignores the value you configure in the Hadoop connection.
Advanced Properties	List of advanced properties that are unique to the Hadoop environment. The properties are common to the Blaze and Spark engines. The advanced properties include a list of default properties. You can configure run-time properties for the Hadoop environment in the Data Integration Service, the Hadoop connection, and in the mapping. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Hadoop connection or in the mapping. The Data Integration Service processes property overrides based on the following priorities: 1. Mapping custom properties set using <code>infacmd ms runMapping</code> with the <code>-cp</code> option 2. Mapping run-time properties for the Hadoop environment 3. Hadoop connection advanced properties for run-time engines 4. Hadoop connection advanced general properties, environment variables, and classpaths 5. Data Integration Service custom properties Note: Informatica does not recommend changing these property values before you consult with third-party documentation, Informatica documentation, or Informatica Global Customer Support. If you change a value without knowledge of the property, you might experience performance degradation or other unexpected results.

Reject Directory Properties

The following table describes the connection properties that you configure to the Hadoop Reject Directory.

Property	Description
Write Reject Files to Hadoop	If you use the Blaze engine to run mappings, select the check box to specify a location to move reject files. If checked, the Data Integration Service moves the reject files to the HDFS location listed in the property, Reject File Directory. By default, the Data Integration Service stores the reject files based on the RejectDir system parameter.
Reject File Directory	The directory for Hadoop mapping files on HDFS when you run mappings.

Hive Pushdown Configuration

Note: Effective in version 10.2.2, Informatica dropped support for the Hive engine. Do not configure the pushdown properties related to the Hive engine.

Blaze Configuration

The following table describes the connection properties that you configure for the Blaze engine:

Property	Description
Blaze Staging Directory	The HDFS file path of the directory that the Blaze engine uses to store temporary files. Verify that the directory exists. The YARN user, Blaze engine user, and mapping impersonation user must have write permission on this directory. Default is <code>/blaze/workdir</code> . If you clear this property, the staging files are written to the Hadoop staging directory <code>/tmp/blaze_<user name></code> .
Blaze User Name	The owner of the Blaze service and Blaze service logs. When the Hadoop cluster uses Kerberos authentication, the default user is the Data Integration Service SPN user. When the Hadoop cluster does not use Kerberos authentication and the Blaze user is not configured, the default user is the Data Integration Service user.
Minimum Port	The minimum value for the port number range for the Blaze engine. Default is 12300.
Maximum Port	The maximum value for the port number range for the Blaze engine. Default is 12600.
YARN Queue Name	The YARN scheduler queue name used by the Blaze engine that specifies available resources on a cluster.
Blaze Job Monitor Address	The host name and port number for the Blaze Job Monitor. Use the following format: <code><hostname>:<port></code> Where - <code><hostname></code> is the host name or IP address of the Blaze Job Monitor server. - <code><port></code> is the port on which the Blaze Job Monitor listens for remote procedure calls (RPC). For example, enter: <code>myhostname:9080</code>

Property	Description
Blaze YARN Node Label	Node label that determines the node on the Hadoop cluster where the Blaze engine runs. If you do not specify a node label, the Blaze engine runs on the nodes in the default partition. If the Hadoop cluster supports logical operators for node labels, you can specify a list of node labels. To list the node labels, use the operators && (AND), (OR), and ! (NOT).
Advanced Properties	List of advanced properties that are unique to the Blaze engine. The advanced properties include a list of default properties. You can configure run-time properties for the Hadoop environment in the Data Integration Service, the Hadoop connection, and in the mapping. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Hadoop connection or in the mapping. The Data Integration Service processes property overrides based on the following priorities: 1. Mapping custom properties set using <code>infacmd ms runMapping</code> with the <code>-cp</code> option 2. Mapping run-time properties for the Hadoop environment 3. Hadoop connection advanced properties for run-time engines 4. Hadoop connection advanced general properties, environment variables, and classpaths 5. Data Integration Service custom properties Note: Informatica does not recommend changing these property values before you consult with third-party documentation, Informatica documentation, or Informatica Global Customer Support. If you change a value without knowledge of the property, you might experience performance degradation or other unexpected results.

Spark Configuration

The following table describes the connection properties that you configure for the Spark engine:

Property	Description
Spark Staging Directory	The HDFS file path of the directory that the Spark engine uses to store temporary files for running jobs. The YARN user, Data Integration Service user, and mapping impersonation user must have write permission on this directory. If you do not specify a file path, by default, the temporary files are written to the Hadoop staging directory <code>/tmp/SPARK_<user name></code>. When you run Sqoop jobs on the Spark engine, the Data Integration Service creates a Sqoop staging directory within the Spark staging directory to store temporary files: <code><Spark staging directory>/sqoop_staging</code>
Spark Event Log Directory	Optional. The HDFS file path of the directory that the Spark engine uses to log events.

Property	Description
YARN Queue Name	The YARN scheduler queue name used by the Spark engine that specifies available resources on a cluster. The name is case sensitive.
Advanced Properties	List of advanced properties that are unique to the Spark engine. The advanced properties include a list of default properties. You can configure run-time properties for the Hadoop environment in the Data Integration Service, the Hadoop connection, and in the mapping. You can override a property configured at a high level by setting the value at a lower level. For example, if you configure a property in the Data Integration Service custom properties, you can override it in the Hadoop connection or in the mapping. The Data Integration Service processes property overrides based on the following priorities: 1. Mapping custom properties set using <code>infacmd ms runMapping</code> with the <code>-cp</code> option 2. Mapping run-time properties for the Hadoop environment 3. Hadoop connection advanced properties for run-time engines 4. Hadoop connection advanced general properties, environment variables, and classpaths 5. Data Integration Service custom properties Note: Informatica does not recommend changing these property values before you consult with third-party documentation, Informatica documentation, or Informatica Global Customer Support. If you change a value without knowledge of the property, you might experience performance degradation or other unexpected results.

HDFS Connection Properties

Use a Hadoop File System (HDFS) connection to access data in the Hadoop cluster. The HDFS connection is a file system type connection. You can create and manage an HDFS connection in the Administrator tool, Analyst tool, or the Developer tool. HDFS connection properties are case sensitive unless otherwise noted.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes HDFS connection properties:

Property	Description
Name	Name of the connection. The name is not case sensitive and must be unique within the domain. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 765 characters.
Location	The domain where you want to create the connection. Not valid for the Analyst tool.
Type	The connection type. Default is Hadoop File System.

Property	Description
User Name	User name to access HDFS.
NameNode URI	The URI to access the storage system. You can find the value for <code>fs.defaultFS</code> in the <code>core-site.xml</code> configuration set of the cluster configuration. Note: If you create connections when you import the cluster configuration, the NameNode URI property is populated by default, and it is updated each time you refresh the cluster configuration. If you manually set this property or override the value, the refresh operation does not update this property.

Accessing Multiple Storage Types

Use the NameNode URI property in the connection parameters to connect to various storage types. The following table lists the storage type and the NameNode URI format for the storage type:

Storage	NameNode URI Format
HDFS	<code>hdfs://<namenode>:<port></code> where: - <code><namenode></code> is the host name or IP address of the NameNode. - <code><port></code> is the port that the NameNode listens for remote procedure calls (RPC). <code>hdfs://<nameservice></code> in case of NameNode high availability.
MapR-FS	<code>maprfs:///</code>
WASB in HDInsight	<code>wasb://<container_name>@<account_name>.blob.core.windows.net/<path></code> where: - <code><container_name></code> identifies a specific Azure Storage Blob container. Note: <code><container_name></code> is optional. - <code><account_name></code> identifies the Azure Storage Blob object. Example: <code>wasb://infabdmoffering1storage.blob.core.windows.net/infabdmoffering1cluster/mr-history</code>
ADLS in HDInsight	<code>adl://home</code>

When you create a cluster configuration from an Azure HDInsight cluster, the cluster configuration uses either ADLS or WASB as the primary storage. You cannot create a cluster configuration with ADLS or WASB as the secondary storage. You can edit the NameNode URI property in the HDFS connection to connect to a local HDFS location.

HBase Connection Properties

Use an HBase connection to access HBase. The HBase connection is a NoSQL connection. You can create and manage an HBase connection in the Administrator tool or the Developer tool. HBase connection properties are case sensitive unless otherwise noted.

The following table describes HBase connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select HBase.
Database Type	Type of database that you want to connect to. Select HBase to create a connection for an HBase table.

HBase Connection Properties for MapR-DB

Use an HBase connection to connect to a MapR-DB table. The HBase connection is a NoSQL connection. You can create and manage an HBase connection in the Administrator tool or the Developer tool. HBase connection properties are case sensitive unless otherwise noted.

The following table describes the HBase connection properties for MapR-DB:

Property	Description
Name	Name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.

Property	Description
Description	Description of the connection. The description cannot exceed 4,000 characters.
Location	Domain where you want to create the connection.
Type	Connection type. Select HBase .
Database Type	Type of database that you want to connect to. Select MapR-DB to create a connection for a MapR-DB table.
Cluster Configuration	The name of the cluster configuration associated with the Hadoop environment.
MapR-DB Database Path	Database path that contains the MapR-DB table that you want to connect to. Enter a valid MapR cluster path. When you create an HBase data object for MapR-DB, you can browse only tables that exist in the MapR-DB path that you specify in the Database Path field. You cannot access tables that are available in sub-directories in the specified path. For example, if you specify the path as <code>/user/customers/</code> , you can access the tables in the <code>customers</code> directory. However, if the <code>customers</code> directory contains a sub-directory named <code>regions</code> , you cannot access the tables in the following directory: <code>/user/customers/regions</code>

Hive Connection Properties

Use the Hive connection to access Hive data. A Hive connection is a database type connection. You can create and manage a Hive connection in the Administrator tool, Analyst tool, or the Developer tool. Hive connection properties are case sensitive unless otherwise noted.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes Hive connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: <code>~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /</code>
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 4000 characters.

Property	Description
Location	The domain where you want to create the connection. Not valid for the Analyst tool.
Type	The connection type. Select Hive.
LDAP username	LDAP user name of the user that the Data Integration Service impersonates to run mappings on a Hadoop cluster. The user name depends on the JDBC connection string that you specify in the Metadata Connection String or Data Access Connection String for the native environment. If the Hadoop cluster uses Kerberos authentication, the principal name for the JDBC connection string and the user name must be the same. Otherwise, the user name depends on the behavior of the JDBC driver. With Hive JDBC driver, you can specify a user name in many ways and the user name can become a part of the JDBC URL. If the Hadoop cluster does not use Kerberos authentication, the user name depends on the behavior of the JDBC driver. If you do not specify a user name, the Hadoop cluster authenticates jobs based on the following criteria: - The Hadoop cluster does not use Kerberos authentication. It authenticates jobs based on the operating system profile user name of the machine that runs the Data Integration Service. - The Hadoop cluster uses Kerberos authentication. It authenticates jobs based on the SPN of the Data Integration Service. LDAP username will be ignored.
Password	Password for the LDAP username.
Environment SQL	SQL commands to set the Hadoop environment. In native environment type, the Data Integration Service executes the environment SQL each time it creates a connection to a Hive metastore. If you use the Hive connection to run profiles on a Hadoop cluster, the Data Integration Service executes the environment SQL at the beginning of each Hive session. The following rules and guidelines apply to the usage of environment SQL in both connection modes: - Use the environment SQL to specify Hive queries. - Use the environment SQL to set the classpath for Hive user-defined functions and then use environment SQL or PreSQL to specify the Hive user-defined functions. You cannot use PreSQL in the data object properties to specify the classpath. If you use Hive user-defined functions, you must copy the .jar files to the following directory: <pre><Informatica installation directory>/services/shared/hadoop/ <Hadoop distribution name>/extras/hive-auxjars</pre> - You can use environment SQL to define Hadoop or Hive parameters that you want to use in the PreSQL commands or in custom queries. - If you use multiple values for the Environment SQL property, ensure that there is no space between the values.
SQL Identifier Character	The type of character used to identify special characters and reserved SQL keywords, such as WHERE. The Data Integration Service places the selected character around special characters and reserved SQL keywords. The Data Integration Service also uses this character for the Support mixed-case identifiers property.

Properties to Access Hive as Source or Target

The following table describes the connection properties that you configure to access Hive as a source or target:

Property	Description
JDBC Driver Class Name	Name of the Hive JDBC driver class. If you leave this option blank, the Developer tool uses the default Apache Hive JDBC driver shipped with the distribution. If the default Apache Hive JDBC driver does not fit your requirements, you can override the Apache Hive JDBC driver with a third-party Hive JDBC driver by specifying the driver class name.
Metadata Connection String	The JDBC connection URI used to access the metadata from the Hadoop server. You can use PowerExchange for Hive to communicate with a HiveServer service or HiveServer2 service. To connect to HiveServer, specify the connection string in the following format: <pre>jdbc:hive2://<hostname>:<port>/<db></pre> Where - <hostname> is name or IP address of the machine on which HiveServer2 runs. - <port> is the port number on which HiveServer2 listens. - <db> is the database name to which you want to connect. If you do not provide the database name, the Data Integration Service uses the default database details. To connect to HiveServer2, use the connection string format that Apache Hive implements for that specific Hadoop Distribution. For more information about Apache Hive connection string formats, see the Apache Hive documentation. For user impersonation, you must add <code>hive.server2.proxy.user=<xyz></code> to the JDBC connection URI. If you do not configure user impersonation, the current user's credentials are used connect to the HiveServer2. If the Hadoop cluster uses SSL or TLS authentication, you must add <code>ssl=true</code> to the JDBC connection URI. For example: <code>jdbc:hive2://<hostname>:<port>/<db>;ssl=true</code> If you use self-signed certificate for SSL or TLS authentication, ensure that the certificate file is available on the client machine and the Data Integration Service machine. For more information, see the <i>Informatica Big Data Management Integration Guide</i>.
Bypass Hive JDBC Server	JDBC driver mode. Select the check box to use the embedded JDBC driver mode. To use the JDBC embedded mode, perform the following tasks: - Verify that Hive client and Informatica services are installed on the same machine. - Configure the Hive connection properties to run mappings on a Hadoop cluster. If you choose the non-embedded mode, you must configure the Data Access Connection String. Informatica recommends that you use the JDBC embedded mode.
Fine Grained Authorization	When you select the option to observe fine grained authorization in a Hive source, the mapping observes the following: - Row and column level restrictions. Applies to Hadoop clusters where Sentry or Ranger security modes are enabled. - Data masking rules. Applies to masking rules set on columns containing sensitive data by Dynamic Data Masking. If you do not select the option, the Blaze and Spark engines ignore the restrictions and masking rules, and results include restricted or sensitive data.

Property	Description
Data Access Connection String	The connection string to access data from the Hadoop data store. To connect to HiveServer, specify the non-embedded JDBC mode connection string in the following format: <pre>jdbc:hive2://<hostname>:<port>/<db></pre> Where - <hostname> is name or IP address of the machine on which HiveServer2 runs. - <port> is the port number on which HiveServer2 listens. - <db> is the database to which you want to connect. If you do not provide the database name, the Data Integration Service uses the default database details. To connect to HiveServer2, use the connection string format that Apache Hive implements for the specific Hadoop Distribution. For more information about Apache Hive connection string formats, see the Apache Hive documentation. For user impersonation, you must add <code>hive.server2.proxy.user=<xyz></code> to the JDBC connection URI. If you do not configure user impersonation, the current user's credentials are used connect to the HiveServer2. If the Hadoop cluster uses SSL or TLS authentication, you must add <code>ssl=true</code> to the JDBC connection URI. For example: <code>jdbc:hive2://<hostname>:<port>/<db>;ssl=true</code> If you use self-signed certificate for SSL or TLS authentication, ensure that the certificate file is available on the client machine and the Data Integration Service machine. For more information, see the <i>Informatica Big Data Management Integration Guide</i>.
Hive Staging Directory on HDFS	HDFS directory for Hive staging tables. You must grant execute permission to the Hadoop impersonation user and the mapping impersonation users. This option is applicable and required when you write data to a Hive target in the native environment.
Hive Staging Database Name	Namespace for Hive staging tables. Use the name <code>default</code> for tables that do not have a specified database name. This option is applicable when you run a mapping in the native environment to write data to a Hive target. If you run the mapping on the Blaze or Spark engine, you do not need to configure the Hive staging database name in the Hive connection. The Data Integration Service uses the value that you configure in the Hadoop connection.

JDBC Connection Properties

You can use a JDBC connection to access tables in a database. You can create and manage a JDBC connection in the Administrator tool, the Developer tool, or the Analyst tool.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes JDBC connection properties:

Property	Description
Database Type	The database type.
Name	Name of the connection. The name is not case sensitive and must be unique within the domain. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 765 characters.
User Name	The database user name.
Password	The password for the database user name.
JDBC Driver Class Name	Name of the JDBC driver class. The following list provides the driver class name that you can enter for the applicable database type: - DataDirect JDBC driver class name for Oracle: <code>com.informatica.jdbc.oracle.OracleDriver</code> - DataDirect JDBC driver class name for IBM DB2: <code>com.informatica.jdbc.db2.DB2Driver</code> - DataDirect JDBC driver class name for Microsoft SQL Server: <code>com.informatica.jdbc.sqlserver.SQLServerDriver</code> - DataDirect JDBC driver class name for Sybase ASE: <code>com.informatica.jdbc.sybase.SybaseDriver</code> - DataDirect JDBC driver class name for Informix: <code>com.informatica.jdbc.informix.InformixDriver</code> - DataDirect JDBC driver class name for MySQL: <code>com.informatica.jdbc.mysql.MySQLDriver</code> For more information about which driver class to use with specific databases, see the vendor documentation.
Connection String	Connection string to connect to the database. Use the following connection string: <code>jdbc:<subprotocol>:<subname></code> For more information about the connection string to use with specific drivers, see the vendor documentation.
Environment SQL	Optional. Enter SQL commands to set the database environment when you connect to the database. The Data Integration Service executes the connection environment SQL each time it connects to the database. Note: If you enable Sqoop, Sqoop ignores this property.
Transaction SQL	Optional. Enter SQL commands to set the database environment when you connect to the database. The Data Integration Service executes the transaction environment SQL at the beginning of each transaction. Note: If you enable Sqoop, Sqoop ignores this property.

Property	Description
SQL Identifier Character	Type of character that the database uses to enclose delimited identifiers in SQL queries. The available characters depend on the database type. Select (None) if the database uses regular identifiers. When the Data Integration Service generates SQL queries, the service does not place delimited characters around any identifiers. Select a character if the database uses delimited identifiers. When the Data Integration Service generates SQL queries, the service encloses delimited identifiers within this character. Note: If you enable Sqoop, Sqoop ignores this property.
Support Mixed-case Identifiers	Enable if the database uses case-sensitive identifiers. When enabled, the Data Integration Service encloses all identifiers within the character selected for the SQL Identifier Character property. When the SQL Identifier Character property is set to none, the Support Mixed-case Identifiers property is disabled. Note: If you enable Sqoop, Sqoop honors this property when you generate and execute a DDL script to create or replace a target at run time. In all other scenarios, Sqoop ignores this property.
Use Sqoop Connector	Enables Sqoop connectivity for the data object that uses the JDBC connection. The Data Integration Service runs the mapping in the Hadoop run-time environment through Sqoop. You can configure Sqoop connectivity for relational data objects, customized data objects, and logical data objects that are based on a JDBC-compliant database. Select Sqoop v1.x to enable Sqoop connectivity. Default is None.
Sqoop Arguments	Enter the arguments that Sqoop must use to connect to the database. Separate multiple arguments with a space. To run the mapping on the Blaze engine with the Teradata Connector for Hadoop (TDCH) specialized connectors for Sqoop, you must define the TDCH connection factory class in the Sqoop arguments. The connection factory class varies based on the TDCH Sqoop Connector that you want to use. - To use Cloudera Connector Powered by Teradata, configure the following Sqoop argument: - <code>Dsqoop.connection.factories=com.cloudera.connector.teradata.TeradataManagerFactory</code> - To use Hortonworks Connector for Teradata (powered by the Teradata Connector for Hadoop), configure the following Sqoop argument: - <code>Dsqoop.connection.factories=org.apache.sqoop.teradata.TeradataManagerFactory</code> To run the mapping on the Spark engine, you do not need to define the TDCH connection factory class in the Sqoop arguments. The Data Integration Service invokes the Cloudera Connector Powered by Teradata and Hortonworks Connector for Teradata (powered by the Teradata Connector for Hadoop) by default. Note: To run the mapping with a generic JDBC connector instead of the specialized Cloudera or Hortonworks connector, you must define the <code>--driver</code> and <code>--connection-manager</code> Sqoop arguments in the JDBC connection. If you define the <code>--driver</code> and <code>--connection-manager</code> arguments in the Read or Write transformation of the mapping, Sqoop ignores the arguments. If you do not enter Sqoop arguments, the Data Integration Service constructs the Sqoop command based on the JDBC connection properties.

JDBC Connection String

Specify the connection string in the JDBC connection to connect to the database.

Specify the connection string in the following format:

```
jdbc:<subprotocol>:<subname>
```

Use the following sample connection strings that you can enter for the applicable database type:

- DataDirect Oracle JDBC driver:

```
jdbc:informatica:oracle://<host>:<port>;SID=<value>
```

- Oracle JDBC driver:

```
jdbc:oracle:thin:@//<host>:<port>:<SID>
```

- DataDirect IBM DB2 JDBC driver:

```
jdbc:informatica:db2://<host>:<port>;DatabaseName=<value>
```

- IBM DB2 JDBC driver:

```
jdbc:db2://<host>:<port>/<database_name>
```

- DataDirect Microsoft SQL Server JDBC driver:

```
jdbc:informatica:sqlserver://<host>;DatabaseName=<value>
```

- Microsoft SQL Server JDBC driver:

```
jdbc:sqlserver://<host>;DatabaseName=<value>
```

- Netezza JDBC driver:

```
jdbc:netezza://<host>:<port>/<database_name>
```

- Pivotal Greenplum driver:

```
jdbc:pivotal:greenplum://<host>:<port>;/database_name=<value>
```

- Postgres Greenplum driver:

```
jdbc:postgresql://<host>:<port>/<database_name>
```

- Teradata JDBC driver:

```
jdbc:teradata://<host>/database_name=<value>,tmode=<value>,charset=<value>
```

Use the following sample connection strings that you can enter for an SSL-enabled applicable database type:

- DataDirect Oracle JDBC driver:

```
jdbc:informatica:oracle://  
<host_name>:<port>;CatalogOptions=6;ServiceName=<service_name>;trustStorePassword=<tru  
ststore_password>;  
keyStorePassword=<keystore_password>;CryptoProtocolVersion=TLSv1.2;keyStore=<keystore_  
location_of_ewallet.p12_file>;  
trustStore=<truststore_location_of_truststore.p12_file>;HostNameInCertificate=<databas  
e_host_name>;encryptionMethod=SSL;  
ValidateServerCertificate=True;
```

- Oracle JDBC driver:

```
jdbc:oracle:thin:@(DESCRIPTION=(ADDRESS=(PROTOCOL=TCPS)(HOST=<host>)  
(PORT=<port_number>))  
(CONNECT_DATA=(SERVICE_NAME=<service_name>))) "
```

- DataDirect Microsoft SQL Server JDBC driver:

```
jdbc:informatica:sqlserver://  
<host_name>:<port>;databaseName=<database_name>;EncryptionMethod=SSL;  
CryptoProtocolVersion=<TLSv1.2_or_TLSv1.1_or_TLSv1>;ValidateServerCertificate=false;Tr  
ustStore=<truststore_location>;  
TrustStorePassword=<truststore_password>
```

- Microsoft SQL Server JDBC driver:

```
jdbc:sqlserver://  
<host_name>:<port>;databaseName=<database_name>;integratedSecurity=false;encrypt=true;  
trustServerCertificate=true;  
TrustStore=/  
<truststore_location>;TrustStorePassword=<truststore_password>;user=<user_name>;passwo  
rd=<password>
```

For more information about the connection string to use with specific drivers, see the vendor documentation.

Sqoop Connection-Level Arguments

In the JDBC connection, you can define the arguments that Sqoop must use to connect to the database. The Data Integration Service merges the arguments that you specify with the default command that it constructs based on the JDBC connection properties. The arguments that you specify take precedence over the JDBC connection properties.

If you want to use the same driver to import metadata and run the mapping, and do not want to specify any additional Sqoop arguments, select **Sqoop v1.x** from the **Use Sqoop Version** list and leave the **Sqoop Arguments** field empty in the JDBC connection. The Data Integration Service constructs the Sqoop command based on the JDBC connection properties that you specify.

However, if you want to use a different driver for run-time tasks or specify additional run-time Sqoop arguments, select **Sqoop v1.x** from the **Use Sqoop Version** list and specify the arguments in the **Sqoop Arguments** field.

You can configure the following Sqoop arguments in the JDBC connection:

driver

Defines the JDBC driver class that Sqoop must use to connect to the database.

Use the following syntax:

```
--driver <JDBC driver class>
```

For example, use the following syntax depending on the database type that you want to connect to:

- **Aurora:** `--driver com.mysql.jdbc.Driver`
- **Greenplum:** `--driver org.postgresql.Driver`
- **IBM DB2:** `--driver com.ibm.db2.jcc.DB2Driver`
- **IBM DB2 z/OS:** `--driver com.ibm.db2.jcc.DB2Driver`
- **Microsoft SQL Server:** `--driver com.microsoft.sqlserver.jdbc.SQLServerDriver`
- **Netezza:** `--driver org.netezza.Driver`
- **Oracle:** `--driver oracle.jdbc.driver.OracleDriver`
- **Teradata:** `--driver com.teradata.jdbc.TeraDriver`

connect

Defines the JDBC connection string that Sqoop must use to connect to the database. The JDBC connection string must be based on the driver that you define in the driver argument.

Use the following syntax:

```
--connect <JDBC connection string>
```

For example, use the following syntax depending on the database type that you want to connect to:

- **Aurora:** `--connect "jdbc:mysql://<host_name>:<port>/<schema_name>"`
- **Greenplum:** `--connect jdbc:postgresql://<host_name>:<port>/<database_name>`
- **IBM DB2:** `--connect jdbc:db2://<host_name>:<port>/<database_name>`
- **IBM DB2 z/OS:** `--connect jdbc:db2://<host_name>:<port>/<database_name>`
- **Microsoft SQL Server:** `--connect jdbc:sqlserver://<host_name>:<port> or named_instance>;databaseName=<database_name>`
- **Netezza:** `--connect "jdbc:netezza://<database_server_name>:<port>/<database_name>;schema=<schema_name>"`

- **Oracle:** `--connect jdbc:oracle:thin:@<database_host_name>:<database_port>:<database_SID>`
- **Teradata:** `--connect jdbc:teradata://<host_name>/database=<database_name>`

Use the following syntax to connect to an SSL-enabled database:

```
--connect <JDBC connection string>
```

For example, use the following syntax depending on the database type that you want to connect to:

- **Microsoft SQL Server:** `--connect jdbc:sqlserver://<host_name>:<port>;databaseName=<database_name>;integratedSecurity=false;encrypt=true;trustServerCertificate=true;TrustStore=/
<truststore_location>;TrustStorePassword=<truststore_password>;user=<user_name>;password=<password>`
- **Oracle:** `--connect jdbc:oracle:thin:@(DESCRIPTION=(ADDRESS=(PROTOCOL=TCPS)(HOST=<host>)(PORT=<port_number>))(CONNECT_DATA=(SERVICE_NAME=<service_name>)))"`

connection-param-file

Defines the extra JDBC parameters through a property file that Sqoop must use to connect to the database. The contents of this file are parsed as standard Java properties and passed into the driver when you create a connection.

Use the following syntax:

```
--connection-param-file <parameter file name>
```

For example, use the following syntax to use the parameter file when you connect to the Oracle database.

```
--connection-param-file param_file
```

connection-manager

Defines the connection manager class name that Sqoop must use to connect to the database.

Use the following syntax:

```
--connection-manager <connection manager class name>
```

For example, use the following syntax to use the generic JDBC manager class name:

```
--connection-manager org.apache.sqoop.manager.GenericJdbcManager
```

direct

When you read data from or write data to Oracle, you can configure the direct argument to enable Sqoop to use OraOop. OraOop is a specialized Sqoop plug-in for Oracle that uses native protocols to connect to the Oracle database. When you configure OraOop, the performance improves.

You can configure OraOop when you run Sqoop mappings on the Spark engine.

Use the following syntax:

```
--direct
```

When you use OraOop, you must use the following syntax to specify multiple arguments:

```
-D<argument=value> -D<argument=value>
```

Note: If you specify multiple arguments and include a space character between -D and the argument name-value pair, Sqoop considers only the first argument and ignores the remaining arguments.

If you do not direct the job to a specific queue, the Spark engine uses the default queue.

-Dsqoop.connection.factories

To run the mapping on the Blaze engine with the Teradata Connector for Hadoop (TDCH) specialized connectors for Sqoop, you must configure the `-Dsqoop.connection.factories` argument. Use the argument to define the TDCH connection factory class that Sqoop must use. The connection factory class varies based on the TDCH Sqoop Connector that you want to use.

- To use Cloudera Connector Powered by Teradata, configure the `-Dsqoop.connection.factories` argument as follows:
`-Dsqoop.connection.factories=com.cloudera.connector.teradata.TeradataManagerFactory`
- To use Hortonworks Connector for Teradata (powered by the Teradata Connector for Hadoop), configure the `-Dsqoop.connection.factories` argument as follows:
`-Dsqoop.connection.factories=org.apache.sqoop.teradata.TeradataManagerFactory`

Note: To run the mapping on the Spark engine, you do not need to configure the `-Dsqoop.connection.factories` argument. The Data Integration Service invokes Cloudera Connector Powered by Teradata and Hortonworks Connector for Teradata (powered by the Teradata Connector for Hadoop) by default.

--infaoptimize

Use this argument to disable the performance optimization of Sqoop pass-through mappings on the Spark engine.

When you run a Sqoop pass-through mapping on the Spark engine, the Data Integration Service optimizes mapping performance in the following scenarios:

- You read data from a Sqoop source and write data to a Hive target that uses the Text format.
- You read data from a Sqoop source and write data to an HDFS target that uses the Flat, Avro, or Parquet format.

If you want to disable the performance optimization, set the `--infaoptimize` argument to false. For example, if you see data type issues after you run an optimized Sqoop mapping, you can disable the performance optimization.

Use the following syntax:

```
--infaoptimize false
```

For a complete list of the Sqoop arguments that you can configure, see the Sqoop documentation.

Kafka Connection Properties

The Kafka connection is a messaging connection. Use the Kafka connection to access Kafka as a message target. You can create and manage a Kafka connection in the Developer tool or through infacmd.

The following table describes the general connection properties for the Kafka connection:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ : ; " ' < , > . ? /
ID	The string that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. Enter a string that you can use to identify the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type.

The following table describes the Kafka broker properties for the Kafka connection:

Property	Description
Kafka Broker List	Comma-separated list of Kafka brokers which maintains the configuration of the Kafka messaging broker. To specify a Kafka broker, use the following format: <IP Address>:<port>
ZooKeeper Host Port List	Optional. Comma-separated list of Apache ZooKeeper which maintains the configuration of the Kafka messaging broker. To specify the ZooKeeper, use the following format: <IP Address>:<port>
Retry Timeout	Number of seconds the Integration Service attempts to reconnect to the Kafka broker to write data. If the source or target is not available for the time you specify, the mapping execution stops to avoid any data loss.
Kafka Broker Version	Configure the Kafka messaging broker version to 0.10.1.x-2.0.0.

Microsoft Azure Blob Storage Connection Properties

Use a Microsoft Azure SQL Blob Storage connection to access a Microsoft Azure Blob Storage.

Note: The order of the connection properties might vary depending on the tool where you view them.

You can create and manage a Microsoft Azure Blob Storage connection in the Administrator tool or the Developer tool. The following table describes the Microsoft Azure Blob Storage connection properties:

Property	Description
Name	Name of the Microsoft Azure Blob Storage connection.
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Description of the connection.
Location	The domain where you want to create the connection.
Type	Type of connection. Select AzureBlob.

The **Connection Details** tab contains the connection attributes of the Microsoft Azure Blob Storage connection. The following table describes the connection attributes:

Property	Description
Account Name	Name of the Microsoft Azure Storage account.
Account Key	Microsoft Azure Storage access key.
Container Name	The root container or sub-folders with the absolute path.
Endpoint Suffix	Type of Microsoft Azure end-points. You can select any of the following end-points: - <code>core.windows.net</code>: Default - <code>core.usgovcloudapi.net</code>: To select the US government Microsoft Azure end-points - <code>core.chinacloudapi.cn</code>: Not applicable

Microsoft Azure Cosmos DB SQL API Connection Properties

Use a Microsoft Azure Cosmos DB connection to connect to the Cosmos DB database. When you create a Microsoft Azure Cosmos DB connection, you enter information for metadata and data access.

The following table describes the Microsoft Azure Cosmos DB connection properties:

Property	Description
Name	Name of the Cosmos DB connection.
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Description of the connection. The description cannot exceed 765 characters.
Location	The project or folder in the Model repository where you want to store the Cosmos DB connection.
Type	Select Microsoft Azure Cosmos DB SQL API.
Cosmos DB URI	The URI of Microsoft Azure Cosmos DB account.
Key	The primary and secondary key to which provides you complete administrative access to the resources within Microsoft Azure Cosmos DB account.
Database	Name of the database that contains the collections from which you want to read or write JSON documents.

Note: You can find the Cosmos DB URI and Key values in the **Keys** settings on Azure portal. Contact your Azure administrator for more details.

Microsoft Azure Data Lake Store Connection Properties

Use a Microsoft Azure Data Lake Store connection to access a Microsoft Azure Data Lake Store.

Note: The order of the connection properties might vary depending on the tool where you view them.

You can create and manage a Microsoft Azure SQL Data Warehouse connection in the Administrator tool or the Developer tool. The following table describes the Microsoft Azure Data Lake Store connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { }] \ : ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Microsoft Azure Data Lake Store.

The following table describes the properties for metadata access:

Property	Description
ADLS Account Name	The name of the Microsoft Azure Data Lake Store.
ClientID	The ID of your application to complete the OAuth Authentication in the Active Directory.
Client Secret	The client secret key to complete the OAuth Authentication in the Active Directory.
Directory	The Microsoft Azure Data Lake Store directory that you use to read data or write data. The default is root directory.
AuthEndpoint	The OAuth 2.0 token endpoint from where access code is generated based on based on the Client ID and Client secret is completed.

For more information about creating a client ID, client secret, and auth end point, contact the Azure administrator or see Microsoft Azure Data Lake Store documentation.

Microsoft Azure SQL Data Warehouse Connection Properties

Use a Microsoft Azure SQL Data Warehouse connection to access a Microsoft Azure SQL Data Warehouse.

Note: The order of the connection properties might vary depending on the tool where you view them.

You can create and manage a Microsoft Azure SQL Data Warehouse connection in the Administrator tool or the Developer tool. The following table describes the Microsoft Azure SQL Data Warehouse connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~`!\$%^&*()-+=[{]} \:;'"<, > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. It must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Microsoft Azure SQL Data Warehouse.

The following table describes the properties for metadata access:

Property	Description
Azure DW JDBC URL	Microsoft Azure Data Warehouse JDBC connection string. For example, you can enter the following connection string: jdbc:sqlserver:// <Server>.database.windows.net:1433;database=<Database>. The Administrator can download the URL from Microsoft Azure portal.
Azure DW JDBC Username	User name to connect to the Microsoft Azure SQL Data Warehouse account. You must have permission to read, write, and truncate data in Microsoft Azure SQL Data Warehouse.
Azure DW JDBC Password	Password to connect to the Microsoft Azure SQL Data Warehouse account.
Azure DW Schema Name	Name of the schema in Microsoft Azure SQL Data Warehouse.
Azure Blob Account Name	Name of the Microsoft Azure Storage account to stage the files.
Azure Blob Account Key	The key that authenticates the access to the Blob storage account.
Blob End-point	Type of Microsoft Azure end-points. You can select any of the following end-points: - core.windows.net: Default - core.usgovcloudapi.net: To select the US government Microsoft Azure end-points - core.chinacloudapi.cn: Not applicable You can configure the US government Microsoft Azure end-points when a mapping runs in the native environment and on the Spark engine.

Snowflake Connection Properties

When you set up a Snowflake connection, you must configure the connection properties.

Note: The order of the connection properties might vary depending on the tool where you view them.

The following table describes the Snowflake connection properties:

Property	Description
Name	The name of the connection. The name is not case sensitive and must be unique within the domain. You can change this property after you create the connection. The name cannot exceed 128 characters, contain spaces, or contain the following special characters: ~ ` ! \$ % ^ & * () - + = { [] } \ ; " ' < , > . ? /
ID	String that the Data Integration Service uses to identify the connection. The ID is not case sensitive. The ID must be 255 characters or less and must be unique in the domain. You cannot change this property after you create the connection. Default value is the connection name.
Description	Optional. The description of the connection. The description cannot exceed 4,000 characters.
Location	The domain where you want to create the connection.
Type	The connection type. Select Snowflake.
Username	The user name to connect to the Snowflake account.
Password	The password to connect to the Snowflake account.
Account	The name of the Snowflake account.
Warehouse	The Snowflake warehouse name.
Role	The Snowflake role assigned to the user.
Additional JDBC URL Parameters	Enter one or more JDBC connection parameters in the following format: <pre><param1>=<value>&<param2>=<value>&<param3>=<value>...</pre> For example: <pre>user=jon&warehouse=mywh&db=mydb&schema=public</pre> To access Snowflake through Okta SSO authentication, enter the web-based IdP implementing SAML 2.0 protocol in the following format: <pre>authenticator=https://<Your_Okta_Account_Name>.okta.com</pre> Note: Microsoft ADFS is not supported. For more information about configuring Okta authentication, see the following website: https://docs.snowflake.net/manuals/user-guide/admin-security-fed-auth-configure-snowflake.html#configuring-snowflake-to-use-federated-authentication

Creating a Connection to Access Sources or Targets

Create connections before you import data objects, preview data, and profile data.

1. Within the Administrator tool click **Manage > Connections**.
2. Select **Actions > New > Connection**.
3. Select the type of connection that you want to create:
 - To select an HBase connection, select **NoSQL > HBase**.
 - To select an HDFS connection, select **File Systems > Hadoop File System**.
 - To select a Hive connection, select **Database > Hive**.
 - To select a JDBC connection, select **Database > JDBC**.
4. Click **OK**.
5. Enter a connection name, ID, and optional description.
6. Configure the connection properties. For a Hive connection, you must choose the **Access Hive as a source or target** option to use Hive as a source or a target.
7. Click **Test Connection** to verify the connection.
8. Click **Finish**.

Creating a Hadoop Connection

Create a Hadoop connection before you run a mapping in the Hadoop environment.

1. Click **Window > Preferences**.
2. Select **Informatica > Connections**.
3. Expand the domain in the **Available Connections** list.
4. Select the **Cluster** connection type in the **Available Connections** list and click **Add**.
The **New Cluster Connection** dialog box appears.

5. Enter the general properties for the connection.

New Cluster Connection

Cluster Connection

Provide the connection details.

Name:

ID:

Description:

Location:

Type:

6. Click **Next**.
7. Enter the Hadoop cluster properties, common properties, and the reject directory properties.
8. Click **Next**.
9. Click **Next**.
- Effective in version 10.2.2, Informatica dropped support for the Hive engine. Do not enter Hive configuration properties.
10. Enter configuration properties for the Blaze engine and click **Next**.
11. Enter configuration properties for the Spark engine and click **Finish**.

Configuring Hadoop Connection Properties

When you create a Hadoop connection, default values are assigned to cluster environment variables, cluster path properties, and advanced properties. You can add or edit values for these properties. You can also reset to default values.

You can configure the following Hadoop connection properties based on the cluster environment and functionality that you use:

- Cluster Environment Variables
- Cluster Library Path

- Common Advanced Properties
- Blaze Engine Advanced Properties
- Spark Engine Advanced Properties

Note: Informatica does not recommend changing these property values before you consult with third-party documentation, Informatica documentation, or Informatica Global Customer Support. If you change a value without knowledge of the property, you might experience performance degradation or other unexpected results.

To reset to default values, delete the property values. For example, if you delete the values of an edited Cluster Library Path property, the value resets to the default \$DEFAULT_CLUSTER_LIBRARY_PATH.

Cluster Environment Variables

Cluster Environment Variables property lists the environment variables that the cluster uses. Each environment variable contains a name and a value. You can add environment variables or edit environment variables.

To edit the property in the text box, use the following format with &: to separate each name-value pair:

```
<name1>=<value1>[&:<name2>=<value2>...&:<nameN>=<valueN>]
```

Configure the following environment variables in the **Cluster Environment Variables** property:

HADOOP_NODE_JDK_HOME

Represents the directory from which you run the cluster services and the JDK version that the cluster nodes use. Required to run the Java transformation in the Hadoop environment and Sqoop mappings on the Blaze engine. Default is /usr/java/default. The JDK version that the Data Integration Service uses must be compatible with the JDK version on the cluster.

Set to <cluster JDK home>/jdk<version>.

For example, HADOOP_NODE_JDK_HOME=<cluster JDK home>/jdk<version>.

Cluster Library Path

Cluster Library Path property is a list of path variables for shared libraries on the cluster. You can add or edit library path variables.

To edit the property in the text box, use the following format with : to separate each path variable:

```
<variable1>[:<variable2>...:<variableN>]
```

Configure the library path variables in the **Cluster Library Path** property.

Common Advanced Properties

Common advanced properties are a list of advanced or custom properties that are unique to the Hadoop environment. The properties are common to the Blaze and Spark engines. Each property contains a name and a value. You can add or edit advanced properties.

To edit the property in the text box, use the following format with &: to separate each name-value pair:

```
<name1>=<value1>[&:<name2>=<value2>...&:<nameN>=<valueN>]
```

Configure the following property in the **Advanced Properties** of the common properties section:

infapdo.java.opts

List of Java options to customize the Java run-time environment. The property contains default values.

If mappings in a MapR environment contain a Consolidation transformation or a Match transformation, change the following value:

- `-Xmx512M`. Specifies the maximum size for the Java virtual memory. Default is 512 MB. Increase the value to at least 700 MB.
For example, `infapdo.java.opts=-Xmx700M`

Blaze Engine Advanced Properties

Blaze advanced properties are a list of advanced or custom properties that are unique to the Blaze engine. Each property contains a name and a value. You can add or edit advanced properties.

To edit the property in the text box, use the following format with `&:` to separate each name-value pair:

```
<name1>=<value1>[&:<name2>=<value2>...&:<nameN>=<valueN>]
```

Configure the following properties in the **Advanced Properties** of the Blaze configuration section:

infagrid.cadi.namespace

Namespace for the Data Integration Service to use. Required to set up multiple Blaze instances.

Set to `<unique namespace>`.

For example, `infagrid.cadi.namespace=TestUser1_namespace`

infagrid.blaze.console.jsfport

JSF port for the Blaze engine console. Use a port number that no other cluster processes use. Required to set up multiple Blaze instances.

Set to `<unique JSF port value>`.

For example, `infagrid.blaze.console.jsfport=9090`

infagrid.blaze.console.httpport

HTTP port for the Blaze engine console. Use a port number that no other cluster processes use. Required to set up multiple Blaze instances.

Set to `<unique HTTP port value>`.

For example, `infagrid.blaze.console.httpport=9091`

infagrid.node.local.root.log.dir

Path for the Blaze service logs. Default is `/tmp/infra/logs/blaze`. Required to set up multiple Blaze instances.

Set to `<local Blaze services log directory>`.

For example, `infagrid.node.local.root.log.dir=<directory path>`

infacal.hadoop.logs.directory

Path in HDFS for the persistent Blaze logs. Default is `/var/log/hadoop-yarn/apps/informatica`. Required to set up multiple Blaze instances.

Set to `<persistent log directory path>`.

For example, `infacal.hadoop.logs.directory=<directory path>`

infagrid.node.hadoop.local.root.log.dir

Path in the Hadoop connection for the service log directory.

Set to `<service log directory path>`.

For example, `infagrid.node.local.root.log.dir=$HADOOP_NODE_INFA_HOME/blazeLogs`

Spark Advanced Properties

Spark advanced properties are a list of advanced or custom properties that are unique to the Spark engine. Each property contains a name and a value. You can add or edit advanced properties. Each property contains a name and a value. You can add or edit advanced properties.

Configure the following properties in the **Advanced Properties** of the Spark configuration section:

To edit the property in the text box, use the following format with `&`: to separate each name-value pair:

```
<name1>=<value1>[&:<name2>=<value2>...&:<nameN>=<valueN>]
```

spark.authenticate

Enables authentication for the Spark service on Hadoop. Required for Spark encryption.

Set to `TRUE`.

For example, `spark.authenticate=TRUE`

spark.authenticate.enableSaslEncryption

Enables encrypted communication when SASL authentication is enabled. Required if Spark encryption uses SASL authentication.

Set to `TRUE`.

For example, `spark.authenticate.enableSaslEncryption=TRUE`

spark.executor.cores

Indicates the number of cores that each executor process uses to run tasklets on the Spark engine.

Set to: `spark.executor.cores=1`

spark.executor.instances

Indicates the number of instances that each executor process uses to run tasklets on the Spark engine.

Set to: `spark.executor.instances=1`

spark.executor.memory

Indicates the amount of memory that each executor process uses to run tasklets on the Spark engine.

Set to: `spark.executor.memory=3G`

infaspark.driver.cluster.mode.extraJavaOptions

List of extra Java options for the Spark driver that runs inside the cluster. Required for streaming mappings to read from or write to a Kafka cluster that uses Kerberos authentication.

For example, set to:

```
infaspark.driver.cluster.mode.extraJavaOptions=  
-Djava.security.egd=file:/dev/./urandom  
-XX:MaxMetaspaceSize=256M -Djavax.security.auth.useSubjectCredsOnly=true  
-Djava.security.krb5.conf=/<path to keytab file>/krb5.conf  
-Djava.security.auth.login.config=<path to jaas config>/kafka_client_jaas.config
```

To configure the property for a specific user, you can include the following lines of code:

```
infaspark.driver.cluster.mode.extraJavaOptions =  
-Djava.security.egd=file:/dev/./urandom  
-XX:MaxMetaspaceSize=256M -XX:+UseG1GC -XX:MaxGCPauseMillis=500  
-Djava.security.krb5.conf=/etc/krb5.conf
```

infaspark.executor.extraJavaOptions

List of extra Java options for the Spark executor. Required for streaming mappings to read from or write to a Kafka cluster that uses Kerberos authentication.

For example, set to:

```
infaspark.executor.extraJavaOptions=  
-Djava.security.egd=file:/dev/./urandom  
-XX:MaxMetaspaceSize=256M -Djavax.security.auth.useSubjectCredsOnly=true  
-Djava.security.krb5.conf=/<path to krb5.conf file>/krb5.conf  
-Djava.security.auth.login.config=/<path to jAAS config>/kafka_client_jaas.config
```

To configure the property for a specific user, you can include the following lines of code:

```
infaspark.executor.extraJavaOptions =  
-Djava.security.egd=file:/dev/./urandom  
-XX:MaxMetaspaceSize=256M -XX:+UseG1GC -XX:MaxGCPauseMillis=500  
-Djava.security.krb5.conf=/etc/krb5.conf
```

infaspark.flatfile.writer.nullValue

When the Databricks Spark engine writes to a target, it converts null values to empty strings (" "). For example, 12, AB,"",23p09udj.

The Databricks Spark engine can write the empty strings to string columns, but when it tries to write an empty string to a non-string column, the mapping fails with a type mismatch.

To allow the Databricks Spark engine to convert the empty strings back to null values and write to the target, configure the following advanced property in the Databricks Spark connection:

```
infaspark.flatfile.writer.nullValue=true
```

spark.hadoop.validateOutputSpecs

Validates if the HBase table exists. Required for streaming mappings to write to a HBase target in an Amazon EMR cluster. Set the value to false.

infaspark.json.parser.mode

Specifies the parser how to handle corrupt JSON records. You can set the value to one of the following modes:

- DROPMALFORMED. The parser ignores all corrupted records. Default mode.
- PERMISSIVE. The parser accepts non-standard fields as nulls in corrupted records.
- FAILFAST. The parser generates an exception when it encounters a corrupted record and the Spark application goes down.

infaspark.json.parser.multiLine

Specifies whether the parser can read a multiline record in a JSON file. You can set the value to true or false. Default is false. Applies only to non-native distributions that use Spark version 2.2.x and above.

infaspark.pythontx.exec

Required to run a Python transformation on the Spark engine for Big Data Management. The location of the Python executable binary on the worker nodes in the Hadoop cluster.

For example, set to:

```
infaspark.pythontx.exec=/usr/bin/python3.4
```

If you use the installation of Python on the Data Integration Service machine, set the value to the Python executable binary in the Informatica installation directory on the Data Integration Service machine.

For example, set to:

```
infaspark.pythontx.exec=INFA_HOME/services/shared/spark/python/lib/python3.4
```

infaspark.pythontx.executorEnv.PYTHONHOME

Required to run a Python transformation on the Spark engine for Big Data Management and Big Data Streaming. The location of the Python installation directory on the worker nodes in the Hadoop cluster.

For example, set to:

```
infaspark.pythontx.executorEnv.PYTHONHOME=/usr
```

If you use the installation of Python on the Data Integration Service machine, use the location of the Python installation directory on the Data Integration Service machine.

For example, set to:

```
infaspark.pythontx.executorEnv.PYTHONHOME=
INFA_HOME/services/shared/spark/python/
```

infaspark.pythontx.executorEnv.LD_PRELOAD

Required to run a Python transformation on the Spark engine for Big Data Streaming. The location of the Python shared library in the Python installation folder on the Data Integration Service machine.

For example, set to:

```
infaspark.pythontx.executorEnv.LD_PRELOAD=
INFA_HOME/services/shared/spark/python/lib/libpython3.6m.so
```

infaspark.pythontx.submit.lib.JEP_HOME

Required to run a Python transformation on the Spark engine for Big Data Streaming. The location of the Jep package in the Python installation folder on the Data Integration Service machine.

For example, set to:

```
infaspark.pythontx.submit.lib.JEP_HOME=
INFA_HOME/services/shared/spark/python/lib/python3.6/site-packages/jep/
```

spark.shuffle.encryption.enabled

Enables encrypted communication when authentication is enabled. Required for Spark encryption.

Set to TRUE.

For example, `spark.shuffle.encryption.enabled=TRUE`

spark.scheduler.maxRegisteredResourcesWaitingTime

The number of milliseconds to wait for resources to register before scheduling a task. Default is 30000. Decrease the value to reduce delays before starting the Spark job execution. Required to improve performance for mappings on the Spark engine.

Set to 15000.

For example, `spark.scheduler.maxRegisteredResourcesWaitingTime=15000`

spark.scheduler.minRegisteredResourcesRatio

The minimum ratio of registered resources to acquire before task scheduling begins. Default is 0.8. Decrease the value to reduce any delay before starting the Spark job execution. Required to improve performance for mappings on the Spark engine.

Set to: 0.5

For example, `spark.scheduler.minRegisteredResourcesRatio=0.5`

APPENDIX B

Data Type Reference

This appendix includes the following topics:

- [Data Type Reference Overview, 280](#)
- [Transformation Data Type Support in a Non-native Environment, 281](#)
- [Complex File and Transformation Data Types, 282](#)
- [Hive Data Types and Transformation Data Types, 285](#)
- [Sqoop Data Types, 287](#)

Data Type Reference Overview

Informatica Developer uses the following data types in mappings that run in a non-native environment:

Native data types

Native data types are specific to the sources and targets used as a physical data object. Native data types appear in the physical data object column properties.

Transformation data types

Transformation data types are set of data types that appear in the transformations. They are internal data types based on ANSI SQL-92 generic data types, which the Data Integration Service uses to move data across platforms. Transformation data types appear in all transformations in a mapping.

Transformation data types include the following data types:

- Primitive data type. Represents a single data value in a single column position.
- Complex data type. Represents multiple data values in a single column position. Use complex data types in mappings that run on the Spark engine to process hierarchical data in complex files.

When the Data Integration Service reads source data, it converts the native data types to the comparable transformation data types before transforming the data. When the Data Integration Service writes to a target, it converts the transformation data types to the comparable native data types.

Transformation Data Type Support in a Non-native Environment

The following table shows the Informatica transformation data type support on non-native run-time engines:

Transformation Data Type	Engine
Array	- Spark - Databricks Spark
Bigint	- Blaze - Spark - Databricks Spark
Binary	- Blaze - Spark
Date/Time	- Blaze - Spark - Databricks Spark
Decimal	- Blaze - Spark - Databricks Spark
Double	- Blaze - Spark - Databricks Spark
Integer	- Blaze - Spark - Databricks Spark
Map	- Spark - Databricks Spark
String	- Blaze - Spark - Databricks Spark
Struct	- Spark - Databricks Spark
Text	- Blaze - Spark - Databricks Spark
timestampWithTZ	Not supported

Complex File and Transformation Data Types

You can use complex data types in mappings to process hierarchical data in complex files.

You can use complex data types in the following complex files in mappings that run on the Spark and Databricks Spark engines:

- Avro
- JavaScript Object Notation (JSON)
- Parquet

Avro and Transformation Data Types

Apache Avro data types map to transformation data types that the Data Integration Service uses to move data across platforms.

The following table compares Avro data types and transformation data types:

Avro	Transformation	Description
Array	Array	Unlimited number of characters.
Boolean	Integer	-2,147,483,648 to 2,147,483,647 Precision of 10, scale of 0
Bytes	Binary	1 to 104,857,600 bytes
Double	Double	Precision of 15 digits
Fixed	Binary	1 to 104,857,600 bytes
Float	Double	Precision of 15 digits
Int	Integer	-2,147,483,648 to 2,147,483,647 Precision of 10, scale of 0
Long	Bigint	-9,223,372,036,854,775,808 to 9,223,372,036,854,775,807 Precision of 19, scale of 0
Map	Map	Unlimited number of characters.
Record	Struct	Unlimited number of characters.
String	String	1 to 104,857,600 characters
Union	Corresponding data type in a union of ["primitive_type complex_type", "null"] or ["null", "primitive_type complex_type"].	Dependent on primitive or complex data type.

Avro Union Data Type

A union indicates that a field might have more than one data type. For example, a union might indicate that a field can be a string or a null. A union is represented as a JSON array containing the data types.

The Developer tool only interprets a union of ["primitive_type|complex_type", "null"] or ["null", "primitive_type|complex_type"]. The Avro data type converts to the corresponding transformation data type. The Developer tool ignores the null.

Unsupported Avro Data Types

The Developer tool does not support the following Avro data types:

- enum
- null

JSON and Transformation Data Types

JavaScript Object Notation data types map to transformation data types that the Data Integration Service uses to move data across platforms.

The following table compares JSON data types and transformation data types:

JSON	Transformation	Description
Array	Array	Unlimited number of characters.
Double	Double	Precision of 15 digits
Integer	Integer	-2,147,483,648 to 2,147,483,647 Precision of 10, scale of 0
Object	Struct	Unlimited number of characters.
String	String	1 to 104,857,600 characters

Unsupported JSON Data Types

The Developer tool does not support the following JSON data types:

- date/timestamp
- enum
- union

Parquet and Transformation Data Types

Apache Parquet data types map to transformation data types that the Data Integration Service uses to move data across platforms.

The following table compares Parquet data types and transformation data types:

Parquet	Transformation	Description
Binary	Binary	1 to 104,857,600 bytes
Binary (UTF8)	String	1 to 104,857,600 characters

Parquet	Transformation	Description
Boolean	Integer	-2,147,483,648 to 2,147,483,647 Precision of 10, scale of 0
Double	Double	Precision of 15 digits
Fixed Length Byte Array	Decimal	Decimal value with declared precision and scale. Scale must be less than or equal to precision. For transformations that support precision up to 38 digits, the precision is 1 to 38 digits, and the scale is 0 to 38. For transformations that support precision up to 28 digits, the precision is 1 to 28 digits, and the scale is 0 to 28. If you specify the precision greater than the maximum number of digits, the Data Integration Service converts decimal values to double in high precision mode.
Float	Double	Precision of 15 digits
group (LIST)	Array	Unlimited number of characters.
Int32	Integer	-2,147,483,648 to 2,147,483,647 Precision of 10, scale of 0
Int64	Bigint	-9,223,372,036,854,775,808 to 9,223,372,036,854,775,807 Precision of 19, scale of 0
Int64 (TIMESTAMP_ MILLIS)	Date/Time	Jan 1, 0001 A.D. to Dec 31, 9999 A.D. Precision of 29, scale of 9 (precision to the nanosecond) Combined date/time value.
Int96	Date/Time	Jan 1, 0001 A.D. to Dec 31, 9999 A.D. Precision of 29, scale of 9 (precision to the nanosecond) Combined date/time value.
Map	Map	Unlimited number of characters.
Struct	Struct	Unlimited number of characters.
Union	Corresponding primitive data type in a union of ["primitive_type", "null"] or ["null", "primitive_type"].	Dependent on primitive data type.

Parquet Union Data Type

A union indicates that a field might have more than one data type. For example, a union might indicate that a field can be a string or a null. A union is represented as a JSON array containing the data types.

The Developer tool only interprets a union of ["primitive_type", "null"] or ["null", "primitive_type"]. The Parquet data type converts to the corresponding transformation data type. The Developer tool ignores the null.

Unsupported Parquet Data Types

The Developer tool does not support the following Parquet data types:

- int96 (TIMESTAMP_MILLIS)
- date
- timestamp

Hive Data Types and Transformation Data Types

The following table lists the Hive data types that Data Integration Service supports and the corresponding transformation data types:

Hive Data Type	Transformation Data Type	Range and Description
Binary	Binary	1 to 104,857,600 bytes. You can read and write data of Binary data type in a Hadoop environment. You can use the user-defined functions to transform the binary data type.
Tiny Int	Integer	-32,768 to 32,767
Integer	Integer	-2,147,483,648 to 2,147,483,647 Precision 10, scale 0
Bigint	Bigint	-9,223,372,036,854,775,808 to 9,223,372,036,854,775,807 Precision 19, scale 0
Decimal	Decimal	For transformations that support precision up to 38 digits, the precision is 1 to 38 digits, and the scale is 0 to 38. For transformations that support precision up to 28 digits, the precision is 1 to 28 digits, and the scale is 0 to 28. If a mapping is not enabled for high precision, the Data Integration Service converts all decimal values to double values. If a mapping is enabled for high precision, the Data Integration Service converts decimal values with precision greater than 38 digits to double values.
Double	Double	Precision 15
Float	Double	Precision 15
String	String	1 to 104,857,600 characters

Hive Data Type	Transformation Data Type	Range and Description
Boolean	Integer	1 or 0 The default transformation type for boolean is integer. You can also set this to string data type with values of True and False.
Array	String	1 to 104,857,600 characters
Struct	String	1 to 104,857,600 characters
Map	String	1 to 104,857,600 characters
Timestamp	datetime	The time stamp format is YYYY-MM-DD HH:MM:SS.ffffffff. Precision 29, scale 9.
Date	datetime	0000-0101 to 999912-31. Hive date format is YYYY-MM-DD. Precision 10, scale 0.
Char	String	1 to 255 characters
Varchar	String	1 to 65355 characters

Hive Complex Data Types

Hive complex data types such as arrays, maps, and structs are a composite of primitive or complex data types. Informatica Developer represents complex data types with the string data type and uses delimiters to separate the elements of the complex data type.

Note: Hive complex data types in a Hive source or Hive target are not supported when you run mappings on a Hadoop cluster.

The following table describes the transformation types and delimiters that are used to represent the complex data types:

Complex Data Type	Description
Array	The elements in the array are of string data type. The elements in the array are delimited by commas. For example, an array of <code>fruits</code> is represented as <code>[apple,banana,orange]</code> .
Map	Maps contain key-value pairs and are represented as pairs of strings and integers delimited by the <code>=</code> character. String and integer pairs are delimited by commas. For example, a map of <code>fruits</code> is represented as <code>[1=apple,2=banana,3=orange]</code> .
Struct	Structs are represented as pairs of strings and integers delimited by the <code>:</code> character. String and integer pairs are delimited by commas. For example, a struct of <code>fruits</code> is represented as <code>[1,apple]</code> .

Sqoop Data Types

When you use Sqoop, some variations apply in the processing. Sqoop supports a subset of data types that database vendors support.

Aurora Data Types

Informatica supports the following Aurora data types when you use Sqoop:

- Binary
- Bit
- Blob (supported only for import)
- Char
- Date
- Datetime
- Decimal
- Double
- Enum
- Float
- Integer
- Numeric
- Real
- Set
- Text
- Time
- Timestamp
- Varbinary
- Varchar

IBM DB2 and DB2 for z/OS Data Types

Informatica supports the following IBM DB2 and DB2 for z/OS data types when you use Sqoop:

- Bigint
- Blob (supported only for import)
- Char
- Clob
- Date
- DBClob
- DecFloat (supported only for import)
- Decimal
- Double (supported only for DB2 for z/OS)
- Float (supported only for DB2)

Graphic
Integer
LongVargraphic (supported only for DB2)
Numeric
Real
Smallint
Time
Timestamp
Varchar
Vargraphic
XML (supported only for import)

Greenplum Data Types

Informatica supports the following Greenplum data types when you use Sqoop:

Bigint
Bigserial
Bytea
Date
Decimal
Double
Integer
Nchar
Numeric
Nvarchar
Real
Serial
Smallint
Text
Time
Timestamp

Microsoft SQL Server Data Types

Informatica supports the following Microsoft SQL Server data types when you use Sqoop:

Bigint
Bit
Char
Datetime
Decimal
Float

INT
Money
Numeric
Real
Smalldatetime
Smallint
Smallmoney
Text
Time
Tinyint
Varchar

Rules and Guidelines for Sqoop Microsoft SQL Server Data Types

Consider the following rules and guidelines when you configure Microsoft SQL Server data types in a Sqoop mapping:

- If you create or replace the target table at run time and run the mapping on the Blaze or Spark engine to export Bigint data, the mapping fails.
- If you run a Sqoop mapping to export time data, Sqoop does not export milliseconds.

Netezza Data Types

Informatica supports the following Netezza data types when you use Sqoop:

Bigint
Blob (supported only for import)
Byteint
Char
Date
Double
Float4
Float8
Number
Timestamp
Varchar

Oracle Data Types

Informatica supports the following Oracle data types when you use Sqoop:

Blob (supported only for import)
Char
Date
Float
Long

Nchar (supported if you configure OraOop)
Nvarchar (supported if you configure OraOop)
Number(P,S)
Timestamp
Varchar
Varchar2

Rules and Guidelines for Sqoop Oracle Data Types

Consider the following rules and guidelines when you configure Oracle data types in a Sqoop mapping:

- If you run a Sqoop mapping on the Blaze engine to export Oracle float data, Sqoop truncates the data.
- If you run a Sqoop mapping on the Blaze engine to export Oracle timestamp data with nanoseconds, Sqoop writes only three digits to the target.
- If you configure OraOop and run a Sqoop mapping on the Spark engine to export Oracle timestamp data, Sqoop writes only three digits to the target.

Teradata Data Types

Informatica supports the following Teradata data types when you use Sqoop:

Bigint (supported only for import)
Blob (supported only for import)
Byteint
Char
Clob
Date
Decimal
Double
Float
Integer
Number
Numeric
Real
Smallint
Time
Timestamp
Varchar

Note: If you run a Sqoop mapping to export data of the Time data type, Sqoop does not export milliseconds.

Teradata Data Types with TDCH Specialized Connectors for Sqoop

Informatica supports the following Teradata data types when you use the Cloudera Connector Powered by Teradata, Hortonworks Connector for Teradata, and MapR Connector for Teradata with Sqoop:

Bigint

Byte (supported only by Hortonworks Connector for Teradata and MapR Connector for Teradata)
Byteint
Character
Date
Decimal
Double Precision/Float/Real
Integer
Number(P,S)
Numeric
Smallint
Time (supported only by Cloudera Connector Powered by Teradata and Hortonworks Connector for Teradata)
Timestamp
Varchar
Varbyte (supported only by Hortonworks Connector for Teradata and MapR Connector for Teradata)

Vertica Data Types

Informatica supports the following Vertica data types when you use Sqoop:

Boolean (supported only on the Spark engine)
Bigint
Char
Date
Datetime
Decimal
Double precision
Float
Float (n)
Float 8
Int
Int8
Integer
Long Varchar
Money
Number
Numeric
Real
Smalldatetime
Smallint
Time
Timestamp

Tinyint

Varchar

APPENDIX C

Function Reference

This appendix includes the following topics:

- [Function Support in a Non-native Environment, 293](#)
- [Function and Data Type Processing, 295](#)

Function Support in a Non-native Environment

Some Informatica transformation language functions that are valid in the native environment might be restricted or unsupported in a non-native environment, such as Hadoop or Databricks. Functions not listed in this table are supported on all engines without restrictions.

Important: When you push a mapping to a non-native environment, the engine that processes the mapping uses a set of rules different from the Data Integration Service. As a result, the mapping results can vary based on the rules that the engine uses.

The following table lists functions and levels of support for functions on different engines in the non-native environments.

Function	Blaze Engine	Spark Engine	Databricks Spark
ABORT	Not supported	Not supported	Not supported
ANY	Supported	Supported	Supported
ARRAY	Not supported	Supported	Supported
AVG	Supported	Supported	Supported
CAST	Not supported	Supported	Supported
COLLECT_LIST	Not supported	Supported	Supported
COLLECT_MAP	Not supported	Supported	Supported
CONCAT_ARRAY	Not supported	Supported	Supported
COUNT	Supported	Supported	Supported
CREATE_TIMESTAMP_TZ	Supported	Not supported	Not supported

Function	Blaze Engine	Spark Engine	Databricks Spark
CUME	Not supported	Not supported	Not supported
DECOMPRESS	Supported	Supported	Supported
ERROR	Not supported	Not supported	Not supported
FIRST	Not supported	Not supported	Not supported
GET_TIMEZONE	Supported	Not supported	Not supported
GET_TIMESTAMP	Supported	Not supported	Not supported
LAST	Not supported	Not supported	Not supported
LAG	Not supported	Supported	Supported
LEAD	Not supported	Supported	Supported
MAP	Not supported	Supported	Supported
MAP_FROM_ARRAYS	Not supported	Supported	Supported
MAP_KEYS	Not supported	Supported	Supported
MAP_VALUES	Not supported	Supported	Supported
MAX (Dates)	Supported	Supported	Supported
MAX (Numbers)	Supported	Supported	Supported
MAX (String)	Supported	Supported	Supported
METAPHONE	Supported	Supported	Supported
MIN (Dates)	Supported	Supported	Supported
MIN (Numbers)	Supported	Supported	Supported
MIN (String)	Supported	Supported	Supported
MOVINGAVG	Not supported	Not supported	Not supported
MOVINGSUM	Not supported	Not supported	Not supported
PERCENTILE	Supported	Supported	Supported
RESPEC	Not supported	Supported	Supported
SIZE	Not supported	Supported	Supported
STDDEV	Supported	Supported	Supported
STRUCT	Not supported	Supported	Supported

Function	Blaze Engine	Spark Engine	Databricks Spark
STRUCT_AS	Not supported	Supported	Supported
SUM	Supported	Supported	Supported
SYSTIMESTAMP	Supported	Supported	Supported
TO_DATE	Supported	Supported	Supported
TO_DECIMAL	Supported	Supported	Supported
TO_DECIMAL38	Supported	Supported	Supported
TO_TIMESTAMP_TZ	Supported	Not supported	Not supported
UUID4	Supported	Supported	Supported
UUID_UNPARSE	Supported	Supported	Supported
VARIANCE	Supported	Supported	Supported

For more information, see the *Informatica Transformation Language Reference*.

Function and Data Type Processing

When you run a mapping in a Hadoop environment, the Hadoop engine might process Informatica functions and data types differently from the Data Integration Service. Also, each of the run-time engines in the Hadoop environment can process Informatica functions and data types differently. Therefore, some variations apply in the processing and validity of functions and data types, and mapping results can vary.

Rules and Guidelines for Spark Engine Processing

Some restrictions and guidelines apply to processing Informatica functions on the Spark engine.

Important: When you push a mapping to the Hadoop environment, the engine that processes the mapping uses a set of rules different from the Data Integration Service. As a result, the mapping results can vary based on the rules that the engine uses. This topic contains some processing differences that Informatica discovered through internal testing and usage. Informatica does not test all the rules of the third-party engines and cannot provide an extensive list of the differences.

Consider the following rules and guidelines for function and data type processing on the Spark engine:

- The Spark engine and the Data Integration Service process overflow values differently. The Spark engine processing rules might differ from the rules that the Data Integration Service uses. As a result, mapping results can vary between the native and Hadoop environment when the Spark engine processes an overflow. Consider the following processing variation for Spark:
 - If an expression results in numerical errors, such as division by zero or SQRT of a negative number, the Spark engine returns NULL. In the native environment, the same expression results in a row error.

- The Spark engine and the Data Integration Service process data type conversions differently. As a result, mapping results can vary between the native and Hadoop environment when the Spark engine performs a data type conversion. Consider the following processing variations for Spark:
 - The results of arithmetic operations on floating point types, such as Decimal, can vary between the native environment and a Hadoop environment. The difference between the results can increase across multiple operations.
 - When you use the TO_DECIMAL or TO_DECIMAL38 function in a mapping that runs in high-precision mode, you must specify a scale argument. If the mapping runs in low-precision mode, the Spark engine ignores the scale argument and returns a double.
 - When the number of fractional digits in a double or decimal value exceeds the scale that is configured in a decimal port, the Spark engine trims trailing digits, rounding the value if necessary.
 - If you use Hive 2.0 or higher, the Spark engine guarantees scale values. For example, when the Spark engine processes the decimal 1.1234567 with scale 9 using Hive 2.0, the output is 1.123456700.
 - The Spark engine cannot process dates to the nanosecond. It can return a precision for date/time data up to the microsecond.
- If data overflow occurs, the Spark engine returns NULL. For example, if you use the expression `TO_DECIMAL(12.34, 2)` in a port that has a precision of 3 and a scale of 2, the return value is NULL. The NULL value will be propagated through the mapping. The mapping might overwrite it using a default value, detect it using the function `IS_NULL`, and write the NULL value to the target.
- If you enable high-precision in a streaming mapping, the Spark engine runs the mapping in low-precision mode.
- If the mapping contains a complex port with an element of a decimal data type, the Spark engine runs the mapping in low-precision mode.
- The Hadoop environment treats `"/n"` values as null values. If an aggregate function contains empty or NULL values, the Hadoop environment includes these values while performing an aggregate calculation.
- Mapping validation fails if you configure SYSTIMESTAMP with a variable value, such as a port name. The function can either include no argument or the precision to which you want to retrieve the timestamp value.
- Mapping validation fails if an output port contains a Timestamp with Time Zone data type.
- Avoid including single and nested functions in an Aggregator transformation. The Data Integration Service fails the mapping in the native environment. It can push the processing to the Hadoop environment, but you might get unexpected results. Informatica recommends creating multiple transformations to perform the aggregation.
- You cannot preview data for a transformation that is configured for windowing.
- The Spark METAPHONE function uses phonetic encoders from the `org.apache.commons.codec.language` library. When the Spark engine runs a mapping, the METAPHONE function can produce an output that is different from the output in the native environment. The following table shows some examples:

String	Data Integration Service	Spark Engine
Might	MFT	MT
High	HF	H

- If you use the TO_DATE function on the Spark engine to process a string written in ISO standard format, you must add `*T*` to the date string and `*"T"*` to the format string. The following expression shows an

example that uses the TO_DATE function to convert a string written in the ISO standard format YYYY-MM-DDTHH24:MI:SS:

```
TO_DATE ( '2017-11-03*T*12:45:00' , 'YYYY-MM-DD*T"*HH24:MI:SS' )
```

The following table shows how the function converts the string:

ISO Standard Format	RETURN VALUE
2017-11-03T12:45:00	Nov 03 2017 12:45:00

- The UUID4 function is supported only when used as an argument in UUID_UNPARSE or ENC_BASE64.
- The UUID_UNPARSE function is supported only when the argument is UUID4().

INDEX

A

- accessing elements
 - array [174](#)
 - map [174](#)
 - struct [174](#)
- Address Validator transformation
 - non-native environment [87](#)
- ADLS
 - complex file sources [61](#), [63](#)
 - complex file targets [74](#), [76](#)
- aggregate window function
 - example [222](#)
 - nesting [224](#)
 - offsets [224](#)
- Aggregator transformation
 - Blaze engine [88](#)
 - Databricks Spark engine [89](#)
 - non-native environment [88](#)
 - Spark engine [89](#)
- Amazon AWS [235](#)
- Amazon Redshift connection
 - properties [240](#)
- Amazon S3
 - complex file sources [62](#)
 - complex file targets [76](#)
- Amazon S3 connection
 - properties [241](#)
- architecture
 - Big Data Management [16](#)
 - Big Data Management with Databricks [22](#)
 - Hadoop environment [17](#)
- array
 - complex data type [159](#)
 - accessing elements [174](#)
 - dimension [159](#)
 - example [159](#)
 - extracting element [175](#)
 - format [159](#)
 - generating [176](#)
 - index [159](#)
 - multi-dimensional [159](#)
 - one-dimensional [159](#)
- array functions
 - ARRAY [176](#)
 - COLLECT_LIST [176](#)
 - CONCAT_ARRAY [176](#)
 - SIZE [176](#)
- array port
 - type configuration [170](#)
- AutoDeploy [123](#)
- Azure
 - configuration [236](#)
- Azure Blob
 - complex file sources [64](#)
 - complex file targets [77](#)

- Azure Blob Storage
 - complex file sources [61](#)
 - complex file targets [74](#)

B

- big data
 - application services [16](#)
 - big data process [23](#)
 - repositories [17](#)
- big data process
 - collect your data [23](#)
- Blaze engine
 - Blaze engine architecture [20](#)
 - connection properties [249](#)
 - mapping properties [145](#)
 - monitoring [141](#), [143](#), [144](#)
 - segment time [144](#)
 - summary report [144–147](#)
 - tasklet execution time [146](#)
 - tasklet information [147](#)
- Blaze execution plan
 - monitoring [138](#)
- Blaze Job Monitor
 - logging [148](#)

C

- Case Converter transformation
 - non-native environment [90](#)
- Cassandra connections
 - properties [243](#)
- Classifier transformation
 - non-native environment [90](#)
- cloud platform clusters [112](#)
- cloud provisioning configuration
 - Databricks properties [239](#)
 - Amazon AWS properties [235](#)
 - Microsoft Azure properties [236](#)
- cluster workflow
 - cloud provisioning connection [234](#)
- cluster workflows
 - Azure job monitoring [124](#)
 - components [113](#), [123](#)
 - Delete Cluster task [123](#)
 - external RDS as the Hive metastore db [121](#)
 - mappings in cluster workflows [123](#)
 - monitoring [124](#)
 - Overview [112](#)
 - running [124](#)
 - Running mappings on Blaze engine [119](#)
 - sample [113](#)
 - workflow tasks [123](#)

- Comparison transformation
 - non-native environment [90](#)
- complex data type
 - array [158](#)
 - map [158](#)
 - struct [158](#)
- complex data type definition
 - creating [167](#)
 - example [165](#)
 - importing [168](#)
 - nested data type definition [167](#)
 - recursive data type definition [167](#)
 - struct [165](#)
 - type definition library [165](#)
- complex data types
 - array [159](#)
 - map [161](#)
 - struct [161](#)
- complex file formats
 - Avro [282](#)
 - JSON [283](#)
 - overview [282](#)
 - Parquet [283](#)
- complex file sources
 - ADLS [61](#), [63](#)
 - Amazon S3 [62](#)
 - Azure Blob [64](#)
 - Azure Blob Storage [61](#)
 - HDFS [65](#)
 - MapR-FS [65](#)
- complex file targets
 - ADLS [74](#), [76](#)
 - Amazon S3 [76](#)
 - Azure Blob [77](#)
 - Azure Blob Storage [74](#)
 - Databricks environment [73](#)
 - HDFS [78](#)
 - MapR-FS [78](#)
- complex functions
 - ARRAY [176](#)
 - CAST [176](#)
 - COLLECT_LIST [176](#)
 - COLLECT_MAP [176](#)
 - CONCAT_ARRAY [176](#)
 - MAP [176](#)
 - MAP_FROM_ARRAYS [176](#)
 - MAP_KEYS [176](#)
 - MAP_VALUES [176](#)
 - RESPEC [176](#)
 - SIZE [176](#)
 - STRUCT [176](#)
 - STRUCT_AS [176](#)
- complex operator
 - dot operator [174](#)
 - example [174](#)
 - subscript operator [174](#)
- complex port
 - array [163](#)
 - creating [165](#)
 - map [163](#)
 - nested [163](#)
 - Read transformation [164](#)
 - struct [163](#)
 - transformations [164](#)
 - type configuration [163](#), [170](#)
 - Write transformation [164](#)
- component architecture
 - clients and tools [16](#)

- configure a mapping
 - non-native [41](#)
- configuring
 - array port properties [170](#)
 - map port properties [172](#)
 - struct port
 - properties [173](#)
- Connection
 - details [267](#)
 - properties [267](#)
- connection properties
 - Databricks [244](#)
- connections
 - properties [234](#), [249](#)
 - HBase [234](#)
 - HDFS [234](#)
 - Hive [234](#)
 - JDBC [234](#)
- Consolidation transformation
 - non-native environment [91](#)
- conversion
 - hierarchical to relational [181](#)
- Cosmos DB connection
 - creating [269](#)
- Create Cluster task
 - properties [115](#)
- creating
 - Cosmos DB connection [269](#)
- creating a column profile
 - profiles [129](#)

D

- Data Masking transformation
 - non-native environment [91](#)
- data object
 - processing data with an intelligent structure [211](#)
- data object profiles
 - creating a single profile [127](#)
 - enterprise discovery [127](#)
- Data Processor transformation
 - Blaze engine [93](#)
 - non-native environment [93](#)
- data type
 - complex data type [158](#)
 - nested data type [158](#)
 - primitive data type [158](#)
- data types
 - complex files [282](#)
 - Hive [285](#)
 - Hive complex data types [286](#)
 - processing in a Hadoop environment [295](#)
 - support [281](#)
- Databricks
 - cloud provisioning configuration [239](#)
- Databricks connection properties [244](#)
- Databricks environment
 - flat file sources [61](#)
 - flat file targets [74](#)
 - complex file targets [73](#)
 - sources [60](#)
- Databricks integration
 - overview [19](#)
- Databricks Spark
 - execution plans [44](#)
- Decision transformation
 - non-native environment [94](#)

Delete Cluster task [123](#)

dot operator

extracting struct element [175](#)

dynamic complex port

input rules [201](#)

E

enterprise discovery

running in Informatica Analyst [130](#)

ephemeral clusters

cloud provisioning connection [234](#)

Overview [112](#)

example

aggregate window function [222](#)

array [159](#)

complex data type definition [165](#)

complex operator [174](#)

dot operator [174](#)

map [161](#)

nested data type definition [167](#)

partition and order keys [219](#)

struct [161](#)

subscript operator [174](#)

windowing [226](#), [228](#)

execution plan

Databricks Spark engine [44](#)

Spark engine [44](#)

Expression transformation

Blaze engine [94](#)

Databricks Spark engine [95](#)

non-native environment [94](#)

Spark engine [94](#)

external RDS as the Hive metastore [121](#)

F

file sources

Hadoop environment [62](#)

file targets

Hadoop environment [75](#)

Filter transformation

Blaze engine [95](#)

non-native environment [95](#)

flat file source [66](#)

flat file sources

in Databricks environment [61](#)

in Hadoop environment [65](#)

flat file targets

in Databricks environment [74](#)

functions

processing in a Hadoop environment [295](#)

G

Google Analytics connections

properties [246](#)

Google BigQuery connection

properties [246](#)

Google Cloud Spanner connection

properties [247](#)

Google Cloud Storage connections

properties [248](#)

grid

description [57](#)

H

Hadoop [234](#)

Hadoop connections

creating [273](#)

Hadoop environment

file sources [62](#)

file targets [75](#)

flat file limitations [65](#)

flat file targets [79](#)

Hive targets [80](#)

logs [135](#)

optimization [45](#)

relational sources [66](#)

Sqoop sources restrictions [70](#)

Sqoop targets restrictions [84](#)

hadoop utilities

Sqoop [18](#)

HBase connections

MapR-DB properties [256](#)

properties [256](#)

HDFS

complex file sources [65](#)

complex file targets [78](#)

HDFS connections

creating [273](#)

properties [254](#)

HDFS mappings

data extraction example [53](#)

description [53](#)

hierarchical conversion

generate nested struct [184](#)

generate struct [182](#)

hierarchical to hierarchical [192](#)

hierarchical to relational [192](#), [194](#)

relational to hierarchical [182](#), [184](#)

hierarchical data

data preview process [178](#)

complex files [155](#)

converting [181](#)

data preview [177](#)

data preview interface [178](#)

data viewer [179](#)

extracting elements [192](#)

flattening elements [194](#)

how to process [156](#)

modifying [192](#), [194](#)

previewing hierarchical data [180](#)

processing [155](#)

processing semi-structured data [209](#)

processing unstructured data [209](#)

rules and guidelines for hierarchical data preview [180](#)

Spark engine [155](#)

hierarchy data

hierarchical type panel [180](#)

high availability

description [58](#)

Hive

target limitations [80](#)

Hive connections

creating [273](#)

properties [257](#)

Hive execution plan

monitoring [138](#)

Hive mappings

description [55](#)

Hive metastore [121](#)

- Hive pushdown
 - connection properties [249](#)
- Hive query plan
 - viewing, for mapping [45](#)
 - viewing, for profile [132](#)
- Hive sources
 - postSQL [67](#)
 - preSQL [67](#)
 - with Blaze engine [68](#)
 - with Informatica mappings [67](#)
- Hive target
 - truncate [81](#)
- Hive targets
 - postSQL [81](#)
 - preSQL [81](#)
 - with Blaze engine [82](#)
- how to
 - process hierarchical data [156](#)

I

- Informatica Big Data Management
 - overview [15](#)
- Informatica engine
 - Informatica engine execution plan [42](#)
- Informatica Intelligent Cloud Services account
 - creating [215](#)
- input rules
 - dynamic array [201](#)
 - dynamic complex port [201](#)
 - dynamic map [201](#)
 - dynamic struct [201](#)
- intelligent structure discovery
 - process [210](#)
- Intelligent Structure Discovery
 - license [212](#)
- intelligent structure mappings
 - data conversion example [214](#)
- intelligent structure model
 - creating [215](#)
 - exporting [216](#)
 - overview [209](#)
- intelligent structures
 - use case [210](#)

J

- Java transformation
 - Blaze engine [96](#)
 - Non-native environment [95](#)
 - Spark engine [96](#)
- JDBC connections
 - connection string [262](#)
 - properties [260](#)
- Joiner transformation
 - Blaze engine [97](#)
 - Databricks Spark engine [98](#)
 - non-native environment [97](#)
 - Spark engine [98](#)

K

- Key Generator transformation
 - Non-native environment [98](#)

L

- Labeler transformation
 - Non-native environment [99](#)
- logging
 - mapping run on Hadoop [148](#)
 - Spark engine [154](#)
- logs
 - Blaze engine [147](#)
 - Hadoop environment [135](#)
 - Spark engine [153](#)
- logs URL
 - YARN web user interface [154](#)
- Lookup transformation
 - Blaze engine [99](#)
 - Databricks Spark engine [101](#)
 - Non-native environment [99](#)
 - Spark engine [99](#)

M

- map
 - complex data type [161](#)
 - accessing elements [174](#)
 - example [161](#)
 - format [161](#)
 - generating [176](#)
 - key-value pair [161](#)
- map functions
 - COLLECT_MAP [176](#)
 - MAP [176](#)
 - MAP_FROM_ARRAYS [176](#)
 - MAP_KEYS [176](#)
 - MAP_VALUES [176](#)
 - SIZE [176](#)
- map port
 - type configuration [172](#)
- mapping
 - intelligent structure [211](#)
- mapping example
 - Hive [55](#)
 - Twitter [56](#)
- mapping execution plans
 - overview [42](#)
- mapping run on Hadoop
 - logging [137](#)
 - monitoring [138](#)
- mapping run on non-native environment
 - overview [27](#)
- MapR-FS
 - complex file sources [65](#)
 - complex file targets [78](#)
- Match transformation
 - Blaze engine [101](#)
 - Non-native environment [101](#)
 - Spark engine [102](#)
- Merge transformation
 - Non-native environment [102](#)
- Microsoft Azure [236](#)
- Microsoft Azure Data Lake Store connection
 - properties [269](#)
- Microsoft Azure SQL Data Warehouse connection
 - properties [270](#)
- monitoring cluster workflow jobs [124](#)
- Monitoring URL
 - Blaze and Spark jobs [137](#)

N

- native environment
 - high availability [58](#)
 - mappings [53](#)
 - optimization [56](#)
 - partitioning [57](#)
- nested data type definition
 - example [167](#)
- nested struct
 - generating [184](#)
- node labeling
 - Blaze [50](#)
- non-native environment
 - transformation support [85](#)
 - valid sources [59](#)
- non-native execution plan
 - description, for mapping [27](#)
 - overview [42](#)
- non-native mapping
 - run-time properties [28](#)
- non-native mappings
 - workflows [40](#)
- Normalizer transformation
 - Non-native environment [102](#)

O

- optimization
 - compress temporary staging tables [47](#)
 - node labeling [49](#)
 - queuing [50](#)
 - scheduling [49](#)
 - truncate partitions [48](#)
- overview
 - processing hierarchical data [155](#)

P

- Parser transformation
 - Non-native environment [102](#)
- partitioning
 - description [57](#)
 - optimization [57](#)
- postSQL
 - for Hive sources [67](#)
 - for Hive targets [81](#)
- preSQL
 - for Hive sources [67](#)
 - for Hive targets [81](#)
- processing hierarchical data
 - overview [155](#)
- processing with an intelligent structure
 - in a mapping [211](#)
 - running on Spark [211](#)
- profile run on Blaze engine
 - Overview [125](#)
- profile run on Hadoop
 - monitoring [132](#)
- profile run on Hive
 - Overview [125](#)
- profiles
 - creating a column profile [129](#)
- Python transformation
 - Non-native environment [103](#)
 - Spark engine [103](#)

R

- Rank transformation
 - Blaze engine [103](#)
 - Databricks Spark engine [104](#)
 - Non-native environment [103](#)
 - Spark engine [104](#)
- registration
 - Informatica Intelligent Cloud Services account [215](#)
- relational to hierarchical [181](#)
- Router transformation
 - Non-native environment [105](#)
- rules and guidelines
 - Spark engine [295](#)
 - window functions [226](#)
 - windowing properties [221](#)
- run-time properties
 - non-native mapping [28](#)

S

- semi-structured data
 - processing [209](#)
- Sequence Generator transformation
 - Blaze engine [105](#)
 - Non-native environment [105](#)
 - Spark engine [105](#)
- Snowflake connection
 - properties [272](#)
- social media mappings
 - description [55](#)
- Sorter transformation
 - Blaze engine [106](#)
 - Databricks Spark engine [107](#)
 - Non-native environment [105](#)
 - Spark engine [106](#)
- source file name
 - for flat file sources [66](#)
- sources
 - in a non-native environment [59](#)
 - Databricks environment [60](#)
- Spark
 - execution plans [44](#)
- Spark deploy mode
 - Hadoop connection properties [249](#)
- Spark engine
 - data type processing [295](#)
 - function processing [295](#)
 - rules and guidelines [295](#)
 - connection properties [249](#)
 - hierarchical data [155](#)
 - monitoring [149](#)
- Spark Event Log directory
 - Hadoop connection properties [249](#)
- Spark execution parameters
 - Hadoop connection properties [249](#)
- Spark HDFS staging directory
 - Hadoop connection properties [249](#)
- Sqoop arguments
 - parameterization [38](#)
- Sqoop configuration
 - mapping properties [38](#)
 - profiling [126](#)
- Sqoop connection arguments
 - Dsqoop.connection.factories [264](#)
 - connect [264](#)
 - direct [264](#)

Sqoop connection arguments (*continued*)

driver [264](#)

Sqoop connectivity

supported data types [287](#)

Sqoop data types

Aurora [287](#)

Greenplum [288](#)

IBM DB2 [287](#)

IBM DB2 for z/OS [287](#)

Microsoft SQL Server [288](#)

Netezza [289](#)

Oracle [289](#)

Teradata [290](#)

Teradata Data Types with TDCH Sqoop Specialized Connectors [290](#)

Vertica [291](#)

Sqoop incremental data extraction

–infa-incremental-key [37](#)

–infa-incremental-type [37](#)

–infa-incremental-value [37](#)

–infa-incremental-value-format [37](#)

configuring [37](#)

Sqoop mapping arguments

batch [35](#)

infaoptimize [36](#)

infaownername [36](#)

m [34](#)

num-mappers [34](#)

schema [36](#)

split-by [35](#)

verbose [37](#)

Sqoop mappings

overview [34](#)

supported engines [34](#)

Sqoop sources

importing data from Vertica [70](#)

in Hadoop environment [70](#)

Sqoop targets

in Hadoop environment [84](#)

Standardizer transformation

Non-native environment [108](#)

stateful computing

overview [217](#)

struct

complex data type [161](#)

accessing elements [174](#)

changing data type [176](#)

changing schema [176](#)

complex data type definition [165](#)

example [161](#)

extracting element [175](#)

format [161](#)

generating [176](#), [182](#)

name-type pair [161](#)

schema [161](#)

struct functions

CAST [176](#)

RESPEC [176](#)

STRUCT [176](#)

STRUCT_AS [176](#)

struct port

complex data type definition [173](#)

type configuration [173](#)

subscript operator

extracting array element [175](#)

T

targets

flat files in Hadoop mapping [79](#)

TDCH connection factory

-Dscoop.connection.factories [264](#)

third-party tools

hadoop cluster [18](#)

transformations

in a non-native environment [85](#)

type configuration

array [170](#)

map [170](#)

struct [170](#)

struct port [173](#)

type definition library

complex data type definition [165](#)

mapping [165](#)

mapplet [165](#)

nested data type definition [167](#)

U

Union transformation

Non-native environment [108](#)

unstructured data

processing [209](#)

Update Strategy transformation

Blaze engine [109](#)

Non-native environment [108](#)

Spark engine [109](#)

utilities

hadoop cluster [18](#)

V

validation environments

non-native mapping [28](#)

W

Weighted Average transformation

Non-native environment [111](#)

window functions

aggregate functions [222](#)

LAG [222](#)

LEAD [222](#)

overview [221](#)

rules and guidelines [226](#)

windowing

example [226](#), [228](#), [230](#)

overview [217](#)

properties [218](#)

windowing properties

frame [218](#), [224](#)

order [218](#), [219](#)

partition [218](#), [219](#)

rules and guidelines [221](#)

workflows

non-native mappings [40](#)

Y

YARN

queues [49](#)

scheduler

capacity [49](#)

YARN (*continued*)

scheduler (*continued*)

fair [49](#)

YARN web user interface

description [136](#)