

Machine Learning 1.5.0

Command Line Reference

Date published: 2020-07-16

Date modified: 2023-01-31

CLOUDERA

<https://docs.cloudera.com/>

Legal Notice

© Cloudera Inc. 2025. All rights reserved.

The documentation is and contains Cloudera proprietary information protected by copyright and other intellectual property rights. No license under copyright or any other intellectual property right is granted herein.

Unless otherwise noted, scripts and sample code are licensed under the Apache License, Version 2.0.

Copyright information for Cloudera software may be found within the documentation accompanying each component in a particular release.

Cloudera software includes software from various open source or other third party projects, and may be released under the Apache Software License 2.0 (“ASLv2”), the Affero General Public License version 3 (AGPLv3), or other license terms. Other software included may be released under the terms of alternative open source licenses. Please review the license and notice files accompanying the software for additional licensing information.

Please visit the Cloudera software product page for more information on Cloudera software. For more information on Cloudera support services, please visit either the Support or Sales page. Feel free to contact us directly to discuss your specific needs.

Cloudera reserves the right to change any products at any time, and without notice. Cloudera assumes no responsibility nor liability arising from the use of products, except as expressly agreed to in writing by Cloudera.

Cloudera, Cloudera Altus, HUE, Impala, Cloudera Impala, and other Cloudera marks are registered or unregistered trademarks in the United States and other countries. All other trademarks are the property of their respective owners.

Disclaimer: EXCEPT AS EXPRESSLY PROVIDED IN A WRITTEN AGREEMENT WITH CLOUDERA, CLOUDERA DOES NOT MAKE NOR GIVE ANY REPRESENTATION, WARRANTY, NOR COVENANT OF ANY KIND, WHETHER EXPRESS OR IMPLIED, IN CONNECTION WITH CLOUDERA TECHNOLOGY OR RELATED SUPPORT PROVIDED IN CONNECTION THEREWITH. CLOUDERA DOES NOT WARRANT THAT CLOUDERA PRODUCTS NOR SOFTWARE WILL OPERATE UNINTERRUPTED NOR THAT IT WILL BE FREE FROM DEFECTS NOR ERRORS, THAT IT WILL PROTECT YOUR DATA FROM LOSS, CORRUPTION NOR UNAVAILABILITY, NOR THAT IT WILL MEET ALL OF CUSTOMER’S BUSINESS REQUIREMENTS. WITHOUT LIMITING THE FOREGOING, AND TO THE MAXIMUM EXTENT PERMITTED BY APPLICABLE LAW, CLOUDERA EXPRESSLY DISCLAIMS ANY AND ALL IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO IMPLIED WARRANTIES OF MERCHANTABILITY, QUALITY, NON-INFRINGEMENT, TITLE, AND FITNESS FOR A PARTICULAR PURPOSE AND ANY REPRESENTATION, WARRANTY, OR COVENANT BASED ON COURSE OF DEALING OR USAGE IN TRADE.

Contents

Command Line Tools in CML.....	4
cdswctl Command Line Interface Client.....	4
Download and Configure cdswctl.....	4
(Optional) Generate an SSH Public/Private Key.....	5
Download cdswctl and Add an SSH Key.....	5
Initialize an SSH Endpoint.....	5
Log into cdswctl.....	7
Prepare to manage models using the model CLI.....	7
Create a model using the CLI.....	8
Build and deployment commands for models.....	9
Deploy a new model with updated resources.....	10
View replica logs for a model.....	10
cdswctl command reference.....	11

Command Line Tools in CML

Cloudera Machine Learning ships with the following command line tools. The purpose of each tool differs.

- CDP CLI for Cloudera Machine Learning - If you prefer to work in a terminal window, you can download and configure the CDP client that gives you access to the CDP CLI tool. The CDP CLI allows you to perform the same actions as can be performed from the management console. Use this CLI to create, delete, upgrade, and manage ML workspaces on CDP.

To view all the available commands, run:

```
cdp ml help
```

To view help for a specific command, run:

```
cdp ml <operation> help
```

If you don't already have the CDP CLI set up, see *Installing the CDP CLI Client*.

- cdswctl - Cloudera Machine Learning also ships with a CLI client that you can download from the Cloudera Machine Learning web UI. This is also referred to as the Model CLI client. The cdswctl client allows you to log in, create an SSH endpoint, launch new sessions, automate model deployment, model updates, and so on.

cdswctl Command Line Interface Client

Cloudera Machine Learning ships with a CLI client that you can download from the Cloudera Machine Learning web UI. The cdswctl client allows you to perform the following tasks:

- Logging in
- Creating an SSH endpoint
- Listing sessions that are starting or running
- Starting or stopping a session
- Creating a model
- Building and deploying models
- Listing model builds and model deployments
- Checking the status of a deployment
- Redeploying a model with updated resources
- Viewing the replica logs for a model

Other actions, such as creating a project, require you to use the Cloudera Machine Learning web UI. For information about the available commands, run the following command:

```
cdswctl --help
```

Download and Configure cdswctl

This topic describes how to download the cdswctl CLI client and configure your SSH public key to authenticate CLI access to sessions.

About this task

Before you begin, ensure that the following prerequisites are met:

- You have an SSH public/private key pair for your local machine.
- You have Contributor permissions for an existing project. Alternatively, create a new project you have access to.
- If you want to configure a third-party editor, make sure the Site Administrator has not disabled remote editing for Cloudera Machine Learning.

(Optional) Generate an SSH Public/Private Key

About this task

This task is optional. If you already have an SSH public/private key pair, skip this task. The steps to create an SSH public/private key pair differ based on your operating system. The following instructions are meant to be an example and are written for macOS using ssh-keygen.

Procedure

1. Open Terminal.
2. Run the following command and complete the fields:

```
ssh-keygen -t rsa -f ~/.ssh/id_rsa
```

Keep the following guidelines in mind:

- Make sure that the SSH key you generate meets the requirements for the local IDE you want to use. For example, PyCharm requires the -m PEM option because PyCharm does not support modern (RFC 4716) OpenSSH keys.
- Provide a passphrase when you generate the key pair. Use this passphrase when prompted for the SSH key passphrase.
- Save the SSH key to the default ~/.ssh location.

Download cdswctl and Add an SSH Key

Procedure

1. Open the Cloudera Machine Learning web UI and go to SettingsRemote Editing for your user account.
2. Download cdswctl client for your operating system.
Unpack it, and optionally, you can add it to the *PATH* environment variable on your system.
3. Add your SSH public key to SSH public keys for session access.

Cloudera Machine Learning uses the SSH public key to authenticate your CLI client session, including the SSH endpoint connection to the Cloudera Machine Learning deployment.

Any SSH endpoints that are running when you add an SSH public key must also be restarted.

Initialize an SSH Endpoint

This topic describes how to establish an SSH endpoint for Cloudera Machine Learning.

About this task

Creating an SSH endpoint is also the first step to configuring a remote editor for an ML workspace.

Procedure

1. Create a configuration file at: \$HOME/.cdsw/config.yaml. The contents of config.yaml should be:

```
username: <USERNAME>
url: <ML_WORKSPACE_URL>
auth:
```

```
authtype: 1
basic: null
apikey: <YOUR_API_KEY>
```

To collect the values for these fields, first log in to your CML workspace using SSO:

- **username:** The username with which you are logged into the CML workspace. Found in the top right corner of your ML workspace.
- **url:** The complete URL used to access the CML workspace. For example: `https://ml-<randomly-generated-cluster-name>`
- **apikey:** Go to **User Settings** **API Keys** . Copy the value of the Legacy API Key to this field.

2. Create a local SSH endpoint to Cloudera Machine Learning. Run the following command:

```
cdswctl ssh-endpoint -p <username>/<project_name> [-c <CPU_cores>] [-m <memory_in_GB>] [-g <number_of_GPUs>] [-r <runtime ID> ]
```

The command uses the following defaults for optional parameters:

- CPU cores: 1
- Memory: 1 GB
- GPUs: 0

For example, the following command starts a session for the user milton under the customerchurn project with .5 cores, .75 GB of memory, 0 GPUs, and the Python3 kernel:

```
cdswctl ssh-endpoint -p customerchurn -c 0.5 -m 0.75
```

To create an SSH endpoint in a project owned by another user or a team, for example finance, prepend the username to the project and separate them with a forward slash:

```
cdswctl ssh-endpoint -p finance/customerchurn -c 0.5 -m 0.75
```

This command creates session in the project customerchurn that belongs to the team finance.

Information for the SSH endpoint appears in the output:

```
...
You can SSH to it using
ssh -p <some_port> cdsw@localhost
...
```

3. Open a new command prompt and run the outputted command from the previous step:

```
ssh -p <some_port> cdsw@localhost
```

For example:

```
ssh -p 9750 cdsw@localhost
```

You will be prompted for the passphrase for the SSH key you entered in the ML workspace web UI.

Once you are connected to the endpoint, you are logged in as the cdsw user and can perform actions as though you are accessing the terminal through the web UI.

4. Test the connection.

If you run `ls`, the project files associated with the session you created are shown. If you run `whoami`, the command returns the cdsw user.

5. Leave the SSH endpoint running as long as you want to use a local IDE.

Log into cdswctl

This topic describes how to log into cdswctl.

Procedure

1. Open the Model CLI client.
2. Run the following command while specifying the actual values for the variables:

```
cdswctl login -u <WORKSPACE_URL> -n <USERNAME> -y <LEGACY_API_KEY>
```

where

- *WORKSPACE_URL* is the workspace URL including the protocol (http(s)://domain.com)
- *USERNAME* is your user name on the workspace
- *LEGACY_API_KEY* is the API key that you can obtain from the Cloudera Machine Learning UI. Go to Settings API Keys and copy the Legacy API Key (and not the API Key).

To see more information about the login command parameters, run

```
cdswctl login --help
```

If all goes well, then you'll see "Login succeeded".

Prepare to manage models using the model CLI

Before you can start using the model CLI to automate model deployment or to perform any other tasks, you must install the scikit-learn machine learning library for Python through the Cloudera Machine Learning web UI.

About this task

You must perform this task through the Cloudera Machine Learning web UI.

Procedure

1. Create a new project with Python through the web UI.
Python provides sample files that you can use to create models using CLI.
2. To start a new session, go to the **Sessions** page from the left navigation panel and click new session.
The **Start the new session** page is displayed.
3. On **Start the new session** page, select Python 3 from the Engine Kernel drop-down menu, and click Launch Session.
A new "Untitled Session" is created.
4. From the input prompt, install the scikit-learn machine learning library for Python by running the following command:

```
!pip3 install sklearn
```

5. Open the fit.py file available within your project from the left navigation panel.
You can use the fit.py file to create a fitted model which creates a model.pkl file that you can use to deploy the actual model.
6. Run the fit.py file by clicking Run Run all .
The model.pkl directory is created that you can see within your project on the left navigation pane.
7. Close the session by clicking Stop.

Create a model using the CLI

This topic describes how to create models using the model CLI.

Procedure

1. Open a terminal window and log into cdswctl.
2. Obtain the project ID as described in the following steps:
 - a) Run the following command:

```
cdswctl projects list
```

The project ID, your username, and the project name are displayed. For example:

1: john-smith/petal-length-predictor

- b) Note the project ID, which is a number in front of your project name.

In this case, it is "1".

3. Run the following command while specifying the project name and note the engine image ID:



Note: The following examples are specific to projects configured to use legacy engines and projects configured to use runtimes. Be sure to use the commands appropriate to your project configuration.

For projects configured to use legacy engines:

```
cdswctl engine-images list -p <PROJECT-NAME>
```

For example,

```
cdswctl engine-images list -p john-smith/petal-length-predictor
```

For projects configured to use runtimes:

```
cdswctl runtimes list
```

Depending on your local setup, you may get a more readable output by post-processing the result with the following command:

```
cdswctl runtimes list | python3 -m json.tool
```

For this example you should pick a runtime with a Python kernel and Workbench editor. Depending on your local setup, you may filter the results using the following command:

```
cdswctl runtimes list | jq '.runtimes[] | select((.editor == "Workbench") and (.kernel | contains("Python")))'
```

4. Create a model by using the following command:



Note: The following examples are specific to projects configured to use legacy engines and projects configured to use runtimes. Be sure to use the commands appropriate to your project configuration.

For projects configured to use legacy engines:

```
cdswctl models create
--kernel="python3"
--targetFilePath="predict.py"
--targetFunctionName="predict"
--name="Petal Length Predictor"
--cpuMillicores=1000
```

```
--memoryMb=2000
--description="Model of the Iris dataset"
--replicationType=fixed
--numReplicas=1
--visibility="private"
--autoBuildModel
--autoDeployModel
--projectId=<PROJECT_ID>
--examples='{"request":{"petal_length":1}}'
--engineImageId=<ENGINE_IMAGE_ID_FROM_BEFORE>
```

For projects configured to use runtimes:

```
cdswctl models create
--targetFilePath="predict.py"
--targetFunctionName="predict"
--name="Petal Length Predictor"
--cpuMillicores=1000
--memoryMb=2000
--description="Model of the Iris dataset"
--replicationType=fixed
--numReplicas=1
--visibility="private"
--autoBuildModel
--autoDeployModel
--projectId=<project ID>
--examples='{"request":{"petal_length":1}}'
--runtimeId=<runtime ID obtained above>
```

If the command runs successfully, the system displays the model details in a JSON format.

5. For more information about the models create command parameters, run the following command:

```
cdswctl models create --help
```

Build and deployment commands for models

Models have separate parameters for builds and deployments. When a model is built, an image is created. Whereas, the deployment is the actual instance of the model. You can list model builds and deployment, and monitor their state using from model CLI client (cdswctl).

Listing a model

To list the models, run the following command:

```
cdswctl models list
```

Monitoring the status of the model

To monitor the status of the build for a particular model, use the following command:

```
cdswctl models listBuild --modelId <MODEL_ID> --projectId <PROJECT_ID>
```

You can use the `--latestModelDeployment` flag to get the build for the latest deployment.

Listing a deployment

To list the deployment for a particular model, run the following command:

```
cdswctl models listDeployments --modelId <MODEL_ID>
```

Checking the status of a deployment

To check the status of your deployment, run the following command:

```
cdswctl models listDeployments --statusSet=deployed
```

Following is a list of arguments for the statusSet parameter:

- deployed
- deploying
- failed
- pending
- stopping
- stopped



Note: You can use the parameter more than once in a command to check multiple statuses of your deployed models. For example,

```
cdswctl models listDeployments --statusSet=deployed --statusSet=stopped --statusSet=failed
```

Deploy a new model with updated resources

You can republish a previously-deployed model in a new serving environment with an updated number of replicas or memory/CPU/GPU allocation by providing the model build ID of the model you want to rebuild.

To deploy a new model, use the following command:

```
cdswctl models deploy --modelBuildId=<BUILD_ID> --cpuMillicores=<NUM_OF_CPU_CORES> --memoryMb=<MEMORY_IN_MB> --numReplicas=<NUM_OF_REPLICAS> --replicationType=<REPLICATION_TYPE>
```

For example:

```
cdswctl models deploy --modelBuildId=<BUILD_ID> --cpuMillicores=1200 --memoryMb=2200 --numReplicas=2 --replicationType=fixed
```



Note: You must specify values for all the non-zero resources, even if you do not wish to update their values. For example, in your existing deployment, if you set the cpuMillicores capacity to 1200 and you do not wish to increase or decrease it, you must still specify cpuMillicores=1200 in the command.

View replica logs for a model

When a model is deployed, Cloudera Machine Learning enables you to specify the number of replicas that must be deployed to serve requests. If a replica crashes or fails to come up, you can diagnose it by viewing the logs for every replica using the model CLI.

Procedure

1. Obtain the modelReplicaId by using the following command:

```
cdswctl models listReplicas --modelDeploymentId=<MODEL_DEPLOYMENT_ID>
```

where the *MODEL_DEPLOYMENT_ID* is the ID of a successfully deployed model.

2. To view the replica logs, run the following command:

```
cdswctl models getReplicaLogs --modelDeploymentId=<MODEL_DEPLOYMENT_ID> --  
modelReplicaId=<REPLICA_ID> --streams=stdout
```

For example:

```
cdswctl models getReplicaLogs --modelDeploymentId=2 --modelReplicaId="petal-length-predictor-1-2-6d6496b467-hp6tz" --streams=stdout
```

The valid values for the streams parameter are stdout and stderr.

cdswctl command reference

You can manage your Cloudera Machine Learning Workbench cluster with the CLI client (cdswctl) that exists within the Cloudera Machine Learning Workbench. Running `cdswctl` without any arguments prints a brief description of each command.

Table 1: Model CLI Command Reference

Command	Description and usage
cdswctl login	Enables you to log into the model CLI client
cdswctl projects list	Lists the projects
cdswctl models create	Creates a model with the specified parameters
cdswctl models list	Lists all models You can refine the search by specifying the modelId
cdswctl models listBuild	Lists the builds for a model You can monitor the status of the build by specifying the modelId and the projectId
cdswctl models listDeployments	List the deployments for a model You can refine the search by specifying the modelId Use the statusSet parameter to check the status of the model being deployed
cdswctl models deploy	Deploys a model with the specified parameters
cdswctl models listReplicas	Enables you to view the list of model replicas You also need this information to obtain replica logs
cdswctl models getReplicaLogs	Enables you to view the logs for a model replica

Command	Description and usage
<code>cdswctl models restart</code>	Restarts a model Usage: <pre>cdswctl models restart --modelDeployment Id=<DEPLOYMENT_ID></pre> Note: Running this command does not change the resources if you previously ran the <code>cdswctl models update</code> command
<code>cdswctl models update</code>	Changes the name, description, or visibility of the model To change a model's resources, use the <code>cdswctl models deploy</code> command
<code>cdswctl models delete</code>	Deletes a model Usage: <pre>cdswctl models delete --id=<MODEL_ID></pre>