
MLBench Helm Documentation

MLBench development team

Dec 07, 2020

MLBENCH

1	mlbench: Distributed Machine Learning Benchmark Helm Chart	1
1.1	Chart Details	1
1.2	Prerequisites	1
1.3	Installing the Chart	1
1.4	Configuration	1
2	Indices and tables	5

MLBENCH: DISTRIBUTED MACHINE LEARNING BENCHMARK HELM CHART

The Helm Chart is used to deploy MLBench to a Kubernetes cluster. The source can be found in the [Helm repository](#).

1.1 Chart Details

This Chart deploys the following:

- 1 x MLBench Dashboard/Master Node with Port 80 exposed (Dashboard and REST API)
- 2 x MLBench Worker Nodes, connecting to the REST API of the Dashboard, with Port 22 (SSH) exposed inside the cluster

1.2 Prerequisites

- [Helm](#)
- Helm needs to be set up with service-account with `cluster-admin` rights:

1.3 Installing the Chart

To install the chart with the release name `my-release` and values file `values.yaml`:

```
$ git clone https://github.com/mlbench/mlbench-helm.git
$ cd mlbench-helm
$ helm install -f values.yaml --name my-release ./
```

1.4 Configuration

The following tables list configurable parameters of the MLBench chart and their default values. Entries without default values are mandatory.

Specify each parameter using the `--set key=value[,key=value]` argument to `helm install`.

Alternatively, a YAML file that specifies the values for the parameters can be provided while installing the chart. For example,

```
$ helm install --name my-release -f values.yaml stable/dask
```

Tip: You can use the default `values.yaml`

1.4.1 Dashboard/Master Node

Parameter	Description	Default
<code>master.enabled</code>	Whether to deploy the master node or not	<code>true</code>
<code>master.name</code>	The name of the node	<code>master</code>
<code>master.image.repository</code>	The Docker Registry to use	<code>mlbench/mlbench_master</code>
<code>master.image.tag</code>	The tag of the image to use	<code>latest</code>
<code>master.image.pullPolicy</code>	The K8s imagePullPolicy	<code>Always</code>
<code>master.service.type</code>	The K8s service type	<code>NodePort</code>
<code>master.service.port</code>	The port to expose in K8s	<code>80</code>

1.4.2 Worker Nodes

Parameter	Description	Default
<code>worker.sshKey.id_rsa</code>	The SSH Private Key	(not shown)
<code>worker.sshKey.id_rsa</code>	The SSH Public Key	(not shown)

1.4.3 Hardware Limits

Important: These values are mandatory.

Parameter	Description	Default
<code>limits.workers</code>	The maximum number of workers that can be commissioned	
<code>limits.cpu</code>	The maximum number of cpu cores that can be commissioned per worker	
<code>limits.gpu</code>	The maximum number of GPUs that can be commissioned per worker	

1.4.4 Google Cloud Storage

If deploying to the Google Cloud, use these to set the shared storage for workers.

Parameter	Description	Default
<code>gcePersistentDisk.enabled</code>	Whether to use Google Cloud Storage	false
<code>gcePersistentDisk.pdName</code>	The name of the persistent Disk to use	

1.4.5 Weave

Settings concerning [WeaveNet](#), a Networking Solution between K8s pods. Necessary in some cases where the SourceIP of a Pod defaults to the IP of the Node it's on, which can cause troubles with MPI execution.

Parameter	Description	Default
<code>weave.enabled</code>	Whether to use WeaveNet	false

1.4.6 NVIDIA Device Plugin

Needed to support NVIDIA GPUs in workers (unless already provided by your K8s provider).

Parameter	Description	Default
<code>nvidiaDevicePlugin.enabled</code>	Whether to use the NVIDIA Device Plugin	false

INDICES AND TABLES

- `genindex`
- `modindex`
- `search`