



Red Hat OpenShift on DGX

User Guide

Table of Contents

- Chapter 1. Introduction..... 1
- Chapter 2. Installing RHEL OpenShift..... 3
 - 2.1. Prerequisites..... 3
 - 2.2. Installing Red Hat CoreOS..... 3
 - 2.3. Installing the NVIDIA GPU Operator..... 4
 - 2.4. Installing NVSM..... 4
- Chapter 3. Using NVSM..... 6
 - 3.1. Retrieving Health Information..... 6
- Chapter 4. Known Issues..... 8

Chapter 1. Introduction

This chapter provides a brief introduction to OpenShift for DGX.

Red Hat OpenShift is an Enterprise-grade container-management solution based on Kubernetes for automating deployment, scaling, and management of containerized applications. It is developed and supported by Red Hat and includes additional security features and tooling for managing complex infrastructures on-premises as well as in hybrid cloud installations.

Red Hat OpenShift 4 is a major release upgrade from version 3 incorporating many technologies from the acquisition of CoreOS. It follows a new paradigm where systems are always reimaged with the latest version with only minimal provisioning. At its core are the immutable Red Hat CoreOS (RHCOS) system images based on Red Hat Enterprise Linux 8. All additional software, drivers, and configuration are ephemeral and provided through kubernetes primitives, such as containers, deployments, and operators. This includes the NVIDIA GPU operator for supporting NVIDIA GPUs and the NVIDIA Network Operator for the ConnectX network interfaces.

While OpenShift 4 still supports Red Hat Enterprise Linux 7 and 8 on the worker nodes, customers are advised to move to the newer Red Hat CoreOS deployments for improved supportability. Refer to the corresponding installation instructions for Red Hat Enterprise Linux on DGX and the OpenShift documentation when you are not planning to use Red Hat CoreOS on DGX.

This user guide provides additional information for installing and configuring OpenShift 4 with Red Hat CoreOS on clusters incorporating DGX worker nodes. It should be seen as a companion document to the official Red Hat OpenShift documentation. The following chapters describe additional configuration steps and best practices that are specific to NVIDIA DGX™ systems. Refer to the [OpenShift Container Platform Documentation](#) for generic information about OpenShift and installation instructions.

Customer Support

Customer support for running OpenShift on DGX systems is provided by Red Hat for OpenShift and NVIDIA for the DGX platform and drivers. For CoreOS and OpenShift support, visit the Red Hat Enterprise Support website: <https://www.redhat.com/en/services/support>. For DGX hardware, firmware / drivers, or NGC application issues, visit the NVIDIA Enterprise Support website: <https://www.nvidia.com/en-us/support/enterprise/>

Additional Documentation

Refer to the following documents for additional Information:

- ▶ [Red Hat OpenShift Product page](#)
- ▶ [OpenShift Container Platform Documentation](#)
- ▶ [GPU Operator on OpenShift](#)
- ▶ [NVIDIA System Management Documentation](#)

Chapter 2. Installing RHEL OpenShift

This chapter describes additional steps that are required or recommended to install OpenShift and CoreOS on DGX worker nodes.

1. [Installing Red Hat CoreOS](#).
2. [Installing the NFD Operator and NVIDIA GPU Operator](#).
3. [Installing and Using NVSM](#).

2.1. Prerequisites

Here are the prerequisites for using RHEL OpenShift.

► **Red Hat Subscription**

Installing and running OpenShift requires a Red Hat account and additional subscriptions. Please refer to [Red Hat OpenShift](#) for more information.

► **OpenShift 4.9.9 or later**

Support for DGX has been added to the GPU operator in version 1.9. The operator requires OpenShift 4.9.9 or later.

► **Helm Management Tool**

[NVIDIA System Management \(NVSM\)](#) uses [Helm](#) for installing NVSM on DGX worker nodes. NVSM is a software framework for collecting health status information and helps users analyze hardware and software issues. Refer to [Installing Helm](#) for instructions installing the Helm tool on the system you use to interact with the OpenShift cluster.

2.2. Installing Red Hat CoreOS

Installing OpenShift and Red Hat CoreOS on clusters with DGX worker nodes is the same as installing on other systems.

Follow the instructions described in [Installing a cluster on bare metal](#) or other methods to create an OpenShift cluster and to install Red Hat CoreOS on the nodes.

2.3. Installing the NVIDIA GPU Operator

The NVIDIA GPU Operator is required to manage and allocate GPU resources to workloads. It uses the operator framework within Kubernetes to automate the management of all NVIDIA software components needed to provision the GPU. These components include the NVIDIA drivers (to enable CUDA), Kubernetes device plugin for GPUs, the NVIDIA Container Toolkit, and DCGM based monitoring.

To install the Node Feature Discovery (NFD) Operator and NVIDIA GPU Operator, follow the instructions in the [GPU Operator on OpenShift](#) user guide. The NFD Operator manages the detection of hardware features and configurations in an OpenShift Container Platform cluster by labeling the nodes with hardware-specific information. These labels are required by the GPU Operator to identify machines with a valid GPU.

2.4. Installing NVSM

NVIDIA System Management (NVSM) is a software framework for monitoring NVIDIA DGX nodes in a data center. It includes active health monitoring, system alerts, and log generation, and also supports a stand-alone mode from the command line to get a quick health report of the system. Running NVSM is typically requested by the NVIDIA Enterprise Support team to resolve a reported problem.

NVSM can be deployed on the DGX nodes with the NVIDIA System Management NGC Container. It allows users to execute the NVSM tool remotely and on-demand in the containers that are deployed to the DGX nodes.

The installation uses the NVSM Helm Chart to create the necessary resources on the cluster. The deployment is limited to systems that are manually labeled with `nvidia.com/gpu.nvsm.deploy=true`.

To deploy NVSM on the DGX worker nodes:

1. **Optional:** Issue the following command to get a list of all DGX nodes in the cluster.

```
oc get nodes --show-labels|grep nvidia.com/gpu.machine=.*DGX[^,]*
```

2. Set the `nvidia.com/gpu.nvsm.deploy=true` flag on the DGX worker nodes on which you want to deploy NVSM (replace `WORKER1` and so on with the actual name of the nodes).

```
oc label node/WORKER1 nvidia.com/gpu.nvsm.deploy=true
oc label node/WORKER2 nvidia.com/gpu.nvsm.deploy=true
...
```

3. Get the Helm chart for deploying NVSM on the cluster.

```
helm fetch https://helm.ngc.nvidia.com/nvidia/cloud-native/charts/nvsm-1.0.1.tgz
--username='${oauthtoken}' --password=<NGC_API_KEY>
```

4. Ensure the file is in your local directory.

```
ls ./nvsm-1.0.1.tgz
```

5. To ensure the default settings are correct for your installation, inspect the contents of `values.yaml` file in the above tar file.

If the settings are not correct, update the `values.yaml` file as per the cluster configuration.

```
helm install --
```

6. Deploy NVSM to the cluster.

The cluster installs the container on all nodes that have been labeled in the previous steps. The following command creates the namespace `nvidia-nvsm` namespace and deploys the resource in the namespace:

```
helm install --set platform.openshift=true --create-namespace -n nvidia-nvsm
nvidia-nvsm ./nvsm-1.0.1.tgz
```

7. Validate that NVSM has been deployed on all selected DGX nodes.

You should see an `nvidia-nvsm-xxxx` pod instance for each node:

```
oc get pods -n nvidia-nvsm -o wide
```

NAME	READY	STATUS	RESTARTS	AGE	IP	NODE	...
nvidia-nvsm-d9d9t	1/1	Running	1	8h	10.128.2.11	worker-0	...
nvidia-nvsm-tt8g5	1/1	Running	1	8h	10.131.0.11	worker-1	...

NVSM is now installed and can be run remotely using `oc exec`.

Chapter 3. Using NVSM

For a maintenance task, or to complete a health analysis, you can now run NVSM remotely inside the deployed containers on one of the DGX worker nodes.

Here is the general NVSM workflow:

1. List all NVSM pod instances and the corresponding worker nodes to find the pod name associated with a specific DGX.

```
oc get pods -n nvidia-nvsm -o wide
NAME                READY   STATUS    RESTARTS   AGE   IP            NODE       ...
nvidia-nvsm-d9d9t   1/1     Running   1           8h    10.128.2.11   worker-0   ...
nvidia-nvsm-tt8g5   1/1     Running   1           8h    10.131.0.11   worker-1   ...
```

2. Use the `oc exec` command to start an interactive shell in the container that is running on that system.

```
oc exec -it <pod-name> -n nvidia-nvsm -- /bin/bash
```

3. You can now use one of the following main NVSM commands.



Note: When you execute NVSM, it can take a couple of minutes to collect system information.

- ▶ To print the software and firmware versions of the DGX system:

```
nvsm show version
```

- ▶ To provide a summary of the system health:

```
nvsm show health
```

- ▶ To create a snapshot of the system components for offline analysis and diagnosis:

```
nvsm dump health
```

This command generates a tar-file in the `/tmp` directory. This command generates a tar-file in the `/tmp` directory. (see [Retrieving Health Information](#) for more information).

4. Exit the interactive shell:

```
nvsm dump health
```

3.1. Retrieving Health Information

This section describes the steps to generate and retrieve health information to debug a system issue offline or when requested by the NVIDIA Enterprise Support organization.

1. List all NVSM pod instances and the corresponding worker nodes to find the pod name that is associated with a DGX system.

```
oc get pods -n nvidia-nvsm -o wide
```

NAME	READY	STATUS	RESTARTS	AGE	IP	NODE	...
nvidia-nvsm-d9d9t	1/1	Running	1	8h	10.128.2.11	worker-0	...
nvidia-nvsm-tt8q5	1/1	Running	1	8h	10.131.0.11	worker-1	...

2. Start an interactive shell in the NVSM pod of the corresponding DGX worker node.

POD is the name of the pd from the list in the table in step 1.

```
oc exec -it <pod-name> -n nvidia-nvsm -- /bin/bash
```

Here is an example:

```
oc exec -it nvidia-nvsm-d9d9t -n nvidia-nvsm -- /bin/bash
```

3. Create the NVSM snapshot file of all system components for offline analysis and diagnostics.

The file is created in the `/tmp` directory in the container.

```
nvsm dump health
```

Unable to find NVML library, Aborting.

Health dump started

This command will collect system configuration and diagnostic information to help diagnose the system.

The output may contain data considered sensitive and should be reviewed before sending to any third party.

The output of this command is written to: /tmp/nvsm-health-nvidia-nvsm-d9d9t-20211211170039.tar.xz

4. Exit the container.

```
exit
```

5. Copy the snapshot file out of the container to the local host for further analysis or to send it to NVIDIA Enterprise Support.

```
oc cp nvidia-nvsm/POD:tmp/<snapshot-file> <target-file>
```

snapshot-file refers to the name of the generated file from step 3, and *target-file* refers to the name of the file on the local host. You need to replace these variables with actual names.

6. Delete the generated snapshot file in the NVSM pod.

```
oc exec -it POD -n nvidia-nvsm -- rm /tmp/<snapshot-file>
```

The generated file can now be used to debug and analyze system issues, or you can send the file to NVIDIA Enterprise Support.

Chapter 4. Known Issues

This section provides a list of known issues for OpenShift.

March 8, 2022

- ▶ OpenShift and Red Hat CoreOS do not currently support upgrading DGX component firmware using the NVIDIA Firmware Update container method.

On the DGX A100, you can use the [DGX A100 Firmware Update ISO](#) method. For other DGX systems, you need to install Red Hat Enterprise Linux or DGX OS to upgrade the firmware before you can return the system to the OpenShift cluster.

- ▶ The current version of the [NVIDIA Network Operator](#) provides only limited support for DGX systems, and should be considered experimental.

Notice

This document is provided for information purposes only and shall not be regarded as a warranty of a certain functionality, condition, or quality of a product. NVIDIA Corporation ("NVIDIA") makes no representations or warranties, expressed or implied, as to the accuracy or completeness of the information contained in this document and assumes no responsibility for any errors contained herein. NVIDIA shall have no liability for the consequences or use of such information or for any infringement of patents or other rights of third parties that may result from its use. This document is not a commitment to develop, release, or deliver any Material (defined below), code, or functionality.

NVIDIA reserves the right to make corrections, modifications, enhancements, improvements, and any other changes to this document, at any time without notice.

Customer should obtain the latest relevant information before placing orders and should verify that such information is current and complete.

NVIDIA products are sold subject to the NVIDIA standard terms and conditions of sale supplied at the time of order acknowledgement, unless otherwise agreed in an individual sales agreement signed by authorized representatives of NVIDIA and customer ("Terms of Sale"). NVIDIA hereby expressly objects to applying any customer general terms and conditions with regards to the purchase of the NVIDIA product referenced in this document. No contractual obligations are formed either directly or indirectly by this document.

NVIDIA products are not designed, authorized, or warranted to be suitable for use in medical, military, aircraft, space, or life support equipment, nor in applications where failure or malfunction of the NVIDIA product can reasonably be expected to result in personal injury, death, or property or environmental damage. NVIDIA accepts no liability for inclusion and/or use of NVIDIA products in such equipment or applications and therefore such inclusion and/or use is at customer's own risk.

NVIDIA makes no representation or warranty that products based on this document will be suitable for any specified use. Testing of all parameters of each product is not necessarily performed by NVIDIA. It is customer's sole responsibility to evaluate and determine the applicability of any information contained in this document, ensure the product is suitable and fit for the application planned by customer, and perform the necessary testing for the application in order to avoid a default of the application or the product. Weaknesses in customer's product designs may affect the quality and reliability of the NVIDIA product and may result in additional or different conditions and/or requirements beyond those contained in this document. NVIDIA accepts no liability related to any default, damage, costs, or problem which may be based on or attributable to: (i) the use of the NVIDIA product in any manner that is contrary to this document or (ii) customer product designs.

No license, either expressed or implied, is granted under any NVIDIA patent right, copyright, or other NVIDIA intellectual property right under this document. Information published by NVIDIA regarding third-party products or services does not constitute a license from NVIDIA to use such products or services or a warranty or endorsement thereof. Use of such information may require a license from a third party under the patents or other intellectual property rights of the third party, or a license from NVIDIA under the patents or other intellectual property rights of NVIDIA.

Reproduction of information in this document is permissible only if approved in advance by NVIDIA in writing, reproduced without alteration and in full compliance with all applicable export laws and regulations, and accompanied by all associated conditions, limitations, and notices.

THIS DOCUMENT AND ALL NVIDIA DESIGN SPECIFICATIONS, REFERENCE BOARDS, FILES, DRAWINGS, DIAGNOSTICS, LISTS, AND OTHER DOCUMENTS (TOGETHER AND SEPARATELY, "MATERIALS") ARE BEING PROVIDED "AS IS." NVIDIA MAKES NO WARRANTIES, EXPRESSED, IMPLIED, STATUTORY, OR OTHERWISE WITH RESPECT TO THE MATERIALS, AND EXPRESSLY DISCLAIMS ALL IMPLIED WARRANTIES OF NON-INFRINGEMENT, MERCHANTABILITY, AND FITNESS FOR A PARTICULAR PURPOSE. TO THE EXTENT NOT PROHIBITED BY LAW, IN NO EVENT WILL NVIDIA BE LIABLE FOR ANY DAMAGES, INCLUDING WITHOUT LIMITATION ANY DIRECT, INDIRECT, SPECIAL, INCIDENTAL, PUNITIVE, OR CONSEQUENTIAL DAMAGES, HOWEVER CAUSED AND REGARDLESS OF THE THEORY OF LIABILITY, ARISING OUT OF ANY USE OF THIS DOCUMENT, EVEN IF NVIDIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGES. Notwithstanding any damages that customer might incur for any reason whatsoever, NVIDIA's aggregate and cumulative liability towards customer for the products described herein shall be limited in accordance with the Terms of Sale for the product.

VESA DisplayPort

DisplayPort and DisplayPort Compliance Logo, DisplayPort Compliance Logo for Dual-mode Sources, and DisplayPort Compliance Logo for Active Cables are trademarks owned by the Video Electronics Standards Association in the United States and other countries.

HDMI

HDMI, the HDMI logo, and High-Definition Multimedia Interface are trademarks or registered trademarks of HDMI Licensing LLC.

OpenCL

OpenCL is a trademark of Apple Inc. used under license to the Khronos Group Inc.

Trademarks

NVIDIA, the NVIDIA logo, DGX, DGX-1, DGX-2, DGX A100, DGX Station, and DGX Station A100 are trademarks and/or registered trademarks of NVIDIA Corporation in the United States and other countries. Other company and product names may be trademarks of the respective companies with which they are associated.

Copyright

© 2022 NVIDIA Corporation. All rights reserved.

