

Intel® Ethernet Adapters and Devices User Guide

Overview

Welcome to the User Guide for Intel® Ethernet Adapters and devices. This guide covers hardware and software installation, setup procedures, and troubleshooting tips for Intel network adapters, connections, and other devices.

Intended Audience

This document is intended for information technology professionals with a high level of knowledge, experience, and competency in Ethernet networking technology.

Supported Devices

Supported 40 Gigabit Network Adapters

- Intel® Ethernet 40G 2P XL710 QSFP+ rNDC
- Intel® Ethernet Converged Network Adapter XL710-Q2



NOTES:

- Devices based on the Intel Ethernet Controller XL710 (4x10 GbE, 1x40 GbE, 2x40 GbE) have an expected total throughput for the entire device of 40 Gb/s in each direction.
- The first port of Intel® Ethernet Controller 700 Series based adapters will display the correct branding string. All other ports on the same device will display a generic branding string.
- In order for an Intel® Ethernet Controller 700 Series based adapter to reach its full potential, you must install it in a PCIe Gen3 x8 slot. Installing it in a shorter slot, or a Gen2 or Gen1 slot, will limit the throughput of the adapter.

Supported 25 Gigabit Network Adapters

- Intel® Ethernet 25G 2P XXV710 Adapter
- Intel® Ethernet 25G 2P XXV710 Mezz
- Intel® Ethernet 25G 2P E810-XXV OCP
- Intel® Ethernet 25G 2P E810-XXV Adapter



NOTES:

- Devices based on the Intel Ethernet Controller XXV710 (2x25 GbE) have a total hardware throughput limit for the entire device of ~96-97% of dual-port 25 GbE line rate in each direction for IPv4 TCP large packets (>1518 bytes) with an MTU size of 1500 bytes. For example, the total payload throughput is limited to ~45.5 Gb/s in each direction. Thus, while single port 25 GbE throughput is not impacted, total simultaneous dual port 25 GbE throughput is expected to be slightly lower than line rate.

The first port of Intel® Ethernet Controller 700 Series based adapters will display the correct branding string. All other ports on the same device will display a generic branding string.

Supported 10 Gigabit Network Adapters

- Intel® Ethernet 10G 2P X520 Adapter
- Intel® Ethernet 10G X520 LOM
- Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz
- Intel® Ethernet 10G 2P X540-t Adapter
- Intel® Ethernet 10G 2P X550-t Adapter
- Intel® Ethernet 10G 4P X550 rNDC
- Intel® Ethernet 10G 4P X550/I350 rNDC
- Intel® Ethernet 10G 4P X540/I350 rNDC
- Intel® Ethernet 10G 4P X520/I350 rNDC
- Intel® Ethernet 10G 2P X520-k bNDC

- Intel® Ethernet 10G 4P X710-k bNDC
- Intel® Ethernet 10G 2P X710-k bNDC
- Intel® Ethernet 10G X710-k bNDC
- Intel® Ethernet Converged Network Adapter X710
- Intel® Ethernet Converged Network Adapter X710-T
- Intel® Ethernet 10G 4P X710/I350 rNDC
- Intel® Ethernet 10G 4P X710 SFP+ rNDC
- Intel® Ethernet 10G X710 rNDC
- Intel® Ethernet Server Adapter X710-DA2 for OCP
- Intel® Ethernet 10G 2P X710 OCP
- Intel® Ethernet 10G 4P X710 OCP
- Intel® Ethernet 10G 2P X710-T2L-t OCP
- Intel® Ethernet 10G 4P X710-T4L-t OCP
- Intel® Ethernet 10G 2P X710-T2L-t Adapter
- Intel® Ethernet 10G 4P X710-T4L-t Adapter

**NOTE:**

The first port of Intel® Ethernet Controller 700 Series based adapters will display the correct branding string. All other ports on the same device will display a generic branding string.

Supported Gigabit Network Adapters and Devices

- Intel® Gigabit 2P I350-t Adapter
- Intel® Gigabit 4P I350-t Adapter
- Intel® Ethernet 1G 4P I350-t OCP
- Intel® Gigabit 4P X550/I350 rNDC
- Intel® Gigabit 4P I350-t rNDC
- Intel® Gigabit 4P X540/I350 rNDC
- Intel® Gigabit 4P X520/I350 rNDC
- Intel® Gigabit 4P I350-t Mezz
- Intel® Gigabit 4P X710/I350 rNDC
- Intel® Gigabit 4P I350 bNDC
- Intel® Ethernet Connection I354 1.0 GbE Backplane
- Intel® Gigabit 2P I350-t LOM
- Intel® Gigabit I350-t LOM
- Intel® Gigabit 2P I350 LOM

Supported Operating Systems

The drivers in this release have been tested with the following operating systems. Additional OSs may function with our drivers but are not tested.

- Microsoft Windows Server 2019, Version 1809
- Microsoft Windows Server 2016
- VMWare® ESXi® 7.0 U1
- VMWare ESXi 6.7 U3
- Red Hat® Enterprise Linux® (RHEL) 8.3
- Red Hat® Enterprise Linux® (RHEL) 8.2
- Red Hat® Enterprise Linux® (RHEL) 7.9
- Novell® SUSE® Linux Enterprise Server (SLES) 15 SP2

Installation

This chapter covers how to install Intel® Ethernet adapters, drivers, and other software.

At a high level, installation involves the following steps, which are covered in more detail later in this chapter.

If you are installing a network adapter, follow this procedure from step 1.

If you are upgrading the driver software, start with step 4.



NOTE: If you update the firmware, you must update the driver software to the same family version.

1. Review [system requirements](#).
2. [Insert the PCI Express Adapter, Mezzanine Card, or Network Daughter Card](#) into your server.
3. Carefully connect the network [copper cable\(s\)](#), fiber cable(s), or [direct attach cables](#)
4. Install the [network drivers and other software](#).
5. [Test the adapter](#).

Hardware Compatibility

Before installing the adapter, check your system for the following:

- The latest BIOS for your system
- One open PCI Express slot (see the [specifications of your card](#) for slot compatibility)

Installing the Adapter

Select the Correct Slot

One open PCI-Express slot, x4, x8, or x16, depending on your adapter.



NOTE: Some systems have physical x8 PCI Express slots that actually only support lower speeds. Please check your system manual to identify the slot.



NOTE: For information on identifying PCI Express slots that support your adapters, see your Dell EMC system guide.

Insert the Adapter into the Computer

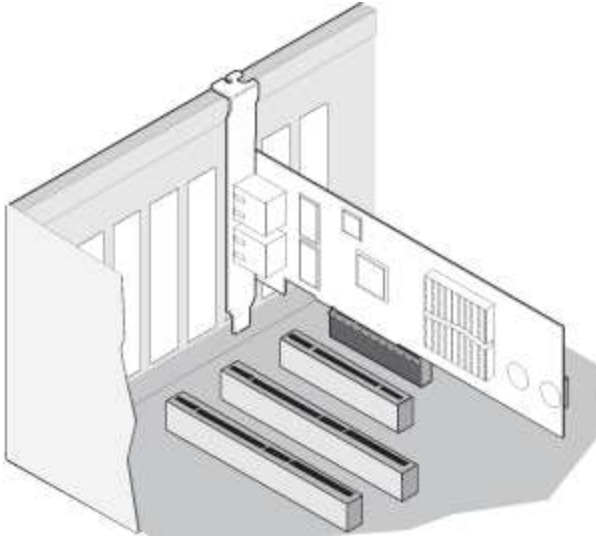
1. If your computer supports PCI Hot Plug, see your computer documentation for special installation instructions.
2. Turn off and unplug your computer. Then remove the cover.



CAUTION: Turn off and unplug the power before removing the computer's cover. Failure to do so could endanger you and may damage the adapter or computer.

3. Remove the cover bracket from an available slot.

4. Insert the adapter, pushing it into the slot until the adapter is firmly seated. You can install a smaller PCI Express adapter in a larger PCI Express slot.



CAUTION: Some PCI Express adapters may have a short connector, making them more fragile than PCI adapters. Excessive force could break the connector. Use caution when pressing the board in the slot.

5. Secure the adapter bracket with a screw, if required.
6. Replace the computer cover and plug in the power cord.
7. Power on the computer.

Install a Mezzanine Card in the Blade Server

See your server documentation for detailed instructions on how to install a Mezzanine card.

1. Turn off the blade server and pull it out of the chassis, then remove its cover.



CAUTION: Failure to turn off the blade server could endanger you and may damage the card or server.

2. Lift the locking lever and insert the card in an available, compatible mezzanine card socket. Push the card into the socket until it is firmly seated.



NOTE: A switch or pass-through module must be present on the same fabric as the card in the chassis to provide a physical connection. For example, if the mezzanine card is inserted in fabric B, a switch must also be present in fabric B of the chassis.

3. Repeat steps 2 for each card you want to install.
4. Lower the locking lever until it clicks into place over the card or cards.
5. Replace the blade server cover and put the blade back into the server chassis.
6. Turn the power on.

Install a Network Daughter Card in a Server

See your server documentation for detailed instructions on how to install a bNDC or rNDC.

1. Turn off the server and then remove its cover.



CAUTION: Failure to turn off the server could endanger you and may damage the card or server.

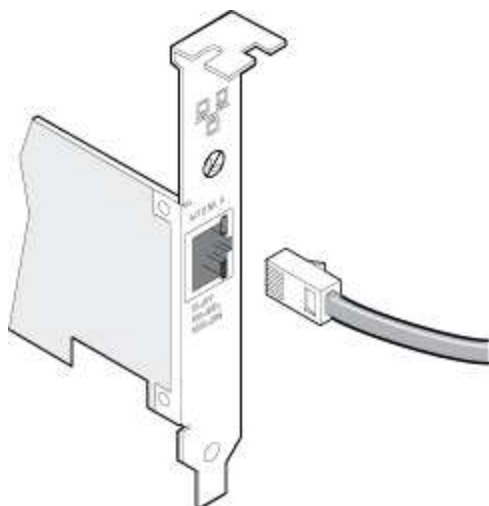
2. Locate the Network Daughter Card connector in your server. See your server's documentation for details.
3. Press the Network Daughter Card into the connector.
4. Tighten the screws on the Network Daughter Card to secure it into place.
5. Replace the server's cover.

Connecting Network Cables

Connect the appropriate network cable, as described in the following sections.

Connect the RJ-45 Network Cable

Connect the RJ-45 network cable as shown:



The following table shows the maximum lengths for each cable type at a given transmission speed.

	Category 5	Category 6	Category 6a	Category 7
1 Gbps	100m	100m	100m	100m
10 Gbps	NA	55m	100m	100m
25 Gbps	NA	NA	NA	50m
40 Gbps	NA	NA	NA	50m



CAUTION: If using less than 4-pair cabling, you must manually configure the speed and duplex setting of the adapter and the link partner. In addition, with 2- and 3-pair cabling the adapter can only achieve speeds of up to 100Mbps.

In all cases:

- The adapter must be connected to a compatible link partner, preferably set to auto-negotiate speed and duplex for Intel gigabit adapters.
- Intel Gigabit and 10 Gigabit Server Adapters using copper connections automatically accommodate either MDI or MDI-X connections. The auto-MDI-X feature of Intel gigabit copper adapters allows you to directly connect two adapters without using a cross-over cable.

Supported SFP+, SFP28, QSFP+, and QSFP28 Modules

Intel® Ethernet Server Adapters only support Intel optics and/or all passive and active limiting direct attach cables that comply with SFF-8431 v4.1 and SFF-8472 v10.4 specifications.

SR transceiver cabling specifications

Laser wavelength: 850 nanometer (not visible)

Connector type: LC or SC

Cable type: Multi-mode fiber with 62.5µm core diameter

- 1 Gbps maximum cable length: 275 meters
- 10 Gbps (and faster) maximum cable length: 33 meters

Cable type: Multi-mode fiber with 50µm core diameter

- 1 Gbps maximum cable length: 550 meters
- 10 Gbps (and faster) maximum cable length: 300 meters

LR transceiver cabling specifications

Laser wavelength: 1310 nanometer (not visible)

Connector type: LC

Cable type: Single-mode fiber with 9.0µm core diameter

- Maximum cable length: 10 kilometers

Most Intel® Ethernet Server Adapters support the following modules:



NOTE: Intel® Ethernet 710 Series based devices do not support third party modules.

Supplier	Type	Part Numbers	Supported Adapters
Dell EMC	Dual Rate 1G/10G SFP+ SR (bailed)	C5RNH ¹ , WTRD1 ^{1,2} , XYD50, Y3KJN ²	X520, X710 ³ , XXV710, E810-XXV
Dell EMC	Dual Rate 10G/25G SFP28	M14MK	XXV710, E810-XXV
Dell EMC	QSFP+ F10 Passive Octopus (QSFP+ to 4xSFP+)	27GG5, JNPF8, P4YPY, P8T4W, TCPM2	X520, X710 ³ , E810-XXV
Dell EMC	SFP+ to 1000BASE-T Transceiver	8T47V, XTY28	X710 ³ , E810-XXV
Dell EMC	SFP+ LR Optic	60F4J, RN84N	X710 ³ , E810-XXV
Dell EMC	Active Optical Cable (AOC)	1DXKP, K0T7R, MT7R2, P9GND, T1KCN, W5G04, YJF03	X710 ³ , XXV710, E810-XXV
Dell EMC	SFP28 Optic	68X15, HHHHC ² , 0YR96, P7D7R ² , W4GPP	XXV710, E810-XXV
Dell EMC	SFP+ F10 Passive	358VV, 53HVN, 5CWK6, C6Y7M, V250M	XXV710, E810-XXV
Dell EMC	SFP28 Passive	2JVDD, 9X8JP, D0R73, VXFJY	XXV710, E810-XXV
Dell EMC	SFP28 Active	3YWG7, 5CMT2, RCPV5, X5DH4	XXV710, E810-XXV
Dell EMC	QSFP28 F10 Passive Octopus (QSFP+ to 4xSFP28)	26FN3, 7R9N9, YFNDD	XXV710, E810-XXV

Dell EMC	QSFP28 Passive Breakout Cables	7VN5T, 8R4VM, D9YM8	XXV710, E810-XXV
Dell EMC	TRIPLE RATE 1G/10G/40G QSFP+ SR (bailed) (1G and 10G not supported on XL710)	5NP8R, 7TCDN, 9GCCD, FC6KV, J90VN, NWGTV, V492M	XL710

¹Not supported on adapters based on the Intel® X520 controller.

²Not supported on adapters based on the Intel® E810-XXV controller.

³The Intel® Ethernet Server Adapter X710-DA2 for OCP only supports modules listed in the table below.

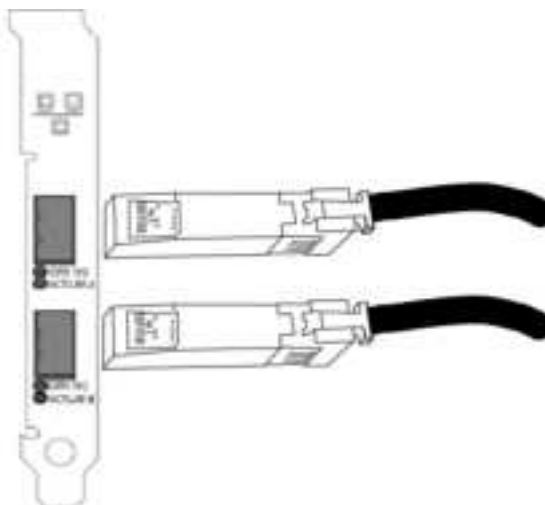
The Intel® Ethernet Server Adapter X710-DA2 for OCP only supports the following modules:

Supplier	Type	Part Numbers
Dell EMC	SFP+ SR High Temp Optics	N8TDR
Dell EMC	QSFP+ F10 Passive Octopus	27GG5, P8T4W, TCPM2

THIRD PARTY OPTIC MODULES AND CABLES REFERRED TO ABOVE ARE LISTED ONLY FOR THE PURPOSE OF HIGHLIGHTING THIRD PARTY SPECIFICATIONS AND POTENTIAL COMPATIBILITY, AND ARE NOT RECOMMENDATIONS OR ENDORSEMENT OR SPONSORSHIP OF ANY THIRD PARTY'S PRODUCT BY INTEL. INTEL IS NOT ENDORSING OR PROMOTING PRODUCTS MADE BY ANY THIRD PARTY AND THE THIRD PARTY REFERENCE IS PROVIDED ONLY TO SHARE INFORMATION REGARDING CERTAIN OPTIC MODULES AND CABLES WITH THE ABOVE SPECIFICATIONS. THERE MAY BE OTHER MANUFACTURERS OR SUPPLIERS, PRODUCING OR SUPPLYING OPTIC MODULES AND CABLES WITH SIMILAR OR MATCHING DESCRIPTIONS. CUSTOMERS MUST USE THEIR OWN DISCRETION AND DILIGENCE TO PURCHASE OPTIC MODULES AND CABLES FROM ANY THIRD PARTY OF THEIR CHOICE. CUSTOMERS ARE SOLELY RESPONSIBLE FOR ASSESSING THE SUITABILITY OF THE PRODUCT AND/OR DEVICES AND FOR THE SELECTION OF THE VENDOR FOR PURCHASING ANY PRODUCT. THE OPTIC MODULES AND CABLES REFERRED TO ABOVE ARE NOT WARRANTED OR SUPPORTED BY INTEL. INTEL ASSUMES NO LIABILITY WHATSOEVER, AND INTEL DISCLAIMS ANY EXPRESS OR IMPLIED WARRANTY, RELATING TO SALE AND/OR USE OF SUCH THIRD PARTY PRODUCTS OR SELECTION OF VENDOR BY CUSTOMERS.

Connect the Direct Attach Cable

Insert the Direct Attach network cable as shown below.



Type of cabling:

- 40 Gigabit Ethernet over SFP+ Direct Attached Cable (Twinaxial)
 - Length is 7 meters max.
- 25 Gigabit Ethernet over SFP28 Direct Attached Cable (Twinaxial)
 - Length is 5 meters max.
 - For optimal performance, must use CA-25G-L with RS-FEC and 25GBASE-CR
- 10 Gigabit Ethernet over SFP+ Direct Attached Cable (Twinaxial)
 - Length is 10 meters max.

Install Drivers and Software

Windows* Operating Systems

You must have administrative rights to the operating system to install the drivers.

1. Download the latest drivers from the [support website](#) and transfer them to the system.
2. If the Found New Hardware Wizard screen is displayed, click **Cancel**.
3. Double-click the downloaded file.
4. Select **Install** from the Dell Update Package screen.
5. Follow the instructions in the install wizard. Be sure to select Intel PROSet for installation.



NOTE: Be sure to select the "iSCSI using Data Center Bridging" install option for systems that have an NPAR capable device installed.

Refer to "Microsoft* Windows* Driver and Software Installation and Configuration" on page 54 for more specific information.

Installing Linux* Drivers from Source Code

1. Download and expand the driver tar file.
2. Compile the driver module.
3. Install the module using the modprobe command.
4. Assign an IP address using the ifconfig command.

Please refer to the [Linux section](#) of this guide for more specific information.

Installing Linux Drivers from RPMs

1. Download and expand the driver tar file.
2. Install the driver using the rpm command.

Please refer to the [Linux section](#) of this guide for more specific information.

Device Features

This chapter describes the features available on Intel Ethernet devices. Major features are organized alphabetically.



NOTE: Available settings are dependent on your device and operating system. Not all settings are available on every device/OS combination.

Adaptive Inter-Frame Spacing

Compensates for excessive Ethernet packet collisions on the network.

The default setting works best for most computers and networks. By enabling this feature, the network adapter dynamically adapts to the network traffic conditions. However, in some rare cases you might obtain better performance by disabling this feature. This setting forces a static gap between packets.

Default	Disabled
Range	• Enabled • Disabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Adaptive Inter-Frame Spacing" -
DisplayValue "Enabled"
```

Data Center Bridging (DCB)

Data Center Bridging provides a lossless data center transport layer for using LANs and SANs in a single unified fabric.

Data Center Bridging includes the following capabilities:

- Priority-based flow control (PFC; IEEE 802.1Qbb)
- Enhanced transmission selection (ETS; IEEE 802.1Qaz)
- Congestion notification (CN)
- Extensions to the Link Layer Discovery Protocol (LLDP) standard (IEEE 802.1AB) that enable Data Center Bridging Capability Exchange Protocol (DCBX)

There are two supported versions of DCBX.

CEE Version: The specification can be found as a link within the following document: <http://www.ieee802.org/1/files/public/docs2008/dcb-baseline-contributions-1108-v1.01.pdf>

IEEE Version: The specification can be found as a link within the following document: <https://standards.ieee.org/findstds/standard/802.1Qaz-2011.html>



NOTE: The OS DCBX stack defaults to the CEE version of DCBX, and if a peer is transmitting IEEE TLVs, it will automatically transition to the IEEE version.

For more information on DCB, including the DCB Capability Exchange Protocol Specification, go to <http://www.ieee802.org/1/pages/dcbbridges.html>

DCB for Windows Configuration



NOTES:

- On systems running a Microsoft Windows Server operating system, enabling *QoS/priority flow control will disable link level flow control.
- If *QOS/DCB is not available, it may be for one of the following reasons:

- The Firmware LLDP (FW-LLDP) agent was disabled from a pre-boot environment (typically UEFI).
- This device is based on the Intel® Ethernet Controller X710 and the current link speed is 2.5 Gbps or 5 Gbps

This setting is found on the Data Center tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

You can use the Intel® PROSet to perform the following tasks:

- **Display Status:**
 - Enhanced Transmission Selection
 - Priority Flow Control

Non-operational status: If the Status indicator shows that DCB is non-operational, there may be a number of possible reasons:

- DCB is not enabled - select the checkbox to enable DCB.
- One or more of the DCB features is in a non-operational state.

A non-operational status is most likely to occur when **Use Switch Settings** is selected or **Using Advanced Settings** is active. This is generally a result of one or more of the DCB features not getting successfully exchanged with the switch. Possible problems include:

- One of the features is not supported by the switch.
- The switch is not advertising the feature.
- The switch or host has disabled the feature (this would be an advanced setting for the host).
- Disable/enable DCB
- Troubleshooting information

Hyper-V (DCB and VMQ)

 **NOTE:** Configuring a device in the VMQ + DCB mode reduces the number of VMQs available for guest OSes.

DCB for Linux

 **NOTE:** DCB is supported on RHEL6 or later or SLES11 SP1 or later. See your operating system documentation for specifics.

Intel Ethernet drivers support firmware-based or software-based DCBX in Linux, depending on the underlying PF device. The following table summarizes DCBX support by driver.

Linux Driver	Firmware-Based DCBX	Software-Based DCBX
ice	Supported	Supported
i40e	Supported	Supported
ixgbe	Not supported	Supported

In **firmware-based** mode, firmware intercepts all LLDP traffic and handles DCBX negotiation transparently for the user. In this mode, the adapter operates in "willing" DCBX mode, receiving DCB settings from the link partner (typically a switch). The local user can only query the negotiated DCB configuration.

In **software-based** mode, LLDP traffic is forwarded to the network stack and user space, where a software agent can handle it. In this mode, the adapter can operate in either "willing" or "nonwilling" DCBX mode and DCB configuration can be both queried and set locally. Software-based mode requires the FW-based LLDP Agent to be disabled, if supported.

NOTES:

- Only one LLDP/DCBX agent can be active on a single interface at a time.
- Software-based and firmware-based DCBX modes are mutually exclusive.
- When the firmware DCBX agent is active, software agents will not be able to receive or transmit LLDP frames. See "Firmware Link Layer Discovery Protocol (FW-LLDP)" on page 13, as well as the Linux driver readme in

your installation, for information on enabling or disabling the FW-LLDP agent.

- In software-based DCBX mode, you can configure DCB parameters using software LLDP/DCBX agents that interface with the Linux kernel's DCB Netlink API. We recommend using OpenLLDP as the DCBX agent when running in software mode. For more information, see the OpenLLDP man pages and <https://github.com/intel/openlldp>.
- For information on configuring DCBX parameters on a switch, please consult the switch manufacturer's documentation.

iSCSI Over DCB

Intel® Ethernet adapters support iSCSI software initiators that are native to the underlying operating system. Data Center Bridging is most often configured at the switch. If the switch is not DCB capable, the DCB handshake will fail but the iSCSI connection will not be lost.



NOTE: DCB does not install in a VM. iSCSI over DCB is only supported in the base OS. An iscsi initiator running in a VM will not benefit from DCB ethernet enhancements.

Microsoft Windows Configuration

iSCSI installation includes the installation of the iSCSI DCB Agent (iscsidcb.exe) user mode service. The Microsoft iSCSI Software Initiator enables the connection of a Windows host to an external iSCSI storage array using an Intel Ethernet adapter. Please consult your operating system documentation for configuration details.

Enable DCB on the adapter by the following:

This setting is found on the Data Center tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

The Data Center Tab provides feedback as to the DCB state, operational or non-operational, as well as providing additional details should it be non-operational.



NOTE: On Microsoft Windows Server operating systems, if you configure Priority using IEEE, the iSCSI policy may not be created automatically. To create the iSCSI policy manually, use Powershell and type:

```
New-NetQosPolicy -Name "UP4" -PriorityValue 8021 Action 4 -iSCSI
```

Linux Configuration

In the case of Open Source distributions, virtually all distributions include support for an Open iSCSI Software Initiator and Intel® Ethernet adapters will support them. Please consult your distribution documentation for additional configuration details on their particular Open iSCSI initiator.

Intel® 82599 and X540-based adapters support iSCSI within a Data Center Bridging cloud. Used in conjunction with switches and targets that support the iSCSI/DCB application TLV, this solution can provide guaranteed minimum bandwidth for iSCSI traffic between the host and target. This solution enables storage administrators to segment iSCSI traffic from LAN traffic. Previously, iSCSI traffic within a DCB supported environment was treated as LAN traffic by switch vendors. Please consult your switch and target vendors to ensure that they support the iSCSI/DCB application TLV.

Direct Memory Access (DMA) Coalescing

DMA (Direct Memory Access) allows the network device to move packet data directly to the system's memory, reducing CPU utilization. However, the frequency and random intervals at which packets arrive do not allow the system to enter a lower power state. DMA Coalescing allows the NIC to collect packets before it initiates a DMA event. This may increase network latency but also increases the chances that the system will consume less energy. Adapters and network devices based on the Intel® Ethernet Controller I350 (and later controllers) support DMA Coalescing.

Higher DMA Coalescing values result in more energy saved but may increase your system's network latency. If you enable DMA Coalescing, you should also set the Interrupt Moderation Rate to 'Minimal'. This minimizes the latency impact imposed by DMA Coalescing and results in better peak network throughput performance. You must enable DMA Coalescing on all active ports in the system. You may not gain any energy savings if it is enabled only on some of the ports in your system. There are also several BIOS, platform, and application settings that will affect your potential energy savings. A white paper containing information on how to best configure your platform is available on the Intel website.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "DMA Coalescing" -DisplayValue "Enabled"
```

Dynamic Device Personalization (DDP)

Adapters based on the Intel® Ethernet Controller 800 Series require a Dynamic Device Personalization (DDP) package file to enable advanced features (such as dynamic tunneling, Flow Director, RSS, and ADQ). DDP allows you to change the packet processing pipeline of a device by applying a profile package to the device at runtime. Profiles can be used to, for example, add support for new protocols, change existing protocols, or change default settings. DDP profiles can also be rolled back without rebooting the system.

The DDP package loads during device initialization. The driver checks to see if the DDP package is present and compatible. If this file exists, the driver will load it into the device. If not, the driver will go into Safe Mode where it will use the configuration contained in the device's NVM.

Safe Mode disables advanced and performance features, and supports only basic traffic and minimal functionality, such as updating the NVM or downloading a new driver or DDP package. For more information, see [Safe Mode](#).

Microsoft Windows and ESX drivers embed the DDP package in the driver itself. Linux loads the package from an external file:

- Linux: Loads the intel/ice/ddp/ice.pkg from your firmware root (typically /lib/firmware/ or /lib/firmware/updates/).



NOTES:

- You cannot update the DDP package if any PF drivers are already loaded. To overwrite a package, unload all PFs and then reload the driver with the new package.
- Except for Linux, you can only use one DDP package per driver, even if you have more than one device installed that uses the driver.
- Only the first loaded PF per device can download a package for that device.
- If you are using DPDK, see the DPDK documentation for installation instructions and more information.

Firmware Link Layer Discovery Protocol (FW-LLDP)

Devices based on the Intel® Ethernet Controller 800 and 700 Series use a Link Layer Discovery Protocol (LLDP) agent that runs in the firmware. When it is running, it prevents the operating system and applications from receiving LLDP traffic from the network adapter.

- The FW-LLDP setting is per port and persists across reboots.
- The FW-LLDP Agent is required for DCB to function.

Adapters Based on the Intel® Ethernet Controller 800 Series

FW-LLDP is disabled in NVM by default. To enable/disable the FW-LLDP Agent:

- **Linux:** Use ethtool to persistently set or show the fw-lldp-agent private flag.
- **ESX:** Use the esxcli command to persistently set or get the fw_lldp_agent setting.
- **Microsoft Windows:** The base driver does not persistently change FW-LLDP. Use the LLDP Agent attribute in UEFI HII to persistently change the FW-LLDP setting. If you enable DCB when FW-LLDP is disabled, the base driver temporarily starts the LLDP Agent while DCB functionality is enabled.

Adapters Based on the Intel® Ethernet Controller 700 Series

FW-LLDP is enabled in NVM by default. To enable/disable the FW-LLDP Agent:

- **Linux:** Use ethtool to set or show the disable-fw-ldp private flag.
- **ESX:** Use the esxcfg-module command to set or get the LLDP module parameter.
- **Microsoft Windows:** Use the LLDP Agent attribute in UEFI HII to change the FW-LLDP setting. Note: You must enable the UEFI HII "LLDP AGENT" attribute for the FW-LLDP setting to take effect. If "LLDP AGENT" is set to disabled in UEFI HII, you cannot enable FW-LLDP from the OS.
- You must enable the LLDP Agent from UEFI HII to use DCB.

Forward Error Correction (FEC) Mode

Allows you to set the Forward Error Correction (FEC) mode. FEC improves link stability, but increases latency. Many high quality optics, direct attach cables, and backplane channels provide a stable link without FEC.

The driver allows you to set the following FEC Modes:

- Auto FEC - Sets the FEC Mode based on the capabilities of the attached cable.
- CL108 RS-FEC - Selects only RS-FEC ability and request capabilities.
- CL74 FC-FEC/BASE-R - Selects only BASE-R ability and request capabilities.
- No FEC - Disables FEC.



NOTES:

- For devices to benefit from this feature, link partners must have FEC enabled.
- Intel® Ethernet Controller 800 Series devices only enable Forward Error Correction (FEC) configurations that are supported by the connected media and which are expected to yield healthy Bit Error Rate (BER) connections.
- If you are having link issues (including no link) at link speeds faster than 10 Gbps, check your switch configuration and/or specifications. Many optical connections and direct attach cables require RS-FEC for connection speeds faster than 10 Gbps. One of the following may resolve the issue:
 - Configure your switch to use RS-FEC mode.
 - Specify a 10 Gbps, or slower, link speed connection.
 - If you are attempting to connect at 25 Gbps, try using an SFP28 CA-S or CS-N Direct Attach cable. These cables do not require RS-FEC.
 - If your switch does not support RS-FEC mode, check with your switch vendor for the availability of a SW or FW upgrade.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "FEC Mode" -DisplayValue "Auto FEC"
```

Flow Control

Enables adapters to more effectively regulate traffic. Adapters generate flow control frames when their receive queues reach a pre-defined limit. Generating flow control frames signals the transmitter to slow transmission. Adapters respond to flow control frames by pausing packet transmission for the time specified in the flow control frame.

By enabling adapters to adjust packet transmission, flow control helps prevent dropped packets. You may improve RDMA performance by enabling flow control on all nodes and on the switch they are connected to.



NOTES:

- For adapters to benefit from this feature, link partners must support flow control frames.
- On systems running a Microsoft Windows Server operating system, enabling *QoS/priority flow control will disable link level flow control.
- Some devices support Auto Negotiation. Selecting this will cause the device to advertise the value stored in

its NVM (usually "Disabled").

- When an adapter is running in NPar mode, Flow Control is limited to the root partition of each port.

Default	Disabled (Microsoft Windows Server 2019 and later) RX & TX Enabled (Microsoft Windows Server 2016 and earlier)
Range	• Disabled • RX Enabled • TX Enabled • RX & TX Enabled • Auto Negotiation (only available on some adapters)

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Flow Control" -DisplayValue "Rx Enabled"
```

Gigabit Master Slave Mode

Determines whether the adapter or link partner is designated as the master. The other device is designated as the slave. By default, the IEEE 802.3ab specification defines how conflicts are handled. Multi-port devices such as switches have higher priority over single port devices and are assigned as the master. If both devices are multi-port devices, the one with higher seed bits becomes the master. This default setting is called "Hardware Default."



NOTE: In most scenarios, it is recommended to keep the default value of this feature.

Setting this to either "Force Master Mode" or "Force Slave Mode" overrides the hardware default.

Default	Auto Detect
Range	• Force Master Mode • Force Slave Mode • Auto Detect



NOTE: Some multi-port devices may be forced to Master Mode. If the adapter is connected to such a device and is configured to "Force Master Mode," link is not established.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Gigabit Master Slave Mode" -DisplayValue "Auto Detect"
```

Interrupt Moderation Rate

Sets the Interrupt Throttle Rate (ITR). This setting moderates the rate at which Transmit and Receive interrupts are generated.

When an event such as packet receiving occurs, the adapter generates an interrupt. The interrupt interrupts the CPU and any application running at the time, and calls on the driver to handle the packet. At greater link speeds, more interrupts are created, and CPU rates also increase. This results in poor system performance. When you use a higher ITR setting, the interrupt rate is lower and the result is better CPU performance.



NOTE: A higher ITR rate also means that the driver has more latency in handling packets. If the adapter is handling many small packets, it is better to lower the ITR so that the driver can be more responsive to incoming and outgoing packets.

Altering this setting may improve traffic throughput for certain network and system configurations, however the default setting is optimal for common network and system configurations. Do not change this setting without verifying that the desired change will have a positive effect on network performance.

Default	Adaptive
Range	• Adaptive • Extreme • High • Medium • Low • Minimal • Off

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Interrupt Moderation Rate" -DisplayValue "Adaptive"
```

IPv4 Checksum Offload

This allows the adapter to compute the IPv4 checksum of incoming and outgoing packets. This feature enhances IPv4 receive and transmit performance and reduces CPU utilization.

With Offloading off, the operating system verifies the IPv4 checksum.

With Offloading on, the adapter completes the verification for the operating system.

Default	RX & TX Enabled
Range	• Disabled • RX Enabled • TX Enabled • RX & TX Enabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "IPv4 Checksum Offload" -DisplayValue "Tx Enabled"
```

Jumbo Frames

Enables or disables Jumbo Packet capability. The standard Ethernet frame size is about 1514 bytes, while Jumbo Packets are larger than this. Jumbo Packets can increase throughput and decrease CPU utilization. However, additional latency may be introduced.

Enable Jumbo Packets only if ALL devices across the network support them and are configured to use the same frame size. When setting up Jumbo Packets on other network devices, be aware that network devices calculate Jumbo Packet sizes differently. Some devices include the frame size in the header information while others do not. Intel adapters do not include frame size in the header information.

Restrictions

- Supported protocols are limited to IP (TCP, UDP).
- Jumbo frames require compatible switch connections that forward Jumbo Frames. Contact your switch vendor for more information.
- When standard-sized Ethernet frames (64 to 1518 bytes) are used, there is no benefit to configuring Jumbo Frames.
- The Jumbo Packets setting on the switch must be set to at least 8 bytes larger than the adapter setting for Microsoft Windows operating systems, and at least 22 bytes larger for all other operating systems.
- Jumbo Frames are not supported over VLANs under Microsoft Windows 10. The only Microsoft operating systems that support Jumbo Frames over VLAN are Microsoft Windows Server 2008 R2, Microsoft Windows Server 2012, and Microsoft Windows Server 2012 R2.

Default	Disabled
Range	Disabled (1514), 4088, or 9014 bytes. (Set the switch 4 bytes higher for CRC, plus 4 bytes if using VLANs.)



NOTES:

- End-to-end hardware must support this capability; otherwise, packets will be dropped.
- Intel adapters that support Jumbo Packets have a frame size limit of 9238 bytes, with a corresponding MTU size limit of 9216 bytes.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Jumbo Packet" -DisplayValue "4088 Bytes"
```

Large Send Offload (IPv4 and IPv6)

Sets the adapter to offload the task of segmenting TCP messages into valid Ethernet frames. The maximum frame size limit for large send offload is set to 64,000 bytes.

Since the adapter hardware is able to complete data segmentation much faster than operating system software, this feature may improve transmission performance. In addition, the adapter uses fewer CPU resources.

Default	Enabled
Range	- Enabled - Disabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Large Send Offload V2 (IPv4)" -DisplayValue "Enabled"
```

Link State on Interface Down

Sets if link is enabled or disabled when the interface is brought down. If this is set to **Disabled** and you bring an interface down (using an administrative tool, or in another way), then the port will lose link. This allows an attached switch to detect that the interface is no longer up. However, if Wake on LAN or manageability is enabled on this port, link will remain up.

Default	Enabled
Range	• Enabled • Disabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Link State on Interface Down" -
DisplayValue "Enabled"
```

Locally Administered Address

Overrides the initial MAC address with a user-assigned MAC address. To enter a new network address, type a 12-digit hexadecimal number in this box.

Default	None
Range	0000 0000 0001 - FFFF FFFF FFFD Exceptions: • Do not use a multicast address (Least Significant Bit of the high byte = 1). For example, in the address 0Y123456789A, "Y" cannot be an odd number. (Y must be 0, 2, 4, 6, 8, A, C, or E.) • Do not use all zeros or all Fs. If you do not enter an address, the address is the original network address of the adapter. For example, Multicast: 0123 4567 8999 Broadcast: FFFF FFFF FFFF Unicast (legal): 0070 4567 8999

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Locally Administered Address" -
DisplayValue "<desired address>"
```

Log Link State Event

This setting is used to enable/disable the logging of link state changes. If enabled, a link up change event or a link down change event generates a message that is displayed in the system event logger. This message contains the link's speed and duplex. Administrators view the event message from the system event log.

The following events are logged.

- The link is up.
- The link is down.
- Mismatch in duplex.
- Spanning Tree Protocol detected.

Default	Enabled
Range	Enabled, Disabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Log Link State Event" -DisplayValue "Enabled"
```

Low Latency Interrupts

LLI enables the network device to bypass the configured interrupt moderation scheme based on the type of data being received. It configures which arriving TCP packets trigger an immediate interrupt, enabling the system to handle the packet more quickly. Reduced data latency enables some applications to gain faster access to network data.

 **NOTE:** When LLI is enabled, system CPU utilization may increase.

LLI can be used for data packets containing a TCP PSH flag in the header or for specified TCP ports.

- **Packets with TCP PSH Flag** - Any incoming packet with the TCP PSH flag will trigger an immediate interrupt. The PSH flag is set by the sending device.
- **TCP Ports** - Every packet received on the specified ports will trigger an immediate interrupt. Up to eight ports may be specified.

Default	Disabled
Range	• Disabled • PSH Flag-Based • Port-Based

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Low Latency Interrupts" -DisplayValue "Port-Based"
```

Malicious Driver Detection (MDD) for VFs

Some Intel Ethernet devices use Malicious Driver Detection (MDD) to detect malicious traffic from the VF and disable Tx/Rx queues or drop the offending packet until a VF driver reset occurs. You can view MDD messages in the PF's event log.

- If the device supports automatic VF resets and the driver detects an MDD event on the receive path, the PF will automatically reset the VF and reen able queues. If automatic VF resets are disabled, the PF will not automatically reset the VF when it detects MDD events. See the table below for supported MDD features.
- If the PF driver logs MDD events from the VF, confirm that the correct VF driver is installed.
- To restore functionality, you can manually reload the VF or VM or, if supported by the device, enable automatic VF resets.

Feature	Intel Ethernet Controller 800 Series Adapters	Intel Ethernet Controller 700 Series Adapters	Intel Ethernet Controller 500 Series Adapters	Intel® I350 Gigabit Network Connection
Automatically resets the VF	If enabled	Yes	Yes	Yes

and reenables queues after MDD events				
Can disable automatic VF reset after MDD events	Yes	No	No	No

MDD Auto Reset VFs

Automatically resets the virtual machine immediately after the adapter detects a Malicious Driver Detection (MDD) event on the receive path.

Default	Disabled
Range	• Disabled • Enabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "MDD Auto Reset VFs" -DisplayValue "Enabled"
```

Max Number of RSS Queues Per Vport

Sets the maximum number of Receive Side Scaling (RSS) queue pairs per VF.

Default	4 Queues
Range	• 2 Queues • 4 Queues • 8 Queues • 16 Queues

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Max Number of RSS Queues Per Vport" -DisplayValue "4 Queues"
```

Network Virtualization using Generic Routing Encapsulation (NVGRE)

Network Virtualization using Generic Routing Encapsulation (NVGRE) increases the efficient routing of network traffic within a virtualized or cloud environment. Some Intel® Ethernet Network devices perform Network Virtualization using Generic Routing Encapsulation (NVGRE) processing, offloading it from the operating system. This reduces CPU utilization.



NOTE: When a port is in NPar mode, NVGRE (the Encapsulated Task Offload setting) is available only on the first partition on the port.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the `Set-IntelNetAdapterSetting` cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "NVGRE Encapsulated Task Offload"
-DisplayValue "Enabled"
```

NIC Partitioning

Network Interface Card (NIC) Partitioning (NPar) allows network administrators to create multiple partitions for each physical port on a network adapter card, and to set different bandwidth allocations on each partition. To the network and operating system, each partition appears as a separate physical port on the adapter. This facilitates the reduction of switch port count and cabling complexity while maintaining network segmentation and isolation. In addition, flexible bandwidth allocation per partition allows for efficient use of the link.

NPar is available in Linux and ESXi, and in Windows Server and Windows Server Core versions starting with 2012 R2.

The following adapters support NPar. Note that NPar supports a maximum of 8 partitions per controller.

- Intel® Ethernet 25G 2P XXV710 Adapter



NOTE: 25GbE adapters do not support NPAR and NPAR-EP on VMWare operating systems.

- Intel® Ethernet 10G 4P X710/I350 rNDC
- Intel® Ethernet 10G 4P X710-k bNDC
- Intel® Ethernet 10G 4P X710 rNDC
- Intel® Ethernet 10G 2P X710-k bNDC
- Intel® Ethernet 10G X710-k bNDC
- Intel® Ethernet Converged Network Adapter X710
- Intel® Converged Network Adapter X710-T
- Intel® Ethernet Server Adapter X710-DA2 for OCP



NOTES:

- Adapters support NPar in NIC (LAN) mode only.
- The following are supported on the first partition of each port only:
 - PXE Boot
 - iSCSIboot
 - Speed and Duplex settings
 - Flow Control
 - Power Management settings
 - SR-IOV
 - NVGRE processing
- Some adapters only support Wake on Lan on the first partition of the first port.
- Resource limits in Microsoft Windows may affect the number of ports that are displayed. If you have several adapters installed in a system, and enable NPar or NParEP on the adapters, Windows Device Manager may not display all of the ports.
- Minimum bandwidth may not be distributed equally between the NIC partitions when changing NPAR/NPAR EP mode. The minimum bandwidth values can be adjusted after changing the NPAR/NPAR EP mode.
- iSCSI Offload is not supported on NIC partitions of Intel X710 based devices. X710 adapters incorrectly show a value of "True" for "iSCSI Offload Support". Enabling "iSCSI Offload Mode" from the [NIC Partitioning Configuration] page enables the partition for iSCSI storage traffic.
- The Loopback diagnostic test is not supported when the device is in NPAR mode.
- When configuring the system for a Microsoft Windows based OS, do not enable iSCSI Offload Mode in the Partition Configuration for Intel® X710 devices either directly in BIOS via HII, or through remote configuration such as racadm or WSMAN.
- If you have NPAR enabled, make sure the "RSS load balancing profile" Advanced setting is set to NUMAScalingStatic.
- NVGRE is not supported when the device is in NPAR mode. If your device has NPAR enabled, NVGRE (the Encapsulated Task Offload setting on the Advanced tab in Windows Device Manager) is not supported.
- With NPAR enabled on Intel® Ethernet Controller 700 series devices, all partitions will lose network con-

nection for 2-3 seconds while the root partition (the first partition of the physical port) is initializing.

NParEP Mode

NParEP Mode is a combination of NPar and PCIe ARI, and increases the maximum number of partitions on an adapter to 16 per controller.

NParEP Platform Support

Dell EMC Platform	OCP Mezz	Rack NDC Slot	PCI Express Slot												
			1	2	3	4	5	6	7	8	9	10	11	12	13
C4130			yes	yes											
C4140		no	yes	no	yes										
C6420	yes		yes												
R230			no	no											
R240			no	no											
R330			no	no											
R340			no	no											
R430			yes	yes											
R440			yes	yes	yes										
R530			yes	yes	yes	no	no								
R530XD			yes	yes	no										
R540			yes	yes	yes	yes	yes	no							
R630		yes	yes	yes	yes										
R640		yes	yes	yes	yes										
R730		yes	yes	yes	yes	yes	yes	yes	yes						
R730XD		yes	yes	yes	yes	yes	yes	yes							
R740		yes	yes	yes	yes	yes	yes	yes	yes	yes					
R740XD2		no	yes	yes	yes	yes	yes	no							
R830		yes	yes	yes	yes	yes	yes	yes							
R840		yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes			
R930		yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes			
R940		yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes
T130			no	no	no	no									
T140			no	no	no	no									
T330			no	no	no	yes									

			PCI Express Slot												
Dell EMC Platform	OCP Mezz	Rack NDC Slot	1	2	3	4	5	6	7	8	9	10	11	12	13
T340			no	no	no	no									
T430			no	no	yes	yes	yes	yes							
T440			no	yes	yes	yes	yes								
T630			yes	no	yes	yes	yes	yes	yes						
T640		yes	yes	yes	yes	yes	yes	yes	yes	yes					

The following Dell EMC Platforms support NParEP mode on all slots.

- C6520
- C6525
- R650
- R650xa
- R6515
- R6525
- R750
- R750xa
- R7525
- R7515

		Mezzanine Slot	
Dell EMC Platform	Blade NDC Slot	B	C
FC430			
FC630	yes		
FC830	yes		
M630	yes		
M630 for VRTX	yes		
M640	yes		
M640 for VRTX	yes		
M830	yes		
M830 for VRTX	yes		
MX740c	yes	yes	
MX840c	yes	yes	

Supported platforms or slots are indicated by "yes." Unsupported are indicated by "no". Not applicable are indicated by blank cells.

Configuring NPar Mode

Configuring NPar from the Boot Manager

When you boot the system, press the **F2** key to enter the **System Setup** menu. Select **Device Settings** from the list under **System Setup Main Menu**, then select your adapter from the list to get to the Device Configuration menu. Select **Device Level Configuration** in the list under **Main Configuration Page**. This brings up the Virtualization settings under **Device Level Configuration**.

There are four options in the **Virtualization Mode** drop down list.

- None: the adapter operates normally
- NPar: allows up to 8 partitions on the adapter. If you select NPar Virtualization Mode, you will then be presented with the option to enable NParEP Mode, which will extend the number of partitions per adapter to a total of 16 by pairing NPar with PCIe ARI.



NOTES:

- When an adapter is running in NPar Mode, it is limited to 8 partitions total. A two-port adapter will have four partitions per port. A four-port adapter will have two partitions per port.
- NParEP Mode can only be enabled when NPar Mode has been enabled.
- When an adapter is running in NParEP Mode, it is limited to 16 partitions total. A two-port adapter will have eight partitions per port. A four port adapter will have four partitions per port.
- SR-IOV: activates SR-IOV on the port
- NPar+SR-IOV: Allows up to 8 partitions (physical functions) for the adapter and activates SR-IOV.



NOTES:

- SR-IOV is limited to the root partition of each port.
- When an adapter is running in NPar mode, virtualization (SR-IOV) settings apply to all ports on the adapter, and to all partitions on each port. Changes made the virtualization settings on one port are applied to all ports on the adapter.

When you have completed your selection, click the **Back** button, and you will return to the **Main Configuration Page**. Click the new item, titled **NIC Partitioning Configuration**, in the configuration list to go to the NIC Partitioning Configuration page, where you will see a list of the NPar (or NParEP) partitions on your adapter.

The Global Bandwidth Allocation page lets you specify the minimum and maximum guaranteed bandwidth allocation for each partition on a port. Minimum TX Bandwidth is the guaranteed minimum data transmission bandwidth, as a percentage of the full physical port link speed, that the partition will receive. The bandwidth the partition is awarded will never fall below the level you specify here. The valid range of values is:

1 to ((100 minus # of partitions on the physical port) plus 1)

For example, if a physical port has 4 partitions, the range would be:

1 to ((100 - 4) + 1 = 97)

The Maximum Bandwidth percentage represents the maximum transmit bandwidth allocated to the partition as a percentage of the full physical port link speed. The accepted range of values is 0-100. The value here can be used as a limiter, should you chose that any one particular partition not be able to consume 100% of a port's bandwidth should it be available. The sum of all the values for Maximum Bandwidth is not restricted, because no more than 100% of a port's bandwidth can ever be used.



NOTE:

- If the sum of the minimum bandwidth percentages does not equal 100, then settings will be automatically adjusted so that the sum equals 100.
- If a partition's maximum bandwidth percentage is set lower than the partition's minimum bandwidth percentage, then the maximum bandwidth percentage will be automatically set to the value of the minimum bandwidth percentage.
- When you attempt to set values for minimum bandwidth percentage via iDRAC with Lifecycle Controller using jobs that do not include the values for all enabled partitions, then the values seen after the jobs have com-

pleted may be different than the values that were supposed to be set. To avoid this issue, set the values for minimum bandwidth percentage on all partitions using a single job and make sure the sum of the values is 100.

Click the **Back** button when you have finished making your bandwidth allocation settings to return to the NIC Partitioning Configuration page. From there you may click on one of the **Partition Configuration** list items under **Global Bandwidth Allocation**. This will bring up a partition configuration information page on a particular port. You can see the NIC Mode, PCI;Device ID, PCI Address, MAC Address, and the virtual MAC Address (if applicable) for all the partitions on any given port by clicking through the items in the Partition Configuration list.

When you have completed the configuration of all the partitions on a port, back out to the Main Configuration Page, click the **Finish** button, then click the **OK** button in the Success (Saving Changes) dialog.

Repeat the partition configuration process for all the ports on your adapter.



NOTE: Once NPar has been enabled on one of the partition on a port, it will appear enabled for all subsequent partitions on that port. If that first setting of NPar included enabling NParEP mode, NParEP Mode will appear enabled on all subsequent partitions on that port as well.

When you have configured all the partitions on all the ports on all the adapters in your server, back out to the System Setup Main Menu, and click the **Finish** button. Then click **Yes** to exit the System Setup Menu and to reboot the system in order to apply your changes.

Once the system has completed the boot process, NPar will remain enabled until you explicitly disable it by turning off the option during a subsequent boot sequence.

Configuring NPar in Microsoft Windows

You can configure an adapter port partition in Windows just like any adapter port. Run Device Manager, select the and open the partition's properties sheets to configure options.



NOTE: On Microsoft® Windows Server® 2019 (and later), you must use Intel® PROSet for Windows PowerShell to configure NPar.

Enabling NPar

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "NIC Partitioning" -DisplayValue "NPAR"
```

Boot Options

On the Boot Options tab, you will be advised that the device is in NPar mode and that legacy preboot protocol settings can only be configured on the root partition. Clicking the **Properties** button will launch the property sheet for the root partition on the adapter.

To set this using Windows PowerShell, find the first partition using the Get-IntelNetAdapter cmdlet. Once you know the port with partition number 0, use that port name in the Bootutil utility for boot option configuration.

Power Management Settings

Power Management settings are allowed only on the first partition of each physical port. If you select the **Power Management** tab while any partition other than the first partition is selected, you will be presented with text in the Power Management dialog stating that Power Management settings cannot be configured on the current connection. Clicking the **Properties** button will launch the property sheet for the root partition on the adapter.



NOTE: Boot options and Power Management settings are only available on the root partition of each physical port.

To set this using Windows PowerShell, find the first partition using the Get-IntelNetAdapter cmdlet. Once you know the port with partition number 0, use that port name with the Get-IntelNetAdapterSetting and Set-IntelNetAdapterSetting cmdlets.

Flow Control

You can change the Flow Control settings for any partition on a given port. However, when a change is made to the Flow Control settings of a partition associated with a port on an adapter operating in NPar mode, the new value will be applied to all partitions on that particular port.

Flow control is reached by selecting the **Advanced** tab, then selecting the **Properties** button, and then selecting **Flow Control** from the list of options in the **Settings** list of the dialog that is displayed.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example, Set-IntelNetAdapterSetting -Name "<adapter name>" -DisplayName "Flow Control" -DisplayValue "Auto Negotiation"

Identifying Port Associations

The Hardware Information dialog in the Intel PROSet property sheets facilitates the identification of the physical port associated with a particular partition. There is an **Identify Adapter** button in the **Link Speed** tab, and clicking that button will cause the ACK/Link light on the port associated with the active partition to blink.

To change this setting in Windows PowerShell, use the Test-IntelNetAdapterSetting cmdlet. For example, Test-IntelNetIdentifyAdapter -Name "<adapter name>" -Seconds 100

Partition Bandwidth Configuration

The Bandwidth Configuration dialog provides an indication of the port for which settings are currently being made, above a list of the partitions on that port and their current bandwidth allocations (Min%, Max%). Partition Bandwidth Configuration is reached by clicking the **Bandwidth Configuration** button on the **Link Speed** tab.

The bandwidth allocated to each partition on the port will never drop below the value set under Min%. For all the partitions on the same physical port, the min bandwidth percentage for all of the partitions must be set to zero, or the sum of all of the minimum bandwidth percentages on each partition must equal 100, where the range of min bandwidth percentages is between 1 and (100-n)%, where n is the number of partitions for a particular port. For example, on a port with four defined partitions:

P1=0	P1=10	P1=20
P2=0	P2=20	P2=80
P3=0	P3=30	P3=0
P4=0	P4=40	P4=0
Valid	Valid	NOT Valid

Valid values for Max% are the value of that partition's "Min%" through "100". For example, if the Min% value for Partition 1 is 50%, the range for that partition's Max% is "50"-"100". If you cause any one partition's Max% value to exceed 100% by incrementing the value with the spinner, an error is displayed and that Max% is decremented to 100%. The *sum* of the Max% values for all partitions on a particular port has no limit.

To change the value for Min% or Max%, select a partition in the displayed list, then use the up or down arrows under “Selected Partition Bandwidth Percentages”.

 NOTE:

- If the sum of the minimum bandwidth percentages does not equal 100, then settings will be automatically adjusted so that the sum equals 100.
- If a partition's maximum bandwidth percentage is set lower than the partition's minimum bandwidth percentage, then the maximum bandwidth percentage will be automatically set to the value of the minimum bandwidth percentage.
- When you attempt to set values for minimum bandwidth percentage via iDRAC with Lifecycle Controller using jobs that do not include the values for all enabled partitions, then the values seen after the jobs have completed may be different than the values that were supposed to be set. To avoid this issue, set the values for minimum bandwidth percentage on all partitions using a single job and make sure the sum of the values is 100.

To change this setting in Windows PowerShell, use the `Set-IntelNetAdapterSetting` cmdlet. For example:

```
Set-IntelNetAdapterMaximumBandwidthPercentage -Name "Intel(R) Ethernet Converged Network Adapter X710" -MaxBwPercent 100
```

```
Set-IntelNetAdapterMinimumBandwidthPercentage -Name "Intel(R) Ethernet Converged Network Adapter X710" -Partition1 25 -Partition2 25 -Partition3 25 -Partition4 25
```

Speed and Duplex Settings

The Speed and Duplex setting for a particular port may be changed from any partition associated with that port. However, because all partitions on a particular port on an adapter operating in NPar mode share the same module that is plugged into the port, changing the Speed and Duplex setting will result in the new value being set across all partitions on that same physical port.

Changing the Speed and Duplex setting for a port on an adapter running in NPar mode will cause the driver for each partition associated with that port to be reloaded. This may result in a momentary loss of link.

Online Diagnostics

Online tests can be performed while in NPar mode without the adapter losing link. The following diagnostics tests are available for all partitions for a particular port while an adapter is running in NPar mode:

- EEPROM
- Register
- NVM Integrity
- Connection

To change this setting in Windows PowerShell, use the `Test-IntelNetAdapterSetting` cmdlet. For example, `Test-IntelNetDiagnostics -Name "<adapter name>" -Test Hardware`

Offline Diagnostics

Offline diagnostics are not supported while an adapter is running in NPar mode. Loopback tests and Cable Offline tests are not allowed in NPar mode.

To change this setting in Windows PowerShell, use the `Test-IntelNetAdapterSetting` cmdlet. For example, `Test-IntelNetDiagnostics -Name "<adapter name>" -Test Hardware`

Virtualization

Settings for virtualization (Virtual Machine Queues and SR-IOV) are found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility, then selecting “Virtualization” from the Settings list.

When an adapter is operating in NPar mode, only the first partition of each physical port may be configured with the virtualization settings.

 **NOTE:** Microsoft* Hyper-V* must be installed on the system in order for virtualization settings to be available. Without Hyper-V* being installed, the Virtualization tab in PROSet will not appear.

To set this using Windows PowerShell, find the first partition using the Get-IntelNetAdapter cmdlet. Once you know the port with partition number 0, use that port name with the Get-IntelNetAdapterSetting and Set-IntelNetAdapterSetting cmdlets.

Configuring NPAR in Linux

On Intel® 710 Series based adapters that support it, you can set up multiple functions on each physical port. You configure these functions through the System Setup/BIOS.

Minimum TX Bandwidth is the guaranteed minimum data transmission bandwidth, as a percentage of the full physical port link speed, that the partition will receive. The bandwidth the partition is awarded will never fall below the level you specify here.

The range for the minimum bandwidth values is:

1 to ((100 minus # of partitions on the physical port) plus 1)

For example, if a physical port has 4 partitions, the range would be

1 to ((100 - 4) + 1 = 97)

The Maximum Bandwidth percentage represents the maximum transmit bandwidth allocated to the partition as a percentage of the full physical port link speed. The accepted range of values is 1-100. The value can be used as a limiter, should you chose that any one particular function not be able to consume 100% of a port's bandwidth (should it be available). The sum of all the values for Maximum Bandwidth is not restricted, because no more than 100% of a port's bandwidth can ever be used.

-  **NOTE:**
- If the sum of the minimum bandwidth percentages does not equal 100, then settings will be automatically adjusted so that the sum equals 100.
 - If a partition's maximum bandwidth percentage is set lower than the partition's minimum bandwidth percentage, then the maximum bandwidth percentage will be automatically set to the value of the minimum bandwidth percentage.
 - When you attempt to set values for minimum bandwidth percentage via iDRAC with Lifecycle Controller using jobs that do not include the values for all enabled partitions, then the values seen after the jobs have completed may be different than the values that were supposed to be set. To avoid this issue, set the values for minimum bandwidth percentage on all partitions using a single job and make sure the sum of the values is 100.

Once the initial configuration is complete, you can set different bandwidth allocations on each function as follows:

1. Make a new directory named /config
2. Edit etc/fstab to include:


```
configs /config configs defaults
```
3. Load (or reload) the i40e driver
4. Mount /config
5. Make a new directory under config for each partition upon which you wish to configure the bandwidth.

Three files will appear under the config/partition directory:

```
- max_bw
- min_bw
- commit
```

Read from max_bw to get display the current maximum bandwidth setting.

Write to `max_bw` to set the maximum bandwidth for this function.

Read from `min_bw` to display the current minimum bandwidth setting.

Write to `min_bw` to set the minimum bandwidth for this function.

Write a '1' to commit to save your changes.



NOTES:

- commit is write only. Attempting to read it will result in an error.
- Writing to commit is only supported on the first function of a given port. Writing to a subsequent function will result in an error.
- Oversubscribing the minimum bandwidth is not supported. The underlying device's NVM will set the minimum bandwidth to supported values in an indeterminate manner. Remove all of the directories under `config` and reload them to see what the actual values are.
- To unload the driver you must first remove the directories created in step 5, above.

Example of Setting the minimum and maximum bandwidth (assume there are four functions on the port `eth6-eth9`, and that `eth6` is the first function on the port):

```
# mkdir /config/eth6
# mkdir /config/eth7
# mkdir /config/eth8
# mkdir /config/eth9
# echo 50 > /config/eth6/min_bw
# echo 100 > /config/eth6/max_bw
# echo 20 > /config/eth7/min_bw
# echo 100 > /config/eth7/max_bw
# echo 20 > /config/eth8/min_bw
# echo 100 > /config/eth8/max_bw
# echo 10 > /config/eth9/min_bw
# echo 25 > /config/eth9/max_bw
# echo 1 > /config/eth6/commit
```

Exiting NPar Mode

NPar mode is disabled in the System Setup menu during a reboot.

Reboot the system, and press the **F2** key to enter the **System Setup** menu. Select **Device Settings** from the list under **System Setup Main Menu**, then select your adapter from the list to get to the Device Configuration menu. Select **Device Level Configuration** in the list under **Main Configuration Page**. This brings up the Virtualization settings under **Device Level Configuration**.

In the Virtualization Mode list, select "None". Then click the **Back** button, which returns you to the Main Configuration Page. There, click the **Finish** button to save your change and reboot the system. When the system completes the reboot, NPar will no longer be active.



NOTE: When NPar has been disabled and the system completes the reboot, any other virtualization-related settings, such as NParEP or SR-IOV will also be disabled.

Performance Options

Optimizing Performance

You can configure Intel network adapter advanced settings to help optimize server performance.

Below the General Optimization section are sections that provide guidance for three server usage models:

- [Optimized for quick response and low latency](#) – useful for video, audio, and High Performance Computing Cluster (HPCC) servers
- [Optimized for throughput](#) – useful for data backup/retrieval and file servers
- [Optimized for CPU utilization](#) – useful for application, web, mail, and database servers



NOTES:

- Linux users, see [the Linux section of this guide](#) and the README file in the Linux driver package for Linux-specific performance enhancement details.
- The recommendations below are guidelines and should be treated as such. Additional factors such as installed applications, bus type, network topology, and operating system also affect system performance.
- These adjustments should be performed by a highly skilled network administrator. They are not guaranteed to improve performance. Not all settings shown here may be available through network driver configuration, operating system or system BIOS.
- When using performance test software, refer to the documentation of the application for optimal results.

General Optimization

- Install the adapter in an appropriate slot.



NOTE: Some PCIe x8 slots are actually configured as x4 slots. These slots have insufficient bandwidth for full line rate with some dual port devices. The driver can detect this situation and will write the following message in the system log: "PCI-Express bandwidth available for this card is not sufficient for optimal performance. For optimal performance a x8 PCI-Express slot is required." If this error occurs, moving your adapter to a true x8 slot will resolve the issue.

- In order for an Intel® 710 Series Network Adapter to reach its full potential, you must install it in a PCIe Gen3 x8 slot. Installing it in a shorter slot, or a Gen2 or Gen1 slot, will impact the throughput the adapter can attain.
- Use the proper cabling for your device.
- Increase the number of TCP and Socket resources from the default value. For Windows based systems, we have not identified system parameters other than the TCP Window Size which significantly impact performance.
- Increase the allocation size of Driver Resources (transmit/receive buffers). However, most TCP traffic patterns work best with the transmit buffer set to its default value, and the receive buffer set to its minimum value.

Jumbo Frames

Enabling jumbo frames may increase throughput. You must enable jumbo frames on all of your network components to get any benefit.

RSS Queues

If you have multiple 10 Gbps (or faster) ports installed in a system, the RSS queues of each adapter port can be adjusted to use non-overlapping sets of processors within the adapter's local NUMA Node/Socket. Change the RSS Base Processor Number for each adapter port so that the combination of the base processor and the max number of RSS processors settings ensure non-overlapping cores. For Microsoft Windows systems, do the following:

1. Identify the adapter ports to be adjusted and inspect at their `RssProcessorArray` using the `Get-NetAdapterRSS` PowerShell cmdlet.
2. Identify the processors with NUMA distance 0. These are the cores in the adapter's local NUMA Node/Socket and will provide the best performance.
3. Adjust the RSS Base processor on each port to use a non-overlapping set of processors within the local set of processors. You can do this manually or using the following PowerShell command:

```
Set-NetAdapterAdvancedProperty -Name <Adapter Name> -DisplayName "RSS Base Processor Number" -DisplayValue <RSS Base Proc Value>
```
4. Use the `Get-NetAdapterAdvancedproperty` cmdlet to check that the right values have been set:

```
Get-NetAdapterAdvancedproperty -Name <Adapter Name>
```

For Example: For a 4 port adapter with Local processors 0, 2, 4, 6, 8, 10, 12, 14, 16, 18, 20, 22, 24, 26, 28, 30, and 'Max RSS processor' of 8, set the RSS base processors to 0, 8, 16 and 24.

CPU Affinity

When passing traffic on multiple network ports using an I/O application that runs on most or all of the cores in your system, consider setting the CPU Affinity for that application to fewer cores. This should reduce CPU utilization and in some cases may increase throughput for the device. The cores selected for CPU Affinity must be local to the affected network device's Processor Node/Group. You can use the PowerShell command `Get-NetAdapterRSS` to list the cores that are local to a device. You may need to increase the number of cores assigned to the application to maximize throughput. Refer to your operating system documentation for more details on setting the CPU Affinity.

Optimized for quick response and low latency

- Minimize or disable Interrupt Moderation Rate.
- Disable Offload TCP Segmentation.
- Disable Jumbo Packets.
- Increase Transmit Descriptors.
- Increase Receive Descriptors.
- Increase RSS Queues.

Optimized for throughput

- Enable Jumbo Packets.
- Increase Transmit Descriptors.
- Increase Receive Descriptors.
- On systems that support NUMA, set the Preferred NUMA Node on each adapter to achieve better scaling across NUMA nodes.

Optimized for CPU utilization

- Maximize Interrupt Moderation Rate.
- Keep the default setting for the number of Receive Descriptors; avoid setting large numbers of Receive Descriptors.
- Decrease RSS Queues.
- In Hyper-V environments, decrease the Max number of RSS CPUs.

Performance Profile

Performance Profiles are supported on Intel® 10GbE adapters and allow you to quickly optimize the performance of your Intel® Ethernet Adapter. Selecting a performance profile will automatically adjust some Advanced Settings to their optimum setting for the selected application. For example, a standard server has optimal performance with only two RSS (Receive-Side Scaling) queues, but a web server requires more RSS queues for better scalability.

You must install Intel® PROSet to use Performance profiles.

Profiles	• Standard Server – This profile is optimized for typical servers. • Web Server – This profile is optimized for IIS and HTTP-based web servers. • Virtualization Server – This profile is optimized for Microsoft's Hyper-V virtualization environment. • Storage Server – This profile is optimized for Fibre Channel over Ethernet or for iSCSI over DCB performance. Selecting this profile will disable SR-IOV and VMQ. • Storage + Virtualization – This profile is optimized for a combination of storage and virtualization requirements. • Low Latency – This profile is optimized to minimize network latency.
-----------------	--



NOTES:

- Not all options are available on all adapter/operating system combinations.
- If you have selected the Virtualization Server profile or the Storage + Virtualization profile, and you uninstall the Hyper-V role, you should select a new profile.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Profile" -DisplayValue "Standard Server"
```

Power Options

The Intel PROSet Power Management tab includes several settings that control the adapter's power consumption. For example, you can set the adapter to reduce its power consumption if the cable is disconnected.

This setting is found on the Power Management tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Wake on Link Settings" -DisplayValue "Enabled"
```

Reduce Power if Cable Disconnected & Reduce Link Speed During Standby

Enables the adapter to reduce power consumption when the LAN cable is disconnected from the adapter and there is no link. When the adapter regains a valid link, adapter power usage returns to its normal state (full power usage).

The Hardware Default option is available on some adapters. If this option is selected, the feature is disabled or enabled based on the system hardware.

Default	The default varies with the operating system and adapter.
Range	The range varies with the operating system and adapter.

Energy Efficient Ethernet

The Energy Efficient Ethernet (EEE) feature allows a capable device to enter Low-Power Idle between bursts of network traffic. Both ends of a link must have EEE enabled for any power to be saved. Both ends of the link will resume full power when data needs to be transmitted. This transition may introduce a small amount of network latency.



NOTES:

- Both ends of the EEE link must automatically negotiate link speed.
- EEE is not supported on every adapter.

Wake on LAN Options

The ability to remotely wake computers is an important development in computer management. This feature has evolved over the last few years from a simple remote power-on capability to a complex system interacting with a variety of device and operating system power states.

Microsoft Windows Server is ACPI-capable. Windows does not support waking from a power-off (S5) state, only from standby (S3) or hibernate (S4). When shutting down the system, these states shut down ACPI devices, including Intel adapters. This disarms the adapter's remote wake-up capability. However, in some ACPI-capable computers, the BIOS may have a setting that allows you to override the operating system and wake from an S5 state anyway. If there is no support for wake from S5 state in your BIOS settings, you are limited to Wake From Standby when using these operating systems in ACPI computers.

The Intel PROSet Power Management tab includes **Wake on Magic Packet** and **Wake on directed packet settings**. These control the type of packets that wake up the system from standby.

For some adapters, the Power Management tab in Intel PROSet includes a setting called **Wake on Magic Packet from power off state**. Enable this setting to explicitly allow wake-up with a Magic Packet* from shutdown under APM power management mode.



NOTES:

- To use the Wake on Directed Packet feature, WoL must first be enabled in the EEPROM using BootUtil.
- If **Reduce speed during standby** is enabled, then **Wake on Magic Packet** and/or **Wake on directed packet** must be enabled. If both of these options are disabled, power is removed from the adapter during standby.
- Wake on Magic Packet from power off state** has no effect on this option.

WoL Supported Devices

[All devices](#) support Wake on LAN on all ports, with the following exceptions:

Device	Adapter Port(s) supporting WoL
Intel® Gigabit 2P I350-t Adapter Intel® Gigabit 4P I350-t Adapter	port 1 only
Intel® Ethernet Converged Network Adapter X710-4 Intel® Ethernet Converged Network Adapter X710-2 Intel® Ethernet Converged Network Adapter X710	port 1 only
Intel® Ethernet 25G 2P E810-XXV Adapter Intel® Ethernet 25G 2P XXV710 Mezz	Not supported

Device	Adapter Port(s) supporting WoL
Intel® Ethernet 25G 2P XXV710 Adapter	
Intel® Ethernet Converged Network Adapter X710-T Intel® Ethernet Converged Network Adapter XL710-Q2	Not supported
Intel® Ethernet 10G 2P X710-T2L-t Adapter Intel® Ethernet 10G 4P X710-T4L-t Adapter	Not Supported
Intel® Ethernet 10G 2P X520 Adapter Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz	Not Supported
Intel® Ethernet 10G 2P X540-t Adapter	Not Supported
Intel® Ethernet 10G 2P X550-t Adapter	Not Supported
Intel® Ethernet 25G 2P XXV710 Mezz	This device only supports waking from a powered off (S5) state. It does not support waking from sleep/hibernate (S3/S4).

Wake on Link Settings

Wakes the computer if the network connection establishes link while the computer is in standby mode. You can enable the feature, disable it, or let the operating system use its default.



NOTES:

- If a copper-based Intel adapter is advertising a speed of one gigabit only, this feature does not work because the adapter cannot identify a gigabit link at a D3 state.
- The network cable must be disconnected when entering into S3/S4 in order to wake the system up by link up event.

Default	Disabled
Range	Disabled OS Controlled Forced

Remote Wake-Up

Remote wake-up can wake your server from a low power or powered off state. If Wake On LAN is enabled, when your system is powered down, the network interface draws standby power and listens for specially designed packet. If it receives such a packet it will wake your system.

Advanced Configuration and Power Interface (ACPI)

ACPI supports a variety of power states. Each state represents a different level of power, from fully powered up to completely powered down, with partial levels of power in each intermediate state.

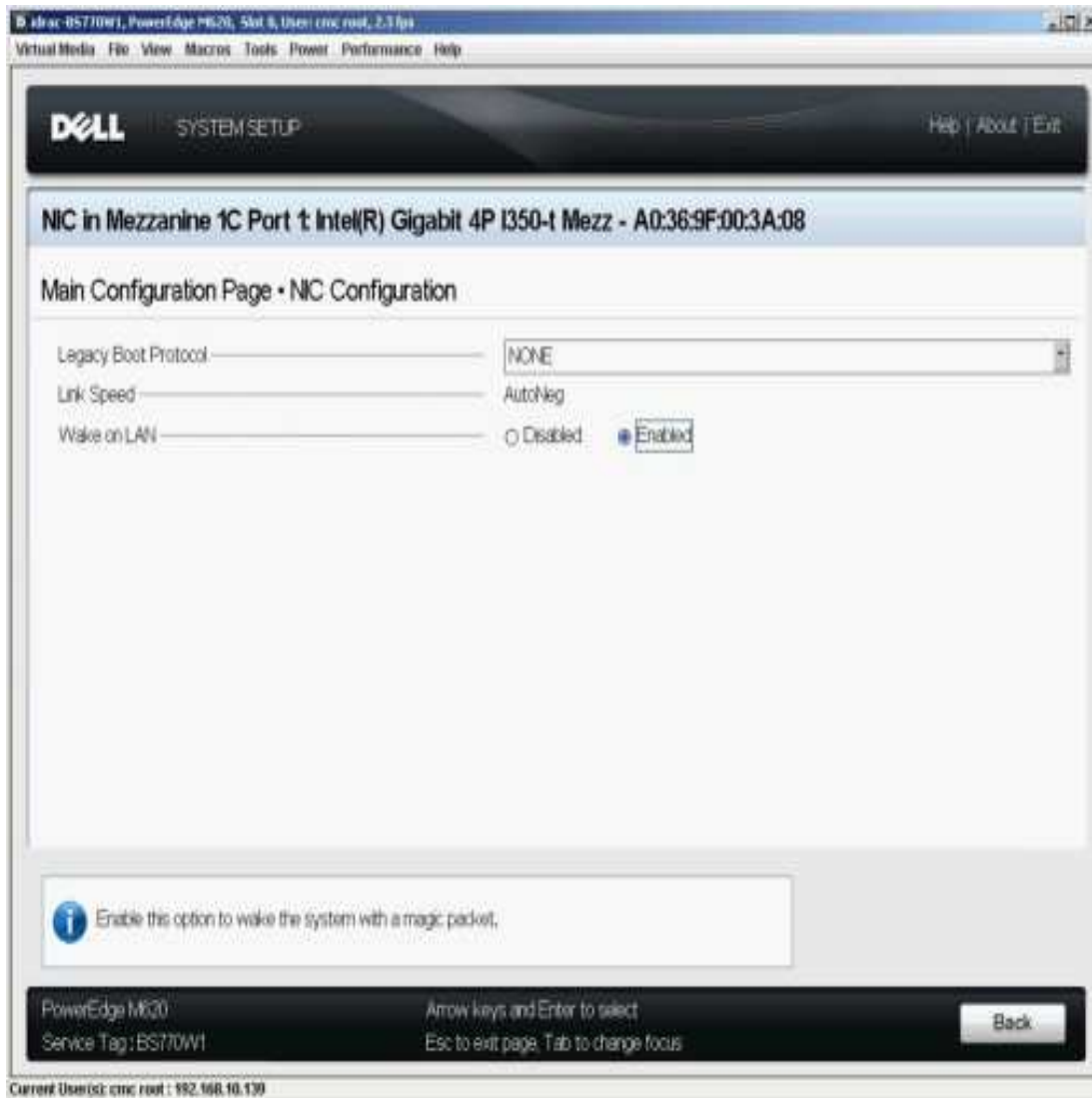
ACPI Power States

Power State	Description
S0	On and fully operational
S1	System is in low-power mode (sleep mode). The CPU clock is stopped, but RAM is powered on and being refreshed.
S2	Similar to S1, but power is removed from the CPU.
S3	Suspend to RAM (standby mode). Most components are shut down. RAM remains operational.
S4	Suspend to disk (hibernate mode). The memory contents are swapped to the disk drive and then reloaded into RAM when the system is awakened.
S5	Power off

Enabling Wake From Power Off

If you want to wake your system from a power off state, you must enable it from the System Setup.

1. Go to System Setup.
2. Choose a port and go to configuration.
3. Enable Wake on LAN.



Wake-Up Address Patterns

Remote wake-up can be initiated by a variety of user selectable packet types and is not limited to the Magic Packet format. For more information about supported packet types, see the [operating system settings](#) section.

The wake-up capability of Intel adapters is based on patterns sent by the OS. You can configure the driver to the following settings using Intel® PROSet. For Linux*, WoL is provided through the ethtool* utility. For more information on ethtool, see the following Web site: <http://sourceforge.net/projects/gkernel>.

- Wake on Directed Packet - accepts only patterns containing the adapter's Ethernet address in the Ethernet header or containing the IP address, assigned to the adapter, in the IP header.
- Wake on Magic Packet - accept only patterns containing 16 consecutive repetitions of the adapter's MAC address.
- Wake on Directed Packet and Wake on Magic Packet - accepts the patterns of both directed packets and magic packets.

Choosing "Wake on directed packet" will also allow the adapter to accept patterns of the Address Resolution Protocol (ARP) querying the IP address assigned to the adapter. If multiple IP addresses are assigned to an adapter, the operating system may request to wake up on ARP patterns querying any of the assigned addresses. However, the adapter will only awaken in response to ARP packets querying the first IP address in the list, usually the first address assigned to the adapter.

Physical Installation Issues

Slot

Some motherboards will only support remote wake-up (or remote wake-up from S5 state) in a particular slot. See the documentation that came with your system for details on remote wake-up support.

Power

Some Intel PRO adapters are 3.3 volt and some are 12 volt. They are keyed to fit either type of slot.

The 3.3 volt standby supply must be capable of supplying at least 0.2 amps for each Intel PRO adapter installed. Turning off the remote wake-up capability on the adapter using the BootUtil utility reduces the power draw to around 50 milliamps (.05 amps) per adapter.

Operating System Settings

Microsoft Windows Operating Systems

Windows Server is ACPI-capable. These operating systems do not support remote wake-up from a powered off state (S5), only from standby. When shutting down the system, they shut down ACPI devices including the Intel Ethernet adapters. This disarms the adapters' remote wake-up capability. However, in some ACPI-capable computers, the BIOS may have a setting that allows you to override the OS and wake from an S5 state anyway. If there is no support for wake from S5 state in your BIOS settings, you are limited to wake from standby when using these operating systems in ACPI computers.

The **Power Management** tab in Intel PROSet includes a setting called Wake on Magic Packet from power off state for some adapters. To explicitly allow wake-up with a Magic Packet from shutdown under APM power management mode, check this box to enable this setting. See Intel PROSet help for more details.

In ACPI-capable versions of Windows, the Intel PROSet advanced settings include a setting called Wake on Settings. This setting controls the type of packets that wake the system from standby. See Intel PROSet help for more details.

If you do not have Intel PROSet installed you will need to do the following:

1. Open the Device Manager or Intel® PROSet Adapter Configuration Utility , then navigate to the **Power Management** tab, check "**Allow this device to bring the computer out of standby.**"
2. On the **Advanced** tab, set the "**Wake on Magic packet**" option to Enabled.

In order to wake from S5 without Intel PROSet, on the **Advanced tab**, set "**Enable PME**" to Enabled.

Other Operating Systems

Remote Wake-Up is also supported in [Linux](#).

Priority & VLAN Tagging

Enables the adapter to offload the insertion and removal of priority and VLAN tags for transmit and receive.

Default	Priority & VLAN Enabled
Range	• Priority & VLAN Disabled • Priority Enabled • VLAN Enabled • Priority & VLAN Enabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To set this in Windows Powershell, first disable DCB, then set priority and VLAN tagging. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "DCB" -DisplayValue "Disabled"
```

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Packet Priority & VLAN" -DisplayValue "VLAN Enabled"
```

PTP Hardware Timestamp

Allows applications that use PTPv2 (Precision Time Protocol) to use hardware generated timestamps to synchronize clocks throughout your network. If this setting is enabled, it takes precedence over the [Software Timestamp](#) setting.

Default	Disabled
Range	• Enabled • Disabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "PTP Hardware Timestamp" -DisplayValue "Enabled"
```

Quality of Service

Quality of Service (QoS) allows the adapter to send and receive IEEE 802.3ac tagged frames. 802.3ac tagged frames include 802.1p priority-tagged frames and 802.1Q VLAN-tagged frames. In order to implement QoS, the adapter must be connected to a switch that supports and is configured for QoS. Priority-tagged frames allow programs that deal with real-time events to make the most efficient use of network bandwidth. High priority packets are processed before lower priority packets.

To implement QoS, the adapter must be connected to a switch that supports and is configured for 802.1p QoS.

QoS Tagging is enabled and disabled in the **Advanced** tab of Intel PROSet for Windows Device Manager.

To set this in Windows Powershell, first disable DCB, then set QoS using the Priority and VLAN tagging DisplayName in the cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "DCB" -DisplayValue "Disabled"
```

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Packet Priority & VLAN" -DisplayValue "VLAN Enabled"
```

Once QoS is enabled, you can specify levels of priority based on IEEE 802.1p/802.1Q frame tagging.

Remote Direct Memory Access (RDMA)

Remote Direct Memory Access, or RDMA, allows a network device to transfer data directly to and from application memory on another system, increasing throughput and lowering latency in certain networking environments.

- Intel® Ethernet Controller 800 Series devices support both iWARP and RoCEv2.

The major difference is that iWARP performs RDMA over TCP, while RoCEv2 uses UDP.

To avoid performance degradation from dropped packets, enable link level flow control or priority flow control on all network interfaces and switches.

 NOTES:

- On systems running a Microsoft Windows Server operating system, enabling *QoS/priority flow control will disable link level flow control.
- Devices based on the Intel® Ethernet Controller 800 Series do not support RDMA when operating in multiport mode with more than 4 ports.
- On Linux systems, RDMA and bonding are not compatible. If RDMA is enabled, bonding will not be functional.

RDMA on Linux

These basic Linux RDMA installation instructions apply for the following devices:

- Intel® Ethernet Controller 800 Series

For detailed installation and configuration information, see the Linux RDMA driver README file in the driver tarball for Intel Ethernet devices that support RDMA.

This example is specific to Red Hat Enterprise Linux. Your operating system specifics may be different.

1. Untar the RDMA driver tarball. For example:

```
# tar xzf irdma-<x.x.x>.tar.gz
```

2. Install the RDMA PF driver. For example:

```
# cd irdma-<x.x.x>
```

```
# ./build.sh
```

```
# modprobe irdma
```



NOTE: By default, the irdma driver is loaded in iWARP mode. It uses the devlink interface to enable RoCEv2 per port. To load all irdma ports in RoCEv2 mode, use the following:

```
# modprobe irdma roce_ena=1
```

3. Uninstall any previously installed version of rdma-core.

```
# yum erase rdma-core
```

4. Download the rdma-core library tarball from GitHub. For example:

```
# wget https://github.com/linux-rdma/rdma-core/releases/download/v27.0/rdma-core-27.0.tar.gz
```



NOTE: Download the rdma-core version that matches the version of the libirdma patch file included with the driver. For example, libirdma-27.0.patch requires rdma-core-27.0.tar.gz.

5. Untar the rdma-core tarball, apply the patch, and change group for the spec file. For example:

```
# tar -xzvf rdma-core-<version>.tar.gz
```

```
# cd rdma-core-<version>
```

```
# patch -p2 <<path-to-component-build>/libirdma-<version>.patch
```

```
# cd ..
```

```
# chgrp -R root <path-to-rdma-core>/redhat
```

```
# tar -zcvf rdma-core-<version>.tgz rdma-core-<version>
```

6. Install rdma-core-intel user space libraries. For example:

```
# mkdir -p ~/rpmbuild/SOURCES
```

```
# mkdir -p ~/rpmbuild/SPECS
```

```
# cp rdma-core-<version>.tgz ~/rpmbuild/SOURCES/
```

```
# cd ~/rpmbuild/SOURCES
```

```
# tar -xzvf rdma-core-<version>.tgz
```

```
# cp ~/rpmbuild/SOURCES/rdma-core-<version>/redhat/rdma-core.spec ~/rpmbuild/SPECS/
```

```
# cd ~/rpmbuild/SPECS/
```

```
# rpmbuild -ba rdma-core.spec
```

```
# cd ~/rpmbuild/RPMS/x86_64
```

```
# yum install *<version>*.rpm
```

7. Enable flow control on your adapter. You can enable link-level or priority flow control; we recommend using priority flow control. See the Linux RDMA driver README for more information on flow control.
8. Enable flow control on the switch your system is connected to. See your switch documentation for details.

RDMA for Microsoft Windows Network Direct (ND) User-Mode Applications

Network Direct (ND) allows user-mode applications to use RDMA features.



NOTE: User mode applications may have prerequisites such as Microsoft HPC Pack or Intel MPI Library, refer to your application documentation for more details.

RDMA User Mode Installation

The Intel® Ethernet User Mode RDMA Provider is supported on Microsoft Windows Server 2016 and later.

Follow the steps below to install user-mode Network Direct features.

1. From the installation media, run Autorun.exe to launch the installer, then choose "Install Drivers and Software" and accept the license agreement.
2. On the Setup Options screen, select "Intel® Ethernet User Mode RDMA Provider".
3. On the RDMA Configuration Options screen, select "Enable RDMA routing across IP Subnets" if desired. Note that this option is displayed during base driver installation even if user mode RDMA was not selected, as this option is applicable to Network Direct Kernel functionality as well.
4. If Windows Firewall is installed and active, select "Create an Intel® Ethernet RDMA Port Mapping Service rule in Windows Firewall" and the networks to which to apply the rule.



NOTE: If Windows Firewall is disabled or you are using a third party firewall, you will need to add this rule manually.

5. Continue with driver and software installation.

RDMA Network Direct Kernel (NDK)

RDMA Network Direct Kernel (NDK) functionality is included in the Intel base networking drivers and requires no additional features to be installed.

RDMA Routing Across IP Subnets

If you want to allow NDK's RDMA functionality across subnets, you will need to select "Enable RDMA routing across IP Subnets" on the RDMA Configuration Options screen during base driver installation.

Enabling Priority Flow Control (PFC) on a Microsoft Windows Server Operating System

To avoid performance degradation from dropped packets enable priority flow control (PFC) or link level flow control on all network interfaces and switches.



NOTE: On systems running a Microsoft Windows Server operating system, enabling *QoS/priority flow control will disable link level flow control.

Use the following PowerShell commands to enable PFC on a Microsoft Windows Server Operating System:

```
Install-WindowsFeature -Name Data-Center-Bridging -IncludeManagementTools
New-NetQoSPolicy "SMB" -NetDirectPortMatchCondition 445 -PriorityValue8021Action 3
Enable-NetQoSFlowControl -Priority 3
Disable-NetQoSFlowControl -Priority 0,1,2,4,5,6,7
New-NetQoSTrafficClass -Name "SMB" -Priority 3 -BandwidthPercentage 60 -Algorithm ETS
Set-NetQoSDbxSetting -Willing $FALSE
Enable-NetAdapterQos -Name "Slot1 4 2 Port 1"
```

Verifying RDMA operation with Microsoft PowerShell

You can check that RDMA is enabled on the network interfaces using the following Microsoft PowerShell command:

```
Get-NetAdapterRDMA
```

Use the following PowerShell command to check if the network interfaces are RDMA capable and multichannel is enabled:

```
Get-SmbClientNetworkInterface
```

Use the following PowerShell command to check if Network Direct is enabled in the operating system:

```
Get-NetOffloadGlobalSetting | Select NetworkDirect
```

Use netstat to make sure each RDMA-capable network interface has a listener at port 445 (Windows Client OSs that support RDMA may not post listeners). For example:

```
netstat.exe -xan | ? {$_ -match "445"}
```

RDMA for Virtualized Environments

To enable RDMA functionality on virtual adapter(s) connected to a VMSwitch, the SRIOV (Single Root IO Virtualization) and VMQ (Virtual Machine Queues) advanced properties must be enabled on each port. Under certain circumstances, these settings may be disabled by default. These options can be set manually in the advanced tab of the adapter properties dialog box, or with the following Powershell commands:

```
Set-NetAdapterAdvancedProperty -Name <nic_name> -RegistryKeyword *SRIOV -RegistryValue 1
```

```
Set-NetAdapterAdvancedProperty -Name <nic_name> -RegistryKeyword *VMQ -RegistryValue 1
```

Configuring RDMA Guest Support (NDK Mode 3)

NDK Mode 3 allows kernel mode Windows components to use RDMA features inside Hyper-V guest partitions. To enable NDK mode 3 on an Intel Ethernet device, do the following:

1. Enable SR-IOV in your system's BIOS or uEFI.
2. Enable the SR-IOV advanced setting on the device.
3. Enable SR-IOV on the VMSwitch bound to the device by performing the following for all physical functions on the same device:

```
New-VMSwitch -Name <switch_name> -NetAdapterName <device_name> -EnableIov $true
```
4. Configure the number of RDMA virtual functions (VFs) on the device by setting the "RdmaMaxVfsEnabled" advanced setting. All physical functions must be set to the same value. The value is the maximum number of VFs that can be capable of RDMA at one time for the entire device. Enabling more VFs will restrict RDMA resources from physical functions (PFs) and other VFs.

```
Set-NetAdapterAdvancedProperty -Name <device_name> -RegistryKeyword RdmaMaxVfsEnabled -RegistryValue <Value: 0 - 32>
```
5. Disable all PF adapters on the host and re-enable them. This is required when the registry keyword "RdmaMaxVfsEnabled" is changed or when creating or destroying a VMSwitch.

```
Get-NetAdapterRdma | Disable-NetAdapter
Get-NetAdapterRdma | Enable-NetAdapter
```
6. Create VM Network Adapters for VMs that require RDMA VF support.

```
Add-VMNetworkAdapter -VMName <vm_name> -VMNetworkAdapterName <device_name> -SwitchName <switch_name>
```
7. If you plan to use Microsoft Windows 10 Creators Update (RS2) or later on a guest partition, set the RDMA weight on the VM Network Adapter by entering the following command on the host:

```
Set-VMNetworkAdapterRdma -VMName <vm_name> -VMNetworkAdapterName <device_name> -RdmaWeight 100
```
8. Set SR-IOV weight on the VM Network Adapter (Note: SR-IOV weight must be set to 0 before setting the RdmaWeight to 0):

```
Set-VMNetworkAdapter -VMName <vm_name> -VMNetworkAdapterName <device_name> -IovWeight 100
```
9. Install the VF network adapter with the PROSET Installer in the VM.

10. Enable RDMA on the VF driver and Hyper-V Network Adapter using PowerShell in the VM:

```
Set-NetAdapterAdvancedProperty -Name <device_name> -RegistryKeyword RdmaVfEnabled -
RegistryValue 1
Get-NetAdapterRdma | Enable-NetAdapterRdma
```

RDMA for NDK Features such as SMB Direct (Server Message Block)

NDK allows Windows components (such as SMB Direct storage) to use RDMA features.

Testing NDK: Microsoft Windows SMB Direct with DiskSPD

This section outlines the recommended way to test RDMA for Intel Ethernet functionality and performance on Microsoft Windows operating systems.

Note that since SMB Direct is a storage workload, the performance of the benchmark may be limited to the speed of the storage device rather than the network interface being tested. Intel recommends using the fastest storage possible in order to test the true capabilities of the network device(s) under test.

Test instructions:

1. Set up and connect at least two servers running a supported Microsoft Windows Server operating system, with at least one RDMA-capable Intel® Ethernet device per server.
2. On the system designated as the SMB server, set up an SMB share. Note that the performance of the benchmark may be limited to the speed of the storage device rather than the network interface being tested. Storage setup is outside of the scope of this document. You can use the following PowerShell command:

```
New-SmbShare -Name <SMBsharename> -Path <SMBsharefilepath> -FullAccess <domain-name>\Administrator,Everyone
```

For Example:

```
New-SmbShare -Name RAMDISKShare -Path R:\RAMDISK -FullAccess group\Administrator,Everyone
```

3. Download and install the Diskspd Microsoft utility from here: <https://gallery.technet.microsoft.com/DiskSpd-a-robust-storage-6cd2f223>
4. Using CMD or Powershell, cd to the DiskSpd folder and run tests. (Refer to Diskspd documentation for more details on parameters)

For Example: Set the block size to 4K, run the test for 60 seconds, disable all hardware and software caching, measure and display latency statistics, leverage 16 overlapped IOs and 16 threads per target, random 0% writes and 100% reads and create a 10GB test file at "\\<SMBserverTestIP>\<SMBsharename>\test.dat" :

```
.\diskspd.exe -b4K -d60 -h -L -o16 -t16 -r -w0 -c10G \\<SMBserver-TestIP>\<SMBsharename>\test.dat
```

5. Verify that RDMA traffic is running using perfmon counters such as "RDMA Activity" and "SMB Direct Connection". Refer to Microsoft documentation for more details.

Accessing Remote NVM Express* drives using RDMA

RDMA provides a high throughput, low latency means to directly access NVM Express* (NVMe*) drives on a remote server. See the *NVM Express* over TCP Using Intel® Ethernet Configuration Guide* and the *NVM Express* over Fabrics Using Intel® Ethernet RDMA Configuration Guide* for details on how to setup and configure your server and client systems. Both guides are available on the [Intel Technical Library](#).

Receive Buffers

Defines the number of Receive Buffers, which are data segments. They are allocated in the host memory and used to store the received packets. Each received packet requires at least one Receive Buffer, and each buffer uses 2KB of memory.

You might choose to increase the number of Receive Buffers if you notice a significant decrease in the performance of received traffic. If receive performance is not an issue, use the default setting appropriate to the adapter.

Default	512, for all adapters.
Range	128-4096, in intervals of 64, for all adapters.
Recommended Value	Using IPSec and/or multiple features: 352

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Receive Buffers" -DisplayValue "256"
```

Receive Side Scaling

When Receive Side Scaling (RSS) is enabled, all of the receive data processing for a particular TCP connection is shared across multiple processors or processor cores. Without RSS all of the processing is performed by a single processor, resulting in less efficient system cache utilization.

LAN RSS

LAN RSS applies to a particular TCP connection.



NOTE: This setting has no effect if your system has only one processing unit.

LAN RSS Configuration

If your adapter does not support RSS, or if the SNP or SP2 is not installed, the RSS setting will not be displayed. If RSS is supported in your system environment, the following will be displayed:

- **Port NUMA Node.** This is the NUMA node number of a device.
- **Starting RSS CPU.** This setting allows you to set the preferred starting RSS processor. Change this setting if the current processor is dedicated to other processes. The setting range is from 0 to the number of logical CPUs - 1.
- **Max number of RSS CPU.** This setting allows you to set the maximum number of CPUs assigned to an adapter and is primarily used in a Hyper-V environment. By decreasing this setting in a Hyper-V environment, the total number of interrupts is reduced which lowers CPU utilization. The default is 8 for Gigabit adapters and 16 for 10 Gigabit, or faster, adapters.
- **Preferred NUMA Node.** This setting allows you to choose the preferred NUMA (Non-Uniform Memory Access) node to be used for memory allocations made by the network adapter. In addition the system will attempt to use the CPUs from the preferred NUMA node first for the purposes of RSS. On NUMA platforms, memory access latency is dependent on the memory location. Allocation of memory from the closest node helps improve performance. The Windows Task Manager shows the NUMA Node ID for each processor.



NOTES:

- This setting only affects NUMA systems. It will have no effect on non-NUMA systems.
- Choosing a value greater than the number of NUMA nodes present in the system selects the NUMA node closest to the device.
- **Receive Side Scaling Queues.** This setting configures the number of RSS queues, which determine the space to buffer transactions between the network adapter and CPU(s).

Default	2 queues for the Intel® 10 Gigabit Server Adapters
Range	• 1 queue is used when low CPU utilization is required. • 2 queues are used when good throughput and low CPU utilization are required.

- 4 or more queues are used for applications that demand maximum throughput and transactions per second.

**NOTES:**

- Not all settings are available on all adapters.
- 8, or more, queues are only available when PROSet for Windows Device Manager or Intel® PROSet Adapter Configuration Utility (Intel® PROSet ACU) is installed. If PROSet is not installed, only 4 queues are available.
- Using 8 or more queues requires the system to reboot.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Receive Side Scaling" -DisplayValue "Enabled"
```

Setting Speed and Duplex

Overview

The Link Speed and Duplex setting lets you choose how the adapter sends and receives data packets over the network.

In the default mode, an Intel network adapter using copper connections will attempt to auto-negotiate with its link partner to determine the best setting. If the adapter cannot establish link with the link partner using auto-negotiation, you may need to manually configure the adapter and link partner to the identical setting to establish link and pass packets. This should only be needed when attempting to link with an older switch that does not support auto-negotiation or one that has been forced to a specific speed or duplex mode.

Auto-negotiation is disabled by selecting a discrete speed and duplex mode in the adapter properties sheet. The settings available when auto-negotiation is disabled are dependent on your device. Not all speeds are available on all devices. Your link partner must match the setting you choose.

**NOTES:**

- Only experienced network administrators should force speed and duplex manually.
- When an adapter is running in NPar mode, Speed settings are limited to the root partition of each port.
- Fiber-based adapters operate only in full duplex at their native speed. You cannot change the speed or duplex of Intel adapters that use fiber cabling.
- Some devices may list 10 Mbps and 100 Mbps in full or half duplex as options. Using those settings is not recommended.
- The Link Speed tab may display a blue informational icon with a mouse-over message "This device is not linked at its maximum capable speed". In that case, if your device is set to auto-negotiate, you can adjust the speed of the device's link partner to the device's maximum speed. If the device is not set to auto-negotiate, you can adjust the device's speed manually, but you must ensure the link partner is set at the same speed.

Manually Configuring Duplex and Speed Settings



CAUTION: The settings at the switch must always match the adapter settings. Adapter performance may suffer, or your adapter might not operate correctly if you configure the adapter differently from your switch.

Configuration is specific to your operating system driver. To set a specific Link Speed and Duplex mode, refer to the section below that corresponds to your operating system.

Windows

The default setting is for auto-negotiation to be enabled. Only change this setting to match your link partner's speed and duplex setting if you are having trouble connecting.

1. In Windows Device Manager or the Intel® PROSet Adapter Configuration Utility, double-click the adapter you want to configure.
2. On the **Link Speed** tab, select a speed and duplex option from the **Speed and Duplex** drop-down menu.
3. Click **OK**.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Speed & Duplex" -DisplayValue "Auto Negotiation"
```

Linux

See [Linux* Driver for the Intel® Gigabit Family of Adapters](#) for information on configuring Speed and Duplex on Linux systems.

Software Timestamp

Allows applications that use PTPv2 (Precision Time Protocol) to use software generated timestamps to synchronize clocks throughout your network. If the [PTP Hardware Timestamp](#) setting is enabled, it takes precedence over this setting.

Default	Disabled
Range	• Disabled • RxAll • TxAll • RxAll & TxAll • TaggedTx • RxAll & TaggedTx

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "PTP Hardware Timestamp" -DisplayValue "Enabled"
```

SR-IOV (Single Root I/O Virtualization)

SR-IOV lets a single network port appear to be several virtual functions in a virtualized environment. If you have an SR-IOV capable NIC, each port on that NIC can assign a virtual function to several guest partitions. The virtual functions bypass the Virtual Machine Manager (VMM), allowing packet data to move directly to a guest partition's memory, resulting in higher throughput and lower CPU utilization. SR-IOV also allows you to move packet data directly to a guest partition's memory. See your operating system documentation for system requirements.

For devices that support it, SR-IOV is enabled in the host partition. Some devices may need to have SR-IOV enabled in a pre-boot environment.



NOTES:

- **Configuring SR-IOV for improved network security:** In a virtualized environment, on Intel® Server Adapters that support SR-IOV, the virtual function (VF) may be subject to malicious behavior. Software-generated layer two frames, like IEEE 802.3x (link flow control), IEEE 802.1Qbb (priority based flow-control), and others of this type, are not expected and can throttle traffic between the host and the virtual switch, reducing performance.

To resolve this issue, and to ensure isolation from unintended traffic streams, configure all SR-IOV enabled ports for VLAN tagging from the administrative interface on the PF. This configuration allows unexpected, and potentially malicious, frames to be dropped.

- SR-IOV must be enabled in the BIOS.
- You must enable VMQ for SR-IOV to function.
- For best performance, on the host use 'Set-VMNetworkAdapter -lovQueuePairsRequested 4' on the VF to allow the virtual network to use 4 queues (maximum supported value) and assign 4 or more virtual CPUs to the connected VM. In the VM, set 'Maximum number of Receive Queues' in the VF's adapter properties to 4.
- Binding more than two virtual functions (VFs) to a virtual machine (VM) is not recommended. Binding more VFs to a VM may cause system instability.
- VMWare ESXi does not support SR-IOV on 1GbE ports.
- Some multiport adapters contain more than one controller. On these adapters, enabling SR-IOV on a port will not enable SR-IOV on all ports. Only ports bound to the same controller will be enabled.
- If SR-IOV is disabled in BIOS or the Boot Manager, enabling SR-IOV from Intel PROSet will require a system reboot.
- When an adapter is running in NPar mode, SR-IOV is limited to the root partition of each port.
- When an adapter is running in NPar mode, virtualization (SR-IOV) settings apply to all ports on the adapter. Changes made the virtualization settings on one port are applied to all ports on the adapter.
- Due to chipset limitations, not all systems or slots support SR-IOV. Below is a chart summarizing SR-IOV support on Dell EMC server platforms.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "SR-IOV" -DisplayValue "Enabled"
```

Enabling SR-IOV on the Server

You must enable SR-IOV in the system's BIOS or HII.

BIOS

1. Enter the system BIOS at POST.
2. Enable Global SR-IOV.
3. Enable Virtualization Technology.
4. Save the changes and exit.

HII

1. During POST, press F2 to enter Device Settings.
2. Navigate to NIC -> Device Level Settings.
3. Set Virtualization Mode to either "SR-IOV" or "NPAR + SR-IOV"
4. Save the changes and exit.

SR-IOV Support on Network Adapters

NDC, LOM, or Adapter	40Gbe	25Gbe	10Gbe	1Gbe
Intel® Ethernet Converged Network Adapter XL710-Q2	Yes			
Intel® Ethernet 40G 2P XL710 QSFP+ rNDC	Yes			
Intel® Ethernet 25G 2P E810-XXV Adapter		Yes		

NDC, LOM, or Adapter	40Gbe	25Gbe	10Gbe	1Gbe
Intel® Ethernet 25G 2P E810-XXV OCP		Yes		
Intel® Ethernet 25G 2P XXV710 Mezz		Yes		
Intel® Ethernet 25G 2P XXV710 Adapter		Yes		
Intel® Ethernet 10G 4P X710-k bNDC			Yes	
Intel® Ethernet 10G 2P X710-k bNDC			Yes	
Intel® Ethernet 10G X710-k bNDC			Yes	
Intel® Ethernet 10G 2P X710-T2L-t Adapter			Yes	
Intel® Ethernet 10G 4P X710-T4L-t Adapter			Yes	
Intel® Ethernet Network Adapter X710-TL			Yes	
Intel® Ethernet Converged Network Adapter X710			Yes	
Intel® Ethernet Converged Network Adapter X710-T			Yes	
Intel® Ethernet 10G 4P X710-T4L-t OCP			Yes	
Intel® Ethernet 10G 2P X710-T2L-t OCP			Yes	
Intel® Ethernet 10G 2P X710 OCP			Yes	
Intel® Ethernet 10G 4P X710 OCP			Yes	
Intel® Ethernet Server Adapter X710-DA2 for OCP			Yes	
Intel® Ethernet 10G 4P X710/I350 rNDC			Yes	No
Intel® Ethernet 10G 4P X710 SFP+ rNDC			Yes	
Intel® Ethernet 10G X710 rNDC			Yes	No
Intel® Ethernet 10G 4P X550 rNDC			Yes	
Intel® Ethernet 10G 4P X550/I350 rNDC			Yes	No
Intel® Ethernet 10G 2P X550-t Adapter			Yes	
Intel® Ethernet 10G 2P X540-t Adapter			Yes	
Intel® Ethernet 10G 4P X540/I350 rNDC			Yes	No
Intel® Ethernet 10G 4P X520/I350 rNDC			Yes	No
Intel® Ethernet 10G 2P X520-k bNDC			Yes	
Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz			Yes	
Intel® Ethernet 1G 4P I350-t OCP				Yes
Intel® Gigabit 4P I350-t rNDC				Yes
Intel® Gigabit 4P I350 bNDC				Yes
Intel® Gigabit 4P I350-t Mezz				Yes
Intel® Gigabit 2P I350-t Adapter				Yes
Intel® Gigabit 4P I350-t Adapter				Yes

NDC, LOM, or Adapter	40Gbe	25Gbe	10Gbe	1Gbe
PowerEdge C4130 LOMs				No
PowerEdge C6320 LOMs			Yes	
PowerEdge C6420 LOMs				No
PowerEdge T620 LOMs				No
PowerEdge T630 LOMs				No
PowerEdge FC430 LOMs			No	Yes
PowerEdge R530XD LOMs				No

Dell EMC Platform		OCP Mezz	Rack NDC	PCI Express Slot												
				1	2	3	4	5	6	7	8	9	10	11	12	13
C4130				yes	yes											
C4140			no	yes	no	yes										
C6320				yes												
C6420		yes		yes												
R230				no	no											
R240				no	no											
R320				no	yes											
R330				no	no											
R340				no	no											
R420	1 x CPU			no	yes											
	2 x CPU			yes	yes											
R430				yes	yes											
R440				yes	yes	yes										
R520	1 x CPU			no	yes	yes	yes									
	2 x CPU			yes	yes	yes	yes									
R530				yes	yes	yes	no	no								
R540				yes	yes	yes	yes	yes	no							
R530XD				yes	yes	no										
R620				yes	yes	yes										
R630				yes	yes	yes										

Dell EMC Platform		OCP Mezz	Rack NDC	PCI Express Slot												
				1	2	3	4	5	6	7	8	9	10	11	12	13
R640			yes	yes	yes	yes										
R720XD			yes	yes	yes	yes	yes	yes	yes							
R720			yes	yes	yes	yes	yes	yes	yes	yes						
R730				yes	yes	yes	yes	yes	yes	yes						
R730XD				yes	yes	yes	yes	yes	yes							
R740			yes	yes	yes	yes	yes	yes	yes	yes	yes					
R740XD2				yes	yes	yes	yes	yes	no							
R820			yes	yes	yes	yes	yes	yes	yes	yes						
R830				yes	yes	yes	yes	yes	yes							
R840			yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes			
R920			yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes			
R930				yes	yes	yes	yes	yes	yes	yes	yes	yes	yes			
R940			yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes	yes
T130				no	no	no	no									
T140				no	no	no	no									
T320				no	no	yes	yes		yes							
T330				no	no	no	no									
T340				no	no	no	no									
T420				no	no	yes	yes	yes	yes							
T430				no	no	yes	yes	yes	yes							
T440				no	yes	yes	yes	yes								
T620				yes	yes	no	yes	yes	yes	yes						
T630				yes	no	yes	yes	yes	yes	yes						
T640			yes	yes	yes	yes	yes	yes	yes	yes	yes					

The following Dell EMC Platforms support SR-IOV on all slots.

- C6520
- C6525
- R650
- R650xa
- R6515
- R6525
- R750
- R750xa
- R7525
- R7515

Dell EMC Platform	Blade NDC	Mezzanine Slot	
		B	C
FC430	yes	yes	yes
FC630	yes	yes	yes
FC830	yes	yes	yes
M420	yes	yes	yes
M520	no	yes	yes
M620	yes	yes	yes
M630	yes	yes	yes
M630 for VRTX	yes		
M640	yes	yes	yes
M640 for VRTX	yes		
M820	yes	yes	yes
M830	yes	yes	yes
M830 for VRTX	yes		
MX740c	yes	yes	yes
MX840c	yes	yes	yes

Supported platforms or slots are indicated by "yes." Unsupported are indicated by "no". Not applicable are indicated by blank cells.

TCP Checksum Offload (IPv4 and IPv6)

Allows the adapter to verify the TCP checksum of incoming packets and compute the TCP checksum of outgoing packets. This feature enhances receive and transmit performance and reduces CPU utilization.

With Offloading off, the operating system verifies the TCP checksum.

With Offloading on, the adapter completes the verification for the operating system.

Default	RX & TX Enabled
Range	• Disabled • RX Enabled • TX Enabled • RX & TX Enabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "TCP Checksum Offload (IPv4)" -
DisplayValue "Tx Enabled"
```

TCP/IP Offloading Options

Thermal Monitoring

Adapters and network controllers based on the Intel® Ethernet Controller I350 (and later controllers) can display temperature data and automatically reduce the link speed if the controller temperature gets too hot.

 **NOTE:** This feature is enabled and configured by the equipment manufacturer. It is not available on all adapters and network controllers. There are no user configurable settings.

Monitoring and Reporting

Temperature information is displayed on the **Link** tab in Intel® PROSet for Windows* Device Manager or Intel® PROSet Adapter Configuration Utility. There are three possible conditions:

- Temperature: Normal
Indicates normal operation.
- Temperature: Overheated, Link Reduced
Indicates that the device has reduced link speed to lower power consumption and heat.
- Temperature: Overheated, Adapter Stopped
Indicates that the device is too hot and has stopped passing traffic so it is not damaged.

If either of the overheated events occur, the device driver writes a message to the system event log.

Transmit Buffers

Defines the number of Transmit Buffers, which are data segments that enable the adapter to track transmit packets in the system memory. Depending on the size of the packet, each transmit packet requires one or more Transmit Buffers.

You might choose to increase the number of Transmit Buffers if you notice a possible problem with transmit performance. Although increasing the number of Transmit Buffers can enhance transmit performance, Transmit Buffers do consume system memory. If transmit performance is not an issue, use the default setting. This default setting varies with the type of adapter.

View the [Adapter Specifications](#) topic for help identifying your adapter.

Default	512, depending on the requirements of the adapter
Range	128-16384, in intervals of 64, for 10 Gigabit Server Adapters. 128-4096, in intervals of 64, for all other adapters.

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Transmit Buffers" -DisplayValue "128"
```

UDP Checksum Offload (IPv4 and IPv6)

Allows the adapter to verify the UDP checksum of incoming packets and compute the UDP checksum of outgoing packets. This feature enhances receive and transmit performance and reduces CPU utilization.

With Offloading off, the operating system verifies the UDP checksum.

With Offloading on, the adapter completes the verification for the operating system.

Default	RX & TX Enabled
Range	• Disabled • RX Enabled • TX Enabled • RX & TX Enabled

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "UDP Checksum Offload (IPv4)" -
DisplayValue "Tx Enabled"
```

UDP Segmentation Offload (IPv4 and IPv6)

Allows the adapter to segmenting UDP packets with payloads up to 64K into valid Ethernet frames. Because the adapter hardware is able to complete data segmentation much faster than operating system software, this feature may improve transmission performance. In addition, the adapter may use fewer CPU resources.

With Offloading off, the operating system segments UDP packets into valid Ethernet frames.

With Offloading on, the adapter segments UDP packets for the operating system.



NOTE: UDP Segmentation Offload requires:

- Microsoft* Windows Server* 2019, Version 1903, or later
- Linux* kernel 4.18, or later

Default	Enabled
Range	• Disabled • Enabled

To change this setting in Windows PowerShell, use the Set-IntelNetAdapterSetting cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "UDP Segmentation Offload (IPv4)" -
-DisplayValue "Enabled"
```

Wait for Link

Determines whether the driver waits for auto-negotiation to be successful before reporting the link state. If this feature is off, the driver does not wait for auto-negotiation. If the feature is on, the driver does wait for auto-negotiation.

If this feature is on and the speed is not set to auto-negotiation, the driver will wait for a short time for link to be established before reporting the link state.

If the feature is set to **Auto Detect**, this feature is automatically set to **On** or **Off** depending on speed and adapter type when the driver is installed. The setting is:

- Off for copper Intel gigabit adapters with a speed of "Auto".
- On for copper Intel gigabit adapters with a forced speed and duplex.
- On for fiber Intel gigabit adapters with a speed of "Auto".

Default	Auto Detect
----------------	-------------

Range	• On • Off • Auto Detect
--------------	--

This setting is found on the Advanced tab of either the device's Device Manager property sheet or the Intel® PROSet Adapter Configuration Utility.

To change this setting in Windows PowerShell, use the `Set-IntelNetAdapterSetting` cmdlet. For example:

```
Set-IntelNetAdapterSetting -Name "<adapter_name>" -DisplayName "Wait for Link" -DisplayValue  
"Off"
```

Microsoft* Windows* Driver and Software Installation and Configuration

Installing Windows Drivers and Software

 **NOTE:** To successfully install or uninstall the drivers or software, you must have administrative privileges on the computer completing installation.

Install the Drivers

-  **NOTES:**
- This will update the drivers for all supported Intel® network adapters in your system.
 - The Roll Back Driver feature of Windows Server (available on the Adapter Properties dialog's **Driver** tab) will not work correctly if Intel PROSet is present on the system. Before you use the Roll Back Driver feature, remove any teams. Then remove Intel PROSet using **Programs and Features** from the Control Panel of Windows. See Installing Intel PROSet for details regarding Intel PROSet.
 - Using Microsoft Windows Update to upgrade or downgrade your Ethernet network drivers is not supported. Please download the latest driver package from the [support website](#).

Before installing or updating the drivers, insert your adapter(s) in the computer and plug in the network cable. When Windows discovers the new adapter, it attempts to find an acceptable Windows driver already installed with the operating system.

If found, the driver is installed without any user intervention. If Windows cannot find the driver, the Found New Hardware Wizard window is displayed.

Regardless of whether Windows finds the driver, it is recommended that you follow the procedure below to install the driver. Drivers for all Intel adapters supported by this software release are installed.

1. Download the latest drivers from the [support website](#) and transfer them to the system.
2. If the Found New Hardware Wizard screen is displayed, click **Cancel**.
3. Double-click the downloaded file.
4. Select **Install** from the Dell Update Package screen.
5. Follow the instructions in the install wizard. Be sure to select Intel PROSet for installation.

 **NOTE:** Be sure to select the "iSCSI using Data Center Bridging" install option for systems that have an NPAR capable device installed.

Dell EMC Update Package (DUP)

The Dell EMC Update Package (DUP) is an executable package that will update the network drivers on your system.

-  **NOTES:**
- If you are installing a driver in a computer with existing Intel adapters, be sure to update all the adapters and ports with the same driver and Intel® PROSet software. This ensures that all adapters will function correctly.

Syntax

Network_Driver_XXXXX_WN64_XX.X.X_A00.exe [/<option1>[=<value1>]] [/<option2>[=<value2>]]...

Command Line Option Descriptions

None	If you do not specify any command line options, the package will guide you through the installation.
------	--

/? or /h	Display the Update Package usage information.
/s	Suppress all graphical user interfaces of the Update Package.
/i	Do a fresh install of the drivers contained in the Update Package.  NOTE: Requires /s option
/e=<path>	Extract the entire Update Package to the folder defined in <path>.  NOTE: Requires /s option
/drivers=<path>	Extract only driver components of the Update Package to the folder defined in <path>.  NOTE: Requires /s option
/driveronly	Install or Update only the driver components of the Update Package.  NOTE: Requires /s option
/passthrough	(Advanced) Sends all text following the /passthrough option directly to the vendor install software of the Update Package. This mode suppresses any provided graphical user interfaces, but not necessarily those of the vendor software.
/capabilities	(Advanced) Returns a coded description of this Update Package's supported features.  NOTE: Requires /s option
/l=<path>	Define a specific path for the Update Package log file.  NOTE: This option can NOT be used in combination with /passthrough or /capabilities
/f	Override a soft dependency error returned from the Update Package.  NOTE: Requires /s option, can NOT be used in combination with /passthrough or /capabilities

Examples

Update the system silently

Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /s

Fresh install silently

Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /s /i

Extract the update contents to the folder C:\mydir

Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /s /e=C:\mydir

Extract the driver components to the folder C:\mydir

Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /s /drivers=C:\mydir

Only install driver components

```
Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /s /driveronly
```

Change from the default log location to C:\my path with spaces\log.txt

```
Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /l="C:\my path with spaces\log.txt"
```

Force update to continue, even on "soft" qualification errors

```
Network_Driver_XXXXX_WN64_XX.X.X_A00.exe /s /f
```

Downgrading Drivers

You can use the /s and /f options to downgrade your drivers. For example, if you have the 17.0.0 drivers loaded and you want to downgrade to 16.5.0, type the following:

```
Network_Driver_XXXXX_WN64_16.5.0_A00.exe /s /f
```

Saving and Restoring an Adapter's Configuration Settings

The Save and Restore Command Line Tool allows you to copy the current adapter settings into a standalone file (such as on a USB drive) as a backup measure. In the event of a hard drive failure, you can reinstate most of your former settings.

The system on which you restore network configuration settings must have the same configuration as the one on which the save was performed. A saved configuration file can be used to restore adapter settings after an operating system upgrade. However, all adapter configuration settings may not be restored depending on the features supported by the new operating system or adapter configuration software.



NOTES:

- Only adapter settings are saved. The adapter's driver is not saved.
- Restore using the script only once. Restoring multiple times may result in unstable configuration.
- Intel® PROSet must be installed for the SaveRestore.ps1 script to run.
- For systems running a 64-bit OS, be sure to run the 64-bit version of Windows PowerShell, not the 32-bit (x86) version, when running the SaveRestore.ps1 script.

Command Line Syntax

```
SaveRestore.ps1 -Action save|restore [-ConfigPath] [-BDF]
```

SaveRestore.ps1 has the following command line options:

Option	Description
-Action	Required. Valid values: save restore. The save option saves adapter settings that have been changed from the default settings. When you restore with the resulting file, any settings not contained in the file are assumed to be the default. The restore option restores the settings.
-ConfigPath	Optional. Specifies the path and filename of the main configuration save file. If not specified, it is the script path and default filename (saved_config.txt).
-BDF	Optional. Default configuration file names are saved_config.txt and Saved_StaticIP.txt. If you specify -BDF during a restore, the script attempts to restore the configuration based on the PCI Bus:Device:Function:Segment values of the saved configuration. If you removed, added, or moved a NIC to a different slot, this may result in the script applying the saved settings to a different device.

**NOTES:**

- If the restore system is not identical to the saved system, the script may not restore any settings when the -BDF option is specified.
- Virtual Function devices do not support the -BDF option.
- If you used Windows to set NPar minimum and maximum bandwidth percentages, you must specify /bdf during save and restore to keep those settings.

Examples

Save Example

To save the adapter settings to a file on a removable media device, do the following.

1. Open a Windows PowerShell Prompt.
2. Navigate to the directory where SaveRestore.ps1 is located (generally c:\Program Files\Intel\Wired Networking\PROSET).
3. Type the following:

```
SaveRestore.ps1 -Action Save -ConfigPath e:\settings.txt
```

Restore Example

To restore the adapter settings from a file on removable media, do the following:

1. Open a Windows PowerShell Prompt.
2. Navigate to the directory where SaveRestore.ps1 is located (generally c:\Program Files\Intel\Wired Networking\PROSET).
3. Type the following:

```
SaveRestore.ps1 -Action Restore -ConfigPath e:\settings.txt
```

Configuring Device Features

Configuring with Intel® PROSet for Windows* Device Manager

Use the additional Windows Device Manager tabs described below to configure Intel Ethernet adapter features.



NOTE: The Intel® PROSet for Windows Device Manager is not available on Microsoft* Windows Server* 2019.

Link Speed tab

The **Link Speed** tab allows you to change the adapter's speed and duplex setting, run diagnostics, and use the identify adapter feature.

Advanced Tab

The settings listed on Intel PROSet for Windows Device Manager's **Advanced** tab allow you to customize how the adapter handles QoS packet tagging, Jumbo Packets, Offloading, and other capabilities. Some of the following features might not be available depending on the operating system you are running, the specific adapters installed, and the specific platform you are using.

Power Management Tab

The Intel® PROSet **Power Management** tab replaces the standard Microsoft* Windows* Power Management tab in Device Manager. The standard Windows power management functionality is included on the Intel PROSet tab.

 NOTES:

- The options available on the Power Management tab are adapter and system dependent. Not all adapters will display all options. There may be BIOS or operating system settings that need to be enabled for your system to wake up. In particular, this is true for Wake from S5 (also referred to as Wake from power off).
- The Intel® 10 Gigabit Network Adapters do not support power management.
- If your system has a Manageability Engine, the Link LED may stay lit even if WoL is disabled.
- When an adapter is running in NPar mode, Power Management is limited to the root partition of each port.

Configuring with Intel® PROSet for Windows PowerShell* software

The Intel® PROSet for Windows PowerShell* software contains several cmdlets that allow you to configure and manage the Intel® Ethernet Adapters and devices present in your system. For a complete list of these cmdlets and their descriptions, type **get-help IntelNetCmdlets** at the Windows PowerShell prompt. For detailed usage information for each cmdlet, type **get-help <cmdlet_name>** at the Windows PowerShell prompt.

 NOTES:

- IntelNetCmdlets are digitally signed. Microsoft* Windows* operating systems check digital signatures online. Depending on your internet connection, this may result in a delay before any cmdlet operation (including get-help). If you have not already done so, make sure you use Import-Module to import the IntelNetCmdlets.
- Online help (get-help -online) is not supported.
- To use the Minihelp property for any cmdlet in the module, append "| Select Minihelp" to the full cmdlet syntax.
For example:

```
Get-IntelNetAdapterSetting -Name "<adapter_name>" -RegistryKeyword *RSS | Select Minihelp
```

The Intel® PROSet for Windows PowerShell* software is installed by default when you install Intel PROSet. After installation, use the Import-Module cmdlet to import the new cmdlets. You may need to restart Windows PowerShell to access the newly installed cmdlets.

To use the Import-Module cmdlet, you must specify the path. For example:

```
PS c:\> Import-Module -Name "C:\Program Files\Intel\Wired Networking\IntelNetCmdlets"
```



NOTE: If you include a trailing backslash ("\") at the end of the Import-Module command, the import operation will fail. In Microsoft Windows Server* 2016, the auto-complete function appends a trailing backslash. If you use auto-complete when entering the Import-Module command, delete the trailing backslash from the path before pressing Return to execute the command.

See Microsoft TechNet for more information about the Import-Module cmdlet.

System requirements

- Microsoft* Windows PowerShell* version 2.0
- .NET version 2.0

Configuring SR-IOV for improved network security

In a virtualized environment, on Intel® Server Adapters that support SR-IOV, the virtual function (VF) may be subject to malicious behavior. Software-generated layer two frames, like IEEE 802.3x (link flow control), IEEE 802.1Qbb (priority based flow-control), and others of this type, are not expected and can throttle traffic between the host and the virtual switch, reducing performance. To resolve this issue, and to ensure isolation from unintended traffic streams, configure all SR-IOV enabled ports for VLAN tagging from the administrative interface on the PF. This configuration allows unexpected, and potentially malicious, frames to be dropped.

Changing Intel PROSet Settings via Microsoft* Windows PowerShell*

You can use the Intel® PROSet for Windows PowerShell* software to change most Intel PROSet settings.

 NOTES:

- The Get-IntelNetAdapterStatus -Status General cmdlet may report the status "Link Up - This device is not linked at its maximum capable speed". In that case, if your device is set to auto-negotiate, you can adjust the speed of the device's link partner to the device's maximum speed. If the device is not set to auto-negotiate, you can adjust the device's speed manually, but you must ensure the link partner is set at the same speed.

Changing Intel PROSet Settings Under Windows Server Core

You can use the Intel® PROSet for Windows PowerShell* software to change most Intel PROSet settings under Windows Server Core. Please refer to the aboutIntelNetCmdlets.hlp.txt help file.

For iSCSI Crash Dump configuration, use the Intel® PROSet for Windows PowerShell* software and refer to the aboutIntelNetCmdlets.hlp.txt help file.

 **NOTE:**Support for the Intel PROSet command line utilities (prosetcl.exe and crashdmp.exe) has been removed, and is no longer installed. This functionality has been replaced by the Intel® PROSet for Windows PowerShell* software. Please transition all of your scripts and processes to use the Intel® PROSet for Windows PowerShell* software.

Configuring with Intel® PROSet Adapter Configuration Utility

The Intel® PROSet Adapter Configuration Utility (Intel® PROSet ACU) is a graphical user interface that allows you to configure and manage supported Intel® Ethernet Adapters.

 **NOTE:**The Intel® PROSet ACU is available for Microsoft* Windows Server* 2019.

The general configuration steps are as follows:

1. Select an adapter in the Adapter Selection panel.
2. Select a setting to configure from the Adapter Settings panel.
3. Select or enter the desired value(s) for the selected setting.
4. Click the "Apply Changes" button.

Linux* Driver Installation and Configuration

Overview

This release includes Linux Base Drivers for Intel® Network Connections. Specific information on building and installation, configuration, and command line parameters for these drivers are located in the following sections:

- [igb Linux Driver for Intel® Ethernet Adapters](#) based on the I350, and I354 controllers
- [ixgbe Linux Driver for Intel® Ethernet Adapters](#) based on the 82599, X540, and X550 controllers
- [i40e Linux Driver for Intel® Ethernet Adapters](#) based on the X710 and XL710 controllers
- [ice Linux Driver for Intel® Ethernet Adapters](#) based on the 800 Series of controllers.

See the [Supported Adapters](#) section below to determine which driver to use.

These drivers are only supported as a loadable module. Intel is not supplying patches against the kernel source to allow for static linking of the driver.

This release also includes support for Single Root I/O Virtualization (SR-IOV) drivers. More detail on SR-IOV can be found [here](#). The following drivers support the listed virtual function devices that can only be activated on kernels that support SR-IOV. SR-IOV requires the correct platform and OS support.

- [ixgbev Linux Driver](#) for 82599, X540, and X550 10 Gigabit Family of Adapters.
- [iavf Linux Driver](#) for adapters based on the following controllers:
 - Intel Ethernet 700 Series
 - Intel Ethernet 800 Series



NOTES:

- On systems running Linux or ESXi, the base driver must be loaded for the Dell EMC FW DUP to function correctly.
- If you plan to direct-assign devices to a VM in Linux, you must enable I/O Memory Management Unit support in order for [SR-IOV](#) to function correctly. Use the kernel boot parameters "intel_iommu=on" for system boards with Intel processors or "amd_iommu=on" for systems boards with AMD processors, and "iommu=pt" to enable IOMMU support. For the best memory protection, use "intel_iommu=on." For the best performance, use both parameters ("intel_iommu=on iommu=pt"). In RedHat and most other Linux distributions, append these parameters to the GRUB_CMDLINE_LINUX entry in the /etc/default/grub configuration file. For systems booting in UEFI mode, run `grub2-mkconfig -o /etc/grub2-efi.cfg`. For systems booting in legacy BIOS mode, run `grub2-mkconfig -o /boot/grub2/grub.cfg`. In SUSE based Linux distributions, add these parameters by opening Yast and then opening the Boot Loader and clicking the Kernel Parameters tab. Add the optional parameters in the Optional Kernel Command Line Parameter field. This adds the options for either boot mode. You will need to reboot for these changes to take effect.

Supported Adapters

The following Intel network adapters are compatible with the drivers in this release:

Devices supported by the igb Linux Base Driver

- Intel® Ethernet 1G 4P I350-t OCP
- Intel® Gigabit 4P X550/I350 rNDC
- Intel® Gigabit 4P I350-t rNDC
- Intel® Gigabit 4P X540/I350 rNDC
- Intel® Gigabit 4P X520/I350 rNDC
- Intel® Gigabit 4P I350-t Mezz
- Intel® Gigabit 4P X710/I350 rNDC
- Intel® Gigabit 4P I350 bNDC
- Intel® Gigabit 2P I350-t Adapter
- Intel® Gigabit 4P I350-t Adapter

- Intel® Ethernet Connection I354 1.0 GbE Backplane
- Intel® Gigabit 2P I350-t LOM
- Intel® Gigabit I350-t LOM
- Intel® Gigabit 2P I350 LOM

Devices Supported by the ixgbe Linux Base Driver

- Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz
- Intel® Ethernet 10G 2P X540-t Adapter
- Intel® Ethernet 10G 2P X550-t Adapter
- Intel® Ethernet 10G 4P X550 rNDC
- Intel® Ethernet 10G 4P X550/I350 rNDC
- Intel® Ethernet 10G 4P X540/I350 rNDC
- Intel® Ethernet 10G 4P X520/I350 rNDC
- Intel® Ethernet 10G 2P X520-k bNDC
- Intel® Ethernet 10G 2P X520 Adapter
- Intel® Ethernet 10G X520 LOM

Devices Supported by the i40e Linux Base Driver

- Intel® Ethernet 10G 4P X710-k bNDC
- Intel® Ethernet 10G 2P X710-k bNDC
- Intel® Ethernet 10G X710-k bNDC
- Intel® Ethernet Converged Network Adapter X710
- Intel® Ethernet Converged Network Adapter X710-T
- Intel® Ethernet 10G 4P X710/I350 rNDC
- Intel® Ethernet 10G 4P X710 SFP+ rNDC
- Intel® Ethernet 10G X710 rNDC
- Intel® Ethernet Server Adapter X710-DA2 for OCP
- Intel® Ethernet 10G 2P X710 OCP
- Intel® Ethernet 10G 4P X710 OCP
- Intel® Ethernet 10G 2P X710-T2L-t OCP
- Intel® Ethernet 10G 4P X710-T4L-t OCP
- Intel® Ethernet 10G 2P X710-T2L-t Adapter
- Intel® Ethernet 10G 4P X710-T4L-t Adapter
- Intel® Ethernet 25G 2P XXV710 Adapter
- Intel® Ethernet 25G 2P XXV710 Mezz
- Intel® Ethernet Converged Network Adapter XL710-Q2
- Intel® Ethernet 40G 2P XL710 QSFP+ rNDC

Devices Supported by the ice Linux Base Driver

- Intel® Ethernet 25G 2P E810-XXV OCP
- Intel® Ethernet 25G 2P E810-XXV Adapter

To verify your adapter is supported, find the board ID number on the adapter. Look for a label that has a barcode and a number in the format 123456-001 (six digits hyphen three digits). Match this to the list of numbers above.

For more information on how to identify your adapter or for the latest network drivers for Linux, see [Customer Support](#).

Supported Linux Versions

Linux drivers are provided for the following distributions (only Intel® 64 versions are supported):

Red Hat Enterprise Linux (RHEL):

- Red Hat* Enterprise Linux* (RHEL) 8.3
- Red Hat* Enterprise Linux* (RHEL) 8.2
- Red Hat* Enterprise Linux* (RHEL) 7.9

SUSE Linux Enterprise Server (SLES):

- Novell* SUSE* Linux Enterprise Server (SLES) 15 SP2

NIC Partitioning

On Intel® 710 Series based adapters that support it, you can set up multiple functions on each physical port. You configure these functions through the System Setup/BIOS.

Minimum TX Bandwidth is the guaranteed minimum data transmission bandwidth, as a percentage of the full physical port link speed, that the partition will receive. The bandwidth the partition is awarded will never fall below the level you specify here.

The range for the minimum bandwidth values is:

1 to ((100 minus # of partitions on the physical port) plus 1)

For example, if a physical port has 4 partitions, the range would be

1 to ((100 - 4) + 1 = 97)

The Maximum Bandwidth percentage represents the maximum transmit bandwidth allocated to the partition as a percentage of the full physical port link speed. The accepted range of values is 1-100. The value can be used as a limiter, should you chose that any one particular function not be able to consume 100% of a port's bandwidth (should it be available). The sum of all the values for Maximum Bandwidth is not restricted, because no more than 100% of a port's bandwidth can ever be used.



NOTE:

- If the sum of the minimum bandwidth percentages does not equal 100, then settings will be automatically adjusted so that the sum equals 100.
- If a partition's maximum bandwidth percentage is set lower than the partition's minimum bandwidth percentage, then the maximum bandwidth percentage will be automatically set to the value of the minimum bandwidth percentage.
- When you attempt to set values for minimum bandwidth percentage via iDRAC with Lifecycle Controller using jobs that do not include the values for all enabled partitions, then the values seen after the jobs have completed may be different than the values that were supposed to be set. To avoid this issue, set the values for minimum bandwidth percentage on all partitions using a single job and make sure the sum of the values is 100.

Once the initial configuration is complete, you can set different bandwidth allocations on each function as follows:

1. Make a new directory named /config
2. Edit etc/fstab to include:


```
configfs /config configfs defaults
```
3. Load (or reload) the i40e driver
4. Mount /config
5. Make a new directory under config for each partition upon which you wish to configure the bandwidth.

Three files will appear under the config/partition directory:

- max_bw
- min_bw
- commit

Read from max_bw to get display the current maximum bandwidth setting.

Write to max_bw to set the maximum bandwidth for this function.

Read from `min_bw` to display the current minimum bandwidth setting.

Write to `min_bw` to set the minimum bandwidth for this function.

Write a '1' to commit to save your changes.



NOTES:

- commit is write only. Attempting to read it will result in an error.
- Writing to commit is only supported on the first function of a given port. Writing to a subsequent function will result in an error.
- Oversubscribing the minimum bandwidth is not supported. The underlying device's NVM will set the minimum bandwidth to supported values in an indeterminate manner. Remove all of the directories under config and reload them to see what the actual values are.
- To unload the driver you must first remove the directories created in step 5, above.

Example of Setting the minimum and maximum bandwidth (assume there are four functions on the port eth6-eth9, and that eth6 is the first function on the port):

```
# mkdir /config/eth6
# mkdir /config/eth7
# mkdir /config/eth8
# mkdir /config/eth9
# echo 50 > /config/eth6/min_bw
# echo 100 > /config/eth6/max_bw
# echo 20 > /config/eth7/min_bw
# echo 100 > /config/eth7/max_bw
# echo 20 > /config/eth8/min_bw
# echo 100 > /config/eth8/max_bw
# echo 10 > /config/eth9/min_bw
# echo 25 > /config/eth9/max_bw
# echo 1 > /config/eth6/commit
```

igb Linux* Driver for the Intel® Gigabit Adapters

igb Overview



NOTE: In a virtualized environment, on Intel® Server Adapters that support SR-IOV, the virtual function (VF) may be subject to malicious behavior. Software-generated layer two frames, like IEEE 802.3x (link flow control), IEEE 802.1Qbb (priority based flow-control), and others of this type, are not expected and can throttle traffic between the host and the virtual switch, reducing performance. To resolve this issue, and to ensure isolation from unintended traffic streams, configure all SR-IOV enabled ports for VLAN tagging from the administrative interface on the PF. This configuration allows unexpected, and potentially malicious, frames to be dropped.



NOTE: To configure VLAN tagging for the ports on an SR-IOV enabled adapter, use the following command. The VLAN configuration should be done before the VF driver is loaded or the VM is booted.

```
# ip link set dev <PF netdev id> vf <id> vlan <vlan id>
```

For example, the following instructions will configure PF eth0 and the first VF on VLAN 10. \$ ip link set dev eth0 vf 0 vlan 10.

The igb driver supports 2.5 Gbps operating speed on 2500BASE-KX only for I354-based network connections.

This driver is only supported as a loadable module

The igb driver supports IEEE time stamping for kernels 2.6.30 and above.

Use ethtool, lspci, or ifconfig to obtain driver information. Instructions on updating the ethtool can be found in the [Additional Configurations](#) section later in this page.

igb Linux Base Driver Supported Devices

The following Intel network adapters are compatible with the igb driver in this release:

- Intel® Ethernet 1G 4P I350-t OCP
- Intel® Gigabit 4P X550/I350 rNDC
- Intel® Gigabit 4P I350-t rNDC
- Intel® Gigabit 4P X540/I350 rNDC
- Intel® Gigabit 4P X520/I350 rNDC
- Intel® Gigabit 4P I350-t Mezz
- Intel® Gigabit 4P X710/I350 rNDC
- Intel® Gigabit 4P I350 bNDC
- Intel® Gigabit 2P I350-t Adapter
- Intel® Gigabit 4P I350-t Adapter
- Intel® Ethernet Connection I354 1.0 GbE Backplane
- Intel® Gigabit 2P I350-t LOM
- Intel® Gigabit I350-t LOM
- Intel® Gigabit 2P I350 LOM

Building and Installation

There are three methods for installing the igb driver:

- [Install from Source Code](#)
- [Install Using KMP RPM](#)
- [Install Using KMOD RPM](#)

Install from Source Code

To build a binary RPM* package of this driver, run 'rpmbuild -tb <filename.tar.gz>'. Replace <filename.tar.gz> with the specific filename of the driver.

 NOTES:

- For the build to work properly it is important that the currently running kernel MATCH the version and configuration of the installed kernel source. If you have just recompiled your kernel, reboot the system.
- RPM functionality has only been tested in Red Hat distributions.

1. Download the base driver tar file to the directory of your choice. For example, use '/home/username/igb' or '/usr/local/src/igb'.
2. Untar/unzip the archive, where <x.x.x> is the version number for the driver tar:

```
# tar xzf igb-<x.x.x>.tar.gz
```

3. Change to the driver src directory, where <x.x.x> is the version number for the driver tar:

```
# cd igb-<x.x.x>/src/
```

4. Compile the driver module:

```
# make install
```

The binary will be installed as:

```
/lib/modules/<KERNEL_VERSION>/kernel/drivers/net/igb/igb.ko
```

The install locations listed above are the default locations. This might differ for various Linux distributions. For more information, see the ldistrib.txt file included in the driver tar.

5. Remove the old driver:

```
# rmmod igb
```

6. Install the module using the modprobe command:

```
# modprobe igb <parameter>=<value>
```

7. Update the system image:

```
dracut -f
```

8. Assign an IP address to and activate the Ethernet interface by entering the following, where <ethx> is the interface name:

```
# ifconfig <ethX> <IP_address> netmask <netmask> up
```

9. Verify that the interface works. Enter the following, where <IP_address> is the IP address for another machine on the same subnet as the interface that is being tested:

```
# ping <IP_address>
```



NOTE: Some systems have trouble supporting MSI and/or MSI-X interrupts. If your system needs to disable this type of interrupt, the driver can be built and installed with the command:

```
# make CFLAGS_EXTRA=-DDISABLE_PCI_MSI install
```

Normally, the driver generates an interrupt every two seconds. If interrupts are not received in cat /proc/interrupts for the ethX device, then this workaround may be necessary.

Install Using KMP RPM

The KMP RPMs update existing igb RPMs currently installed on the system. These updates are provided by SuSE in the SLES release. If an RPM does not currently exist on the system, the KMP will not install.

The RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
intel-<component name>-<component version>.<arch type>.rpm
```

For example, intel-igb-1.3.8.6-1.x86_64.rpm:

- igb is the component name
- 1.3.8.6-1 is the component version
- x86_64 is the architecture type

KMP RPMs are provided for supported Linux distributions. The naming convention for the included KMP RPMs is:

```
intel-<component name>-kmp-<kernel type>-<component version>_<kernel version>.<arch
type>.rpm
```

For example, intel-igb-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm:

- igb is the component name
- default is the kernel type
- 1.3.8.6 is the component version
- 2.6.27.19_5-1 is the kernel version
- x86_64 is the architecture type

To install the KMP RPM, type the following two commands:

```
# rpm -i <rpm filename>
# rpm -i <kmp rpm filename>
```

For example, to install the igb KMP RPM package, type the following:

```
# rpm -i intel-igb-1.3.8.6-1.x86_64.rpm
# rpm -i intel-igb-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm
```

Install Using KMOD RPM

The KMOD RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
kmod-<driver name>-<version>-1.<arch type>.rpm
```

For example, kmod-igb-2.3.4-1.x86_64.rpm:

- igb is the driver name
- 2.3.4 is the version
- x86_64 is the architecture type

To install the KMOD RPM, go to the directory of the RPM and type the following command:

```
# rpm -i <rpm filename>
```

For example, to install the igb KMOD RPM package, type the following:

```
# rpm -i kmod-igb-2.3.4-1.x86_64.rpm
```

Command Line Parameters

If the driver is built as a module, the following optional parameters are used by entering them on the command line with the modprobe command using this syntax:

```
# modprobe igb [<option>=<VAL1>,<VAL2>,...]
```

A value (<VAL#>) must be assigned to each network port in the system supported by this driver. The values are applied to each instance, in function order. For example:

```
# modprobe igb InterruptThrottleRate=16000,16000
```

In this case, there are two network ports supported by igb in the system. The default value for each parameter is generally the recommended setting, unless otherwise noted.

The following table contains parameters and possible values for modprobe commands:

Parameter Name	Valid Range/Settings	Default	Description
InterruptThrottleRate	0, 1, 3, 100-100000	3	0=off 1=dynamic 3=dynamic conservative <min_ITR>-<max_ITR> Interrupt Throttle Rate controls the number of interrupts each interrupt vector can generate per second. Increasing ITR lowers latency at the cost of increased CPU utilization, though it may help throughput in some circumstances. 0 = Setting InterruptThrottleRate to 0 turns off any interrupt moderation and may improve small packet latency. However, this is generally not suitable for bulk throughput traffic due to the increased CPU utilization of the higher interrupt rate. 1 = Setting InterruptThrottleRate to Dynamic mode attempts to moderate interrupts per vector while maintaining very low latency. This can sometimes cause extra CPU utilization. If planning on deploying this driver in a latency sensitive environment, this parameter should be considered. <min_ITR>-<max_ITR> = Setting InterruptThrottleRate to a value greater or equal to <min_ITR> will program the adapter to send at most that many interrupts per second, even if more packets have come in. This reduces interrupt load on the system and can lower CPU utilization under heavy load, but will increase latency as packets are not processed as quickly.  NOTE: Unsupported Adapters: InterruptThrottleRate is NOT supported by 82542, 82543, or 82544-based adapters.
LLI			Low Latency Interrupts (LLI) allow for immediate generation of an interrupt upon processing receive packets that match certain criteria as set by the parameters described below. LLI parameters are not enabled when Legacy interrupts are used. You must be using MSI or MSI-X (see cat /proc/interrupts) to successfully use LLI.
LLIPort	0-65535	0 (disabled)	LLI is configured with the LLIPort command line parameter, which specifies which TCP port should generate Low Latency Interrupts. For example, using LLIPort=80 would cause the board to generate an immediate interrupt upon receipt of any packet sent to TCP port 80 on the local machine.  WARNING: Enabling LLI can result in an excessive number of interrupts/second that may cause problems with the system and in some cases may cause a kernel panic.
LLIPush	0-1	0 (disabled)	LLIPush can be set to enabled or disabled (default). It is most effective in an environment with many small transactions.  NOTE: Enabling LLIPush may allow a denial of service attack.

Parameter Name	Valid Range/Settings	Default	Description																									
LLISize	0-1500	0 (dis-abled)	LLISize causes an immediate interrupt if the board receives a packet smaller than the specified size.																									
IntMode	0-2	2	Interrupt mode controls the allowed load time control over the type of interrupt registered for by the driver. MSI-X is required for multiple queue support. Some kernels and combinations of kernel .config options will force a lower level of interrupt support. 'cat/-proc/interrupts' will show different values for each type of interrupt. 0 = Legacy Interrupts 1 = MSI Interrupts 2 = MSI-X interrupts																									
RSS	0-8	1	0 = Assign up to the lesser value of the number of CPUs or the number of queues X = Assign X queues, where X is less than or equal to the maximum number of queues The maximum number of queues allowed are: • I350-based adapters: 8 queues • 82575-based adapters: 4 queues • 82576-based and newer adapters: 8 queues • I210-based adapters: 4 queues • I211-based adapters: 2 queues This parameter is also affected by the VMDq parameter in that it will limit the queues more. <table><tr><td></td><td colspan="4">VMDQ</td></tr><tr><td>Model</td><td>0</td><td>1</td><td>2</td><td>3+</td></tr><tr><td>82575</td><td>4</td><td>4</td><td>3</td><td>1</td></tr><tr><td>82576</td><td>8</td><td>2</td><td>2</td><td>2</td></tr><tr><td>82580</td><td>8</td><td>1</td><td>1</td><td>1</td></tr></table>		VMDQ				Model	0	1	2	3+	82575	4	4	3	1	82576	8	2	2	2	82580	8	1	1	1
	VMDQ																											
Model	0	1	2	3+																								
82575	4	4	3	1																								
82576	8	2	2	2																								
82580	8	1	1	1																								
VMDQ	0-4 on 82575-based adapters 0-8 for 82576/82580-based adapters	0	Supports enabling VMDq pools, which is needed to support SR-IOV. 0 = Disabled 1 = Sets the netdev as pool 0 2+ = Add additional queues but they currently are not used This parameter is forced to 1 or more if the max_vfs module parameter is used. In addition, the number of queues available for RSS is limited if this is set to 1 or greater.  NOTE: When either SR-IOV mode or VMDq mode is enabled, hardware VLAN filtering and VLAN tag stripping/insertion will remain enabled.																									

Parameter Name	Valid Range/Settings	Default	Description
max_vfs	0-7	0	This parameter adds support for SR-IOV. It causes the driver to spawn up to max_vfs worth of virtual functions. If the value is greater than 0 it will also force the VMDq parameter to be 1 or more. The parameters for the driver are referenced by position. Thus, if you have a dual port adapter, or more than one adapter in your system, and want N virtual functions per port, you must specify a number for each port with each parameter separated by a comma. For example: <pre># modprobe igb max_vfs=4</pre> This will spawn 4 VFs on the first port. <pre># modprobe igb max_vfs=2,4</pre> This will spawn 2 VFs on the first port and 4 VFs on the second port.  NOTES: • Caution must be used in loading the driver with these parameters. Depending on your system configuration, number of slots, etc., it is impossible to predict in all cases where the positions would be on the command line. • Neither the device nor the driver control how VFs are mapped into config space. Bus layout will vary by operating system. On operating systems that support it, you can check sysfs to find the mapping. When either SR-IOV mode or VMDq mode is enabled, hardware VLAN filtering and VLAN tag stripping/insertion will remain enabled. Please remove the old VLAN filter before the new VLAN filter is added. For example: <pre># ip link set eth0 vf 0 vlan 100 // set vlan 100 for VF 0</pre> <pre># ip link set eth0 vf 0 vlan 0 // Delete vlan 100</pre> <pre># ip link set eth0 vf 0 vlan 200 // set a new vlan 200 for VF 0</pre>
QueuePairs	0-1	1	If set to 0, when MSI-X is enabled, the Tx and Rx will attempt to occupy separate vectors. This option can be overridden to 1 if there are not sufficient interrupts available. This can occur if any combination of RSS, VMDQ, and max_vfs results in more than 4 queues being used.
Node	0-n, -1	-1 (off)	0-n: where n is the number of the NUMA node that should be used to allocate memory for this adapter port. -1: uses the driver default of allocating memory on whichever processor is running modprobe.

Parameter Name	Valid Range/Settings	Default	Description
			The Node parameter allows you to choose which NUMA node you want to have the adapter allocate memory from. All driver structures, in-memory queues, and receive buffers will be allocated on the node specified. This parameter is only useful when interrupt affinity is specified; otherwise, part of the interrupt time could run on a different core than where the memory is allocated causing slower memory access and impacting throughput, CPU, or both.
EEE	0-1	1 (enabled)	0 = Disables EEE 1 = Enables EEE A link between two EEE-compliant devices will result in periodic bursts of data followed by periods where the link is in an idle state. This Low Power Idle (LPI) state is supported at 1 Gbps and 100 Mbps link speeds.  NOTES: • EEE support requires auto-negotiation. • Both link partners must support EEE. • EEE is not supported on all Intel® Ethernet Network devices or at all link speeds.
DMAC	0, 250, 500, 1000, 2000, 3000, 4000, 5000, 6000, 7000, 8000, 9000, 10000	0 (disabled)	This parameter enables or disables DMA Coalescing feature. Values are in microseconds and set the internal DMA Coalescing internal timer. DMA (Direct Memory Access) allows the network device to move packet data directly to the system's memory, reducing CPU utilization. However, the frequency and random intervals at which packets arrive do not allow the system to enter a lower power state. DMA Coalescing allows the adapter to collect packets before it initiates a DMA event. This may increase network latency but also increases the chances that the system will enter a lower power state. Turning on DMA Coalescing may save energy with kernel 2.6.32 and later. DMA Coalescing must be enabled across all active ports in order to save platform power.
MDD	0-1	1 (enabled)	0 = Disabled 1 = Enabled This parameter is only relevant for devices operating in SR-IOV mode. When this parameter is set, the driver detects malicious VF driver and disables its Tx/Rx queues until a VF driver reset occurs.

Additional Configurations

ethtool

The driver utilizes the ethtool interface for driver configuration and diagnostics, as well as displaying statistical information. The latest ethtool version is required for this functionality. Download it at: <https://kernel.org/pub/software/network/ethtool/>

Viewing Link Messages

Link messages will not be displayed to the console if the distribution is restricting system messages. In order to see network driver link messages on your console, set dmesg to eight by entering the following:

```
# dmesg -n 8
```



NOTE: This setting is not saved across reboots.

Configuring the Driver on Different Distributions

Configuring a network driver to load properly when the system is started is distribution dependent. Typically, the configuration process involves adding an alias line to `/etc/modules.conf` or `/etc/modprobe.conf` as well as editing other system start up scripts and/or configuration files. Many popular Linux distributions ship with tools to make these changes for you. To learn the proper way to configure a network device for your system, refer to your distribution documentation. If during this process you are asked for the driver or module name, the name for the Linux Base Driver for the device is `igb`.

For example, if you install the `igb` driver for two adapters (`eth0` and `eth1`) and want to set the interrupt mode to MSI-X and MSI, respectively, add the following to `modules.conf` or `/etc/modprobe.conf`:

```
# alias eth0 igb
# alias eth1 igb
# options igb IntMode=2,1
```

Jumbo Frames

Jumbo Frames support is enabled by changing the MTU to a value larger than the default of 1500 bytes. Use the `ifconfig` command to increase the MTU size. For example, enter the following where `<ethX>` is the interface number::

```
# ifconfig <ethX> mtu 9000 up
```

Alternatively, you can use the `ip` command as follows:

```
# ip link set mtu 9000 dev <ethX>
# ip link set up dev <ethX>
```

This setting is not saved across reboots. The setting change can be made permanent by adding 'MTU = 9000' to the following file:

- `/etc/sysconfig/network-scripts/ifcfg-<ethX>` for RHEL
- `/etc/sysconfig/network/<config_file>` for SLES



NOTES:

- The maximum MTU setting for Jumbo Frames is 9216 bytes. This value coincides with the maximum Jumbo Frames size of 9234 bytes.
- Using Jumbo Frames at 10 or 100 Mbps may result in poor performance or loss of link.
- Packet loss may have a greater impact on throughput when you use jumbo frames. If you observe a drop in performance after enabling jumbo frames, enabling flow control may mitigate the issue.

Speed and Duplex Configuration

In addressing speed and duplex configuration issues, you need to distinguish between copper-based adapters and fiber-based adapters.

In the default mode, an Intel® Ethernet Network Adapter using copper connections will attempt to auto-negotiate with its link partner to determine the best setting. If the adapter cannot establish link with the link partner using auto-negotiation, you may need to manually configure the adapter and link partner to identical settings to establish link and pass packets. This should only be needed when attempting to link with an older switch that does not support auto-negotiation or one that has been forced to a specific speed or duplex mode. Your link partner must match the setting you choose. 1 Gbps speeds and higher cannot be forced. Use the autonegotiation advertising setting to manually set devices for 1 Gbps and higher.

Speed, duplex, and autonegotiation advertising are configured through the ethtool utility. ethtool is included with all versions of Red Hat after Red Hat 7.2. To see the speed configurations your device supports, run the following:

```
# ethtool <ethX>
```



CAUTION: Only experienced network administrators should force speed and duplex or change autonegotiation advertising manually. The settings at the switch must always match the adapter settings. Adapter performance may suffer or your adapter may not operate if you configure the adapter differently from your switch.

An Intel® Ethernet Network Adapter using fiber-based connections will not attempt to auto-negotiate with its link partner since those adapters operate only in full duplex and only at their native speed.

Wake on LAN (WoL) Support

Some adapters do not support Wake on LAN (WoL). To determine if your adapter supports WoL, run the following command:

```
# ethtool <ethX>
```

WoL is configured through the ethtool utility. ethtool is included with all versions of Red Hat after Red Hat 7.2. For other Linux distributions, download and install ethtool from the following website: <https://kernel.org/pub/software/network/ethtool/>.

For instructions on enabling WoL with ethtool, refer to the website listed above.

WoL will be enabled on the system during the next shutdown or reboot. For this driver version, in order to enable WoL, the driver must be loaded prior to shutting down or suspending the system.



NOTES:

- Wake on LAN is only supported on port A of multi-port devices.
- Wake On LAN is not supported for the Intel® Gigabit VT Quad Port Server Adapter.

Multiqueue

In this mode, a separate MSI-X vector is allocated for each queue and one for "other" interrupts such as link status change and errors. All interrupts are throttled via interrupt moderation. Interrupt moderation must be used to avoid interrupt storms while the driver is processing one interrupt. The moderation value should be at least as large as the expected time for the driver to process an interrupt. Multiqueue is off by default.

REQUIREMENTS: MSI-X support is required for Multiqueue. If MSI-X is not found, the system will fallback to MSI or to Legacy interrupts. This driver supports multiqueue in kernel versions 2.6.24 and newer. This driver supports receive multiqueue on all kernels that support MSI-X.



NOTES:

- Do not use MSI-X with the 2.6.19 or 2.6.20 kernels.
- On some kernels a reboot is required to switch between single queue mode and multiqueue mode or vice-versa.

Large Receive Offload (LRO)

Large Receive Offload (LRO) is a technique for increasing inbound throughput of high-bandwidth network connections by reducing CPU overhead. It works by aggregating multiple incoming packets from a single stream into a larger buffer before they are passed higher up the networking stack, thus reducing the number of packets that have to be processed. LRO combines multiple Ethernet frames into a single receive in the stack, thereby potentially decreasing CPU utilization for receives.

LRO requires 2.4.22 or later kernel version.

IGB_LRO is a compile time flag. The user can enable it at compile time to add support for LRO from the driver. The flag is used by adding `CFLAGS_EXTRA="-DIGB_LRO"` to the make file when it's being compiled.

```
# make CFLAGS_EXTRA="-DIGB_LRO" install
```

You can verify that the driver is using LRO by looking at these counters in ethtool:

- `lro_aggregated` - counts total packets that were combined
- `lro_flushed` - counts the number of packets flushed out of LRO
- `lro_recycled` - counts the number of buffers returned to the ring from recycling



NOTE: IPv6 and UDP are not supported by LRO.

IEEE 1588 Precision Time Protocol (PTP) Hardware Clock (PHC)

Precision Time Protocol (PTP) is used to synchronize clocks in a computer network. PTP support varies among Intel devices that support this driver. Use `'ethtool -T <ethX>'` to get a definitive list of PTP capabilities supported by the device.



NOTE: PTP requires a 3.0.0 or later kernel version with PTP support enabled in the kernel and a user-space software daemon.

IGB_PTP is a compile time flag. The user can enable it at compile time to add support for PTP from the driver. The flag is used by editing the make file as follows when it is being compiled:

```
# make CFLAGS_EXTRA="-DIGB_PTP" install
```



NOTE: The driver will fail to compile if your kernel does not support PTP.

You can verify that the driver is using PTP by looking at the system log to see whether a PHC was attempted to be registered or not. If you have a kernel and version of ethtool with PTP support, you can check the PTP support in the driver by executing:

```
# ethtool -T <ethX>
```

MAC and VLAN Anti-Spoofing Feature for VFs

When a malicious driver on a Virtual Function (VF) interface attempts to send a spoofed packet, it is dropped by the hardware and not transmitted.

An interrupt is sent to the PF driver notifying it of the spoof attempt. When a spoofed packet is detected, the PF driver will send the following message to the system log (displayed by the `"dmesg"` command):

When a spoofed packet is detected the PF driver will send the following message to the system log (displayed by the `"dmesg"` command):

```
Spoof event(s) detected on VF(n)
```

Where n=the VF that attempted to do the spoofing.

Setting MAC Address, VLAN and Rate Limit Using IProute2 Tool

You can set a MAC address of a Virtual Function (VF), a default VLAN and the rate limit using the IProute2 tool. Download the latest version of the iproute2 tool from Sourceforge if your version does not have all the features you require.

Known Issues

MAC Address of Virtual Function Changes Unexpectedly

If a Virtual Function's MAC address is not assigned in the host, then the VF (virtual function) driver will use a random MAC address. This random MAC address may change each time the VF driver is reloaded. You can assign a static MAC address in the host machine. This static MAC address will survive a VF driver reload.

Hardware Issues

For known hardware and troubleshooting issues, either refer to the "Release Notes" in your User Guide, or for more detailed information, go to <http://www.intel.com>.

In the search box enter your devices controller ID followed by "spec update". The specification update file has complete information on known hardware issues.

Software Issues



NOTE: After installing the driver, if your Intel® Ethernet Network Connection is not working, verify that you have installed the correct driver. Intel® Active Management Technology 2.0, 2.1, and 2.5 are not supported in conjunction with the Linux driver.

Using the igb Driver on 2.4 or Older 2.6 Based Kernels

Due to limited support for PCI-Express in 2.4 kernels and older 2.6 kernels, the igb driver may run into interrupt related problems on some systems, such as no link or hang when bringing up the device.

We recommend the newer 2.6 based kernels, as these kernels correctly configure the PCI-Express configuration space of the adapter and all intervening bridges. If you are required to use a 2.4 kernel, use a 2.4 kernel newer than 2.4.30. For 2.6 kernels we recommend using the 2.6.21 kernel or newer.

Alternatively, on 2.6 kernels you may disable MSI support in the kernel by booting with the "pci=noms" option or permanently disable MSI support in your kernel by configuring your kernel with CONFIG_PCI_MSI unset.

Detected Tx Unit Hang in Quad Port Adapters

In some cases ports 3 and 4 don't pass traffic and report 'Detected Tx Unit Hang' followed by 'NETDEV WATCHDOG: <ethX>: transmit timed out' errors. Ports 1 and 2 do not show any errors and will pass traffic.

This issue may be resolved by updating to the latest kernel and BIOS. You should use an OS that fully supports Message Signaled Interrupts (MSI) and make sure that MSI is enabled in your system's BIOS.

Compiling the Driver

When trying to compile the driver by running make install, the following error may occur: "Linux kernel source not configured - missing version.h"

To solve this issue, create the version.h file by going to the Linux source tree and entering:

```
# make include/linux/version.h
```

Performance Degradation with Jumbo Frames

Degradation in throughput performance may be observed in some Jumbo frames environments. If this is observed, increasing the application's socket buffer size and/or increasing the /proc/sys/net/ipv4/tcp_*mem entry values may help.

See the specific application manual and /usr/src/linux*/Documentation/networking/ip-sysctl.txt for more details.

Jumbo Frames on Foundry BigIron 8000 switch

There is a known issue using Jumbo frames when connected to a Foundry BigIron 8000 switch. This is a 3rd party limitation. If you experience loss of packets, lower the MTU size.

Multiple Interfaces on Same Ethernet Broadcast Network

Due to the default ARP behavior on Linux, it is not possible to have one system on two IP networks in the same Ethernet broadcast domain (non-partitioned switch) behave as expected. All Ethernet interfaces will respond to IP traffic for any IP address assigned to the system. This results in unbalanced receive traffic.

If you have multiple interfaces in a server, either turn on ARP filtering by entering:

```
# echo 1 > /proc/sys/net/ipv4/conf/all/arp_filter
```

This only works if your kernel's version is higher than 2.4.5.



NOTE: This setting is not saved across reboots. The configuration change can be made permanent by adding the following line to the file `/etc/sysctl.conf`:

```
net.ipv4.conf.all.arp_filter = 1
```

Another alternative is to install the interfaces in separate broadcast domains (either in different switches or in a switch partitioned to VLANs).

Disable rx Flow Control with ethtool

In order to disable receive flow control using ethtool, you must turn off auto-negotiation on the same command line:

```
# ethtool -A <ethX> autoneg off rx off
```

Unplugging Network Cable While ethtool -p is Running

In kernel versions 2.5.50 and newer, unplugging the network cable while ethtool -p is running will cause the system to become unresponsive to keyboard commands, except for control-alt-delete. Restarting the system appears to be the only remedy.

Do Not Use LRO When Routing Packets

Due to a known general compatibility issue with LRO and routing, do not use LRO when routing packets.

Rx Page Allocation Errors

'Page allocation failure. order:0' errors may occur under stress with kernels 2.6.25 and newer. This is caused by the way the Linux kernel reports this stressed condition.

Under Red Hat 5.4-GA, the system may crash when closing the guest OS window after loading or unloading the Physical Function (PF) driver. Do not remove the igb driver from Dom0 while Virtual Functions (VFs) are assigned to guests. VFs must first use the `xm "pci-detach"` command to hot-plug the VF device out of the VM it is assigned to or else shut down the VM.

Unloading Physical Function (PF) driver causes the system to reboot when the VM is running and VF is loaded on the VM. Do not unload the PF driver (igb) while VFs are assigned to guests.

Host May Reboot after Removing PF when VF is Active in Guest

Using kernel versions earlier than 3.2, do not unload the PF driver with active VFs. Doing this will cause your VFs to stop working until you reload the PF driver and may cause a spontaneous reboot of your system.

Prior to unloading the PF driver, you must first ensure that all VFs are no longer active. Do this by shutting down all VMs and unloading the VF driver.

ixgbe Linux* Driver for the Intel® 10 Gigabit Server Adapters

ixgbe Overview

	WARNING: By default, the ixgbe driver complies with the Large Receive Offload (LRO) feature enabled. This option offers the lowest CPU utilization for receives but is incompatible with routing/ip forwarding and bridging. If enabling ip forwarding or bridging is a requirement, it is necessary to disable LRO using compile time options as noted in the LRO section later in this section. The result of not disabling LRO when combined with ip forwarding or bridging can be low throughput or even a kernel panic.
	NOTE: Do not unload a port's driver if a Virtual Function (VF) with an active Virtual Machine (VM) is bound to it. Doing so will cause the port to appear to hang. Once the VM shuts down, or otherwise releases the VF, the command will complete.
	NOTE: In a virtualized environment, on Intel® Server Adapters that support SR-IOV, the virtual function (VF) may be subject to malicious behavior. Software-generated layer two frames, like IEEE 802.3x (link flow control), IEEE 802.1Qbb (priority based flow-control), and others of this type, are not expected and can throttle traffic between the host and the virtual switch, reducing performance. To resolve this issue, and to ensure isolation from unintended traffic streams, configure all SR-IOV enabled ports for VLAN tagging from the administrative interface on the PF. This configuration allows unexpected, and potentially malicious, frames to be dropped.

This file describes the Linux* Base Driver for the 10 Gigabit Intel® Network Connections. This driver supports the 2.6.x and newer kernels and includes support for any Linux supported system, including X86_64, i686 and PPC.

This driver is only supported as a loadable module. Intel is not supplying patches against the kernel source to allow for static linking of the driver. A version of the driver may already be included by your distribution or the kernel.

The following features are now available in supported kernels:

- Native VLANs
- Channel Bonding (teaming)
- Generic Receive Offload
- Data Center Bridging

Adapter teaming is now implemented using the native Linux Channel bonding module. This is included in supported Linux kernels. Channel Bonding documentation can be found in the Linux kernel source: /documentation/networking/bonding.txt

Use ethtool, lspci, or ifconfig to obtain driver information. Instructions on updating the ethtool can be found in the [Additional Configurations](#) section later in this page.

ixgbe Linux Base Driver Supported Devices

The following Intel network adapters are compatible with the Linux driver in this release:

- Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz
- Intel® Ethernet 10G 2P X540-t Adapter
- Intel® Ethernet 10G 2P X550-t Adapter
- Intel® Ethernet 10G 4P X550 rNDC
- Intel® Ethernet 10G 4P X550/I350 rNDC
- Intel® Ethernet 10G 4P X540/I350 rNDC
- Intel® Ethernet 10G 4P X520/I350 rNDC
- Intel® Ethernet 10G 2P X520-k bNDC
- Intel® Ethernet 10G 2P X520 Adapter
- Intel® Ethernet 10G X520 LOM

Building and Installation

There are three methods for installing the Linux driver:

- [Install from Source Code](#)
- [Install Using KMP RPM](#)
- [Install Using KMOD RPM](#)

Install from Source Code

To build a binary RPM* package of this driver, run 'rpmbuild -tb <filename.tar.gz>'. Replace <filename.tar.gz> with the specific filename of the driver.



NOTES:

- For the build to work properly it is important that the currently running kernel MATCH the version and configuration of the installed kernel source. If you have just recompiled your kernel, reboot the system.
- RPM functionality has only been tested in Red Hat distributions.

1. Download the base driver tar file to the directory of your choice. For example, use '/home/username/ixgbe' or '/usr/local/src/ixgbe'.
2. Untar/unzip the archive, where <x.x.x> is the version number for the driver tar:

```
# tar xzf ixgbe-<x.x.x>.tar.gz
```

3. Change to the driver src directory, where <x.x.x> is the version number for the driver tar:

```
# cd ixgbe-<x.x.x>/src/
```

4. Compile the driver module:

```
# make install
```

The binary will be installed as:

```
/lib/modules/<KERNEL VERSION>/kernel/drivers/net/ixgbe/ixgbe.ko
```

The install locations listed above are the default locations. This might differ for various Linux distributions. For more information, see the ldistrib.txt file included in the driver tar.



NOTE: IXGBE_NO_LRO is a compile time flag. The user can enable it at compile time to remove support for LRO from the driver. The flag is used by adding 'CFLAGS_EXTRA="-DIXGBE_NO_LRO"' to the make file when it is being compiled. For example:

```
# make CFLAGS_EXTRA="-DIXGBE_NO_LRO" install
```

5. Remove the old driver:

```
# rmmod ixgbe
```

6. Install the module using the modprobe command:

```
# modprobe ixgbe <parameter>=<value>
```

7. Update the system image:

```
dracut -f
```

8. Assign an IP address to and activate the Ethernet interface by entering the following, where <ethx> is the interface name:

```
# ifconfig <ethX> <IP_address> netmask <netmask> up
```

9. Verify that the interface works. Enter the following, where <IP_address> is the IP address for another machine on the same subnet as the interface that is being tested:

```
# ping <IP_address>
```

Install Using KMP RPM

The KMP RPMs update existing ixgbe RPMs currently installed on the system. These updates are provided by SuSE in the SLES release. If an RPM does not currently exist on the system, the KMP will not install.

The RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
intel-<component name>-<component version>.<arch type>.rpm
```

For example, intel-ixgbe-1.3.8.6-1.x86_64.rpm:

- ixgbe is the component name
- 1.3.8.6-1 is the component version
- x86_64 is the architecture type

KMP RPMs are provided for supported Linux distributions. The naming convention for the included KMP RPMs is:

```
intel-<component name>-kmp-<kernel type>-<component version>_<kernel version>.<arch type>.rpm
```

For example, intel-ixgbe-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm:

- ixgbe is the component name
- default is the kernel type
- 1.3.8.6 is the component version
- 2.6.27.19_5-1 is the kernel version
- x86_64 is the architecture type

To install the KMP RPM, type the following two commands:

```
# rpm -i <rpm filename>
# rpm -i <kmp rpm filename>
```

For example, to install the ixgbe KMP RPM package, type the following:

```
# rpm -i intel-ixgbe-1.3.8.6-1.x86_64.rpm
# rpm -i intel-ixgbe-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm
```

Install Using KMOD RPM

The KMOD RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
kmod-<driver name>-<version>-1.<arch type>.rpm
```

For example, kmod-ixgbe-2.3.4-1.x86_64.rpm:

- ixgbe is the driver name
- 2.3.4 is the version
- x86_64 is the architecture type

To install the KMOD RPM, go to the directory of the RPM and type the following command:

```
# rpm -i <rpm filename>
```

For example, to install the ixgbe KMOD RPM package, type the following:

```
# rpm -i kmod-ixgbe-2.3.4-1.x86_64.rpm
```

Command Line Parameters

If the driver is built as a module, the following optional parameters are used by entering them on the command line with the modprobe command using this syntax:

```
# modprobe ixgbe [<option>=<VAL1>,<VAL2>,...]
```

For example:

```
# modprobe ixgbe InterruptThrottleRate=16000,16000
```

The default value for each parameter is generally the recommended setting, unless otherwise noted.

The following table contains parameters and possible values for modprobe commands:

Parameter Name	Valid Range/Settings	Default	Description
RSS	0 - 16	1	0 = Assign up to the lesser value of the number of CPUs or the number of queues X = Assign X queues, where X is less than or equal to the maximum number of queues RSS also effects the number of transmit queues allocated on 2.6.23 and newer kernels with CONFIG_NET_MULTIQUEUE set in the kernel .config file. CONFIG_NETDEVICES_MULTIQUEUE is only supported in kernels 2.6.23 to 2.6.26. For kernels 2.6.27 or newer, other options enable multiqueue.
Multiqueue	0, 1	1	0 = Disables Multiple Queue support 1 = Enables Multiple Queue support (a prerequisite for RSS)
Direct Cache Access (DCA)	0, 1		0 = Disables DCA support in the driver 1 = Enables DCA support in the driver If the driver is enabled for DCA, this parameter allows load-time control of the feature.  NOTE: DCA is not supported on X550-based adapters.
IntMode	0 - 2	2	Interrupt mode controls the allowed load time control over the type of interrupt registered for by the driver. MSI-X is required for multiple queue support. Some kernels and combinations of kernel .config options will force a lower level of interrupt support. 'cat/proc/interrupts' will show different values for each type of interrupt. 0 = Legacy Interrupts 1 = MSI Interrupts 2 = MSI-X interrupts

Parameter Name	Valid Range/Settings	Default	Description
InterruptThrottleRate	956 - 488,281 (0=off, 1=dynamic)	1	0=off 1=dynamic <min_ITR>-<max_ITR> Interrupt Throttle Rate controls the number of interrupts each interrupt vector can generate per second. Increasing ITR lowers latency at the cost of increased CPU utilization, though it may help throughput in some circumstances. • 0 = Setting InterruptThrottleRate to 0 turns off any interrupt moderation and may improve small packet latency. However, this is generally not suitable for bulk throughput traffic due to the increased CPU utilization of the higher interrupt rate. • 1 = Setting InterruptThrottleRate to Dynamic mode attempts to moderate interrupts per vector while maintaining very low latency. This can sometimes cause extra CPU utilization. If planning on deploying this driver in a latency sensitive environment, this parameter should be considered. • <min_ITR>-<max_ITR> = Setting InterruptThrottleRate to a value greater or equal to <min_ITR> will program the adapter to send at most that many interrupts per second, even if more packets have come in. This reduces interrupt load on the system and can lower CPU utilization under heavy load, but will increase latency as packets are not processed as quickly.  On 82599, and X540, and X550-based adapters, disabling InterruptThrottleRate will also result in the driver disabling HW RSC. On 82598-based adapters, disabling InterruptThrottleRate will also result in disabling LRO (Large Receive Offloads).
LLI			Low Latency Interrupts (LLI) allow for immediate generation of an interrupt upon processing receive packets that match certain criteria as set by the parameters described below. LLI parameters are not enabled when Legacy interrupts are used. You must be using MSI or MSI-X (see cat /proc/interrupts) to successfully use LLI.  NOTE: LLI is not supported on X550-based adapters.
LLIPort	0 - 65535	0 (disabled)	LLI is configured with the LLIPort command line parameter, which specifies which TCP port should generate Low Latency Interrupts. For example, using LLIPort=80 would cause the board to generate an immediate interrupt upon receipt of any packet sent to TCP port 80 on the local machine.  WARNING: Enabling LLI can result in an excessive number of interrupts/second that may cause problems with the system and in some cases may cause a kernel panic.

Parameter Name	Valid Range/Settings	Default	Description
			 NOTE: LLI is not supported on X550-based adapters.
LLIPush	0 - 1	0 (disabled)	LLIPush can be set to enabled or disabled (default). It is most effective in an environment with many small transactions.  NOTE: Enabling LLIPush may allow a denial of service attack. LLI is not supported on X550-based adapters.
LLISize	0 - 1500	0 (disabled)	LLISize causes an immediate interrupt if the board receives a packet smaller than the specified size.  NOTE: LLI is not supported on X550-based adapters.
LLIEType	0 - x8FFF	0 (disabled)	This parameter specifies the Low Latency Interrupt (LLI) Ethernet protocol type.  NOTE: LLI is not supported on X550-based adapters.
LLIVLANP	0 - 7	0 (disabled)	This parameter specifies the LLI on VLAN priority threshold.  NOTE: LLI is not supported on X550-based adapters.
FdirPballoc	1 - 3	1 (64k)	Specifies the Flow Director allocated packet buffer size. 1 = 64k 2 = 128k 3 = 256k
AtrSampleRate	0 - 255	20	This parameter is used with the Flow Director and is the software ATR transmit packet sample rate. For example, when AtrSampleRate is set to 20, every 20th packet looks to see if the packet will create a new flow. A value of 0 indicates that ATR should be disabled and no samples will be taken.

Parameter Name	Valid Range/Settings	Default	Description
max_vfs	1 - 63	0	This parameter adds support for SR-IOV. It causes the driver to spawn up to max_vfs worth of virtual functions. If the value is greater than 0 it will also force the VMDq parameter to be 1 or more. NOTE: This parameter is only used on kernel 3.7.x and below. On kernel 3.8.x and above, use sysfs to enable VFs. Also, for Red Hat distributions, this parameter is only used on version 6.6 and older. For version 6.7 and newer, use sysfs. For example, you can create 4 VFs as follows: <pre># echo 4 > /sys/class/net/<ethX>/device/sriov_numvfs</pre> To disable VFs, write 0 to the same file: <pre># echo 0 > /sys/class/net/<ethX>/device/sriov_numvfs</pre> The parameters for the driver are referenced by position. Thus, if you have a dual port adapter, or more than one adapter in your system, and want N virtual functions per port, you must specify a number for each port with each parameter separated by a comma. For example: <pre># modprobe var_err max_vfs=4</pre> This will spawn 4 VFs on the first port. <pre># modprobe var_err max_vfs=2,4</pre> This will spawn 2 VFs on the first port and 4 VFs on the second port.  NOTES: • Caution must be used in loading the driver with these parameters. Depending on your system configuration, number of slots, etc., it is impossible to predict in all cases where the positions would be on the command line. • Neither the device nor the driver control how VFs are mapped into config space. Bus layout will vary by operating system. On operating systems that support it, you can check sysfs to find the mapping. When either SR-IOV mode or VMDq mode is enabled, hardware VLAN filtering and VLAN tag stripping/insertion will remain enabled. Please remove the old VLAN filter before the new VLAN filter is added. For example: <pre># ip link set eth0 vf 0 vlan 100 // set vlan 100 for VF 0</pre> <pre># ip link set eth0 vf 0 vlan 0 // Delete vlan 100</pre> <pre># ip link set eth0 vf 0 vlan 200 // set a new vlan 200 for VF 0</pre>

Parameter Name	Valid Range/Settings	Default	Description
			With kernel 3.6, the driver supports the simultaneous usage of max_vfs and DCB features, subject to the constraints described below. Prior to kernel 3.6, the driver did not support the simultaneous operation of max_vfs greater than 0 and the DCB features (multiple traffic classes utilizing Priority Flow Control and Extended Transmission Selection). When DCB is enabled, network traffic is transmitted and received through multiple traffic classes (packet buffers in the NIC). The traffic is associated with a specific class based on priority, which has a value of 0 through 7 used in the VLAN tag. When SR-IOV is not enabled, each traffic class is associated with a set of receive/transmit descriptor queue pairs. The number of queue pairs for a given traffic class depends on the hardware configuration. When SR-IOV is enabled, the descriptor queue pairs are grouped into pools. The Physical Function (PF) and each Virtual Function (VF) is allocated a pool of receive/transmit descriptor queue pairs. When multiple traffic classes are configured (for example, DCB is enabled), each pool contains a queue pair from each traffic class. When a single traffic class is configured in the hardware, the pools contain multiple queue pairs from the single traffic class. The number of VFs that can be allocated depends on the number of traffic classes that can be enabled. The configurable number of traffic classes for each enabled VF is as follows: • 0 - 15 VFs = Up to 8 traffic classes, depending on device support • 16 - 31 VFs = Up to 4 traffic classes • 32 - 63 VFs = 1 traffic class When VFs are configured, the PF is allocated one pool as well. The PF supports the DCB features with the constraint that each traffic class will only use a single queue pair. When zero VFs are configured, the PF can support multiple queue pairs per traffic class.
LRO	0-1		0=off, 1=on Large Receive Offload (LRO) is a technique for increasing inbound throughput of high-bandwidth network connections by reducing CPU overhead. It works by aggregating multiple incoming packets from a single stream into a larger buffer before they are passed higher up the networking stack, thus reducing the number of packets that have to be processed. LRO combines multiple Ethernet frames into a single receive in the stack, thereby potentially decreasing CPU utilization for receives. This technique is also referred to as Hardware Receive Side Coalescing (HW RSC). 82599 and X540, and X550-based adapters support HW RSC. The LRO parameter controls HW RSC enablement. You can verify that the driver is using LRO by looking at these counters in ethtool: • hw_rsc_aggregated - counts total packets that were com-

Parameter Name	Valid Range/Settings	Default	Description
			binned hw_rsc_flushed - counts the number of packets flushed out of LRO  NOTE: IPv6 and UDP are not supported by LRO.
EEE	0-1		0 = Disables EEE 1 = Enables EEE A link between two EEE-compliant devices will result in periodic bursts of data followed by periods where the link is in an idle state. This Low Power Idle (LPI) state is supported at 1 Gbps and 10 Gbps link speeds.  NOTES: - EEE support requires auto-negotiation. - Both link partners must support EEE. - EEE is not supported on all Intel® Ethernet Network devices or at all link speeds.
DMAC	0, 41-10000		This parameter enables or disables DMA Coalescing feature. Values are in microseconds and set the internal DMA Coalescing internal timer.  NOTE: DMAC is available on Intel® X550 (and later) based adapters. DMA (Direct Memory Access) allows the network device to move packet data directly to the system's memory, reducing CPU utilization. However, the frequency and random intervals at which packets arrive do not allow the system to enter a lower power state. DMA Coalescing allows the adapter to collect packets before it initiates a DMA event. This may increase network latency but also increases the chances that the system will enter a lower power state. Turning on DMA Coalescing may save energy with kernel 2.6.32 and later. DMA Coalescing must be enabled across all active ports in order to save platform power. InterruptThrottleRate (ITR) should be set to dynamic. When ITR=0, DMA Coalescing is automatically disabled. A whitepaper containing information on how to best configure your platform is available on the Intel website.
MDD	0-1	1 (enabled)	0 = Disabled 1 = Enabled This parameter is only relevant for devices operating in SR-IOV mode. When this parameter is set, the driver detects malicious VF driver and disables its Tx/Rx queues until a VF driver reset occurs.

Parameter Name	Valid Range/Settings	Default	Description
AQRate			Devices that support AQRate (X550 and later) will include 2.5 Gbps and 5 Gbps in the speeds that the driver advertises during auto-negotiation, even though ethtool will not display 2.5 Gbps or 5 Gbps as "Supported link modes" or "Advertised link modes." These speeds are only available through unmodified auto-negotiation. You cannot use ethtool -s advertise to force auto-negotiation to advertise 2.5 Gbps or 5 Gbps. If a 2.5 Gbps or 5 Gbps link is created, ethtool will report the correct link speed.

Additional Configurations

ethtool

The driver utilizes the ethtool interface for driver configuration and diagnostics, as well as displaying statistical information. The latest ethtool version is required for this functionality. Download it at: <https://kernel.org/pub/software/network/ethtool/>

Configuring the Driver on Different Distributions

Configuring a network driver to load properly when the system is started is distribution dependent. Typically, the configuration process involves adding an alias line to /etc/modules.conf or /etc/modprobe.conf as well as editing other system start up scripts and/or configuration files. Many popular Linux distributions ship with tools to make these changes for you. To learn the proper way to configure a network device for your system, refer to your distribution documentation. If during this process you are asked for the driver or module name, the name for the Linux Base Driver for the device is ixgbe.

For example, if you install the ixgbe driver for two adapters (eth0 and eth1) and want to set the interrupt mode to MSI-X and MSI, respectively, add the following to modules.conf or /etc/modprobe.conf:

```
# alias eth0 ixgbe
# alias eth1 ixgbe
# options ixgbe IntMode=2,1
```

Viewing Link Messages

Link messages will not be displayed to the console if the distribution is restricting system messages. In order to see network driver link messages on your console, set dmesg to eight by entering the following:

```
# dmesg -n 8
```



NOTE: This setting is not saved across reboots.

Jumbo Frames

Jumbo Frames support is enabled by changing the MTU to a value larger than the default of 1500 bytes. Use the ifconfig command to increase the MTU size. For example, enter the following where <ethX> is the interface number::

```
# ifconfig <ethX> mtu 9000 up
```

Alternatively, you can use the ip command as follows:

```
# ip link set mtu 9000 dev <ethX>
# ip link set up dev <ethX>
```

This setting is not saved across reboots. The setting change can be made permanent by adding 'MTU = 9000' to the following file:

- /etc/sysconfig/network-scripts/ifcfg-`<ethX>` for RHEL
- /etc/sysconfig/network/`<config_file>` for SLES



NOTES:

- The maximum MTU setting for Jumbo Frames is 9710 bytes. This value coincides with the maximum Jumbo Frames size of 9728 bytes.
- This driver will attempt to use multiple page sized buffers to receive each jumbo packet. This should help to avoid buffer starvation issues when allocating receive packets.
- Packet loss may have a greater impact on throughput when you use jumbo frames. If you observe a drop in performance after enabling jumbo frames, enabling flow control may mitigate the issue.
- For 82599-based network connections, if you are enabling jumbo frames in a virtual function (VF), jumbo frames must first be enabled in the physical function (PF). The VF MTU setting cannot be larger than the PF MTU.

Speed and Duplex Configuration

In addressing speed and duplex configuration issues, you need to distinguish between copper-based adapters and fiber-based adapters.

In the default mode, an Intel® Ethernet Network Adapter using copper connections will attempt to auto-negotiate with its link partner to determine the best setting. If the adapter cannot establish link with the link partner using auto-negotiation, you may need to manually configure the adapter and link partner to identical settings to establish link and pass packets. This should only be needed when attempting to link with an older switch that does not support auto-negotiation or one that has been forced to a specific speed or duplex mode. Your link partner must match the setting you choose. 1 Gbps speeds and higher cannot be forced. Use the autonegotiation advertising setting to manually set devices for 1 Gbps and higher.

Speed, duplex, and autonegotiation advertising are configured through the ethtool utility. ethtool is included with all versions of Red Hat after Red Hat 7.2. To see the speed configurations your device supports, run the following:

```
# ethtool <ethX>
```



CAUTION: Only experienced network administrators should force speed and duplex or change autonegotiation advertising manually. The settings at the switch must always match the adapter settings. Adapter performance may suffer or your adapter may not operate if you configure the adapter differently from your switch.

An Intel® Ethernet Network Adapter using fiber-based connections will not attempt to auto-negotiate with its link partner since those adapters operate only in full duplex and only at their native speed.



NOTE: For the Intel® Ethernet Connection X552 10 GbE SFP+, you must specify the desired speed.

Flow Control

Ethernet Flow Control (IEEE 802.3x) can be configured with ethtool to enable receiving and transmitting pause frames for this driver. When transmit is enabled, pause frames are generated when the receive packet buffer crosses a predefined threshold. When receive is enabled, the transmit unit will halt for the time delay specified when a pause frame is received.



NOTES:

- You must have a flow control capable link partner.
- This driver requires flow control on both the port and link partner. If flow control is disabled on one of the sides, the port may appear to hang on heavy traffic.
- For 82598 backplane cards entering 1 gigabit mode, flow control default behavior is changed to off. Flow control in 1 gigabit mode on these devices can lead to transmit hangs.

Use ethtool to change the flow control settings.

To enable or disable Rx or Tx Flow Control:

```
# ethtool -A <ethX> rx <on|off> tx <on|off>
```



NOTE: This command only enables or disables Flow Control if auto-negotiation is disabled. If auto-negotiation is enabled, this command changes the parameters used for auto-negotiation with the link partner.

To enable or disable auto-negotiation:

```
# ethtool -s <ethX> autoneg <on|off>
```



NOTE: Flow Control auto-negotiation is part of link auto-negotiation. Depending on your device, you may not be able to change the auto-negotiation setting.

Intel® Ethernet Flow Director

Note: Intel Ethernet Flow Director parameters are only supported on kernel versions 2.6.30 or newer.

The Intel Ethernet Flow Director performs the following tasks:

- Directs receive packets according to their flows to different queues
- Enables tight control on routing a flow in the platform
- Matches flows and CPU cores for flow affinity
- Supports multiple parameters for flexible flow classification and load balancing (in SFP mode only)

An included script (set_irq_affinity) automates setting the IRQ to CPU affinity.

Intel Ethernet Flow Director masking works in the opposite manner from subnet masking. For instance, in the following command:

```
# ethtool -N eth11 flow-type ip4 src-ip 172.4.1.2 m 255.0.0.0 dst-ip 172.21.1.1 m
255.128.0.0 action 31
```

The src-ip value that is written to the filter will be 0.4.1.2, not 172.0.0.0 as might be expected. Similarly, the dst-ip value written to the filter will be 0.21.1.1, not 172.0.0.0.

To enable or disable the Intel Ethernet Flow Director:

```
# ethtool -K <ethX> ntuple <on|off>
```

When disabling ntuple filters, all the user programmed filters are flushed from the driver cache and hardware. All needed filters must be re-added when ntuple is re-enabled.

Sideband Perfect Filters

Sideband Perfect Filters are used to direct traffic that matches specified characteristics. They are enabled through ethtool's ntuple interface. To enable or disable these filters:

```
# ethtool -K <ethX> ntuple <off|on>
```

To display all of the active filters:

```
# ethtool -u <ethX>
```

To add a new filter:

```
# ethtool -U <ethX> flow-type <type> src-ip <ip> dst-ip <ip> src-port <port> dst-
port <port> action <queue>
```

Where:

- <ethX> - the Ethernet device to program
- <type> - can be ip4, tcp4, udp4, or sctp4

- <ip> - the ip address to match on
- <port> - the port number to match on
- <queue> - the queue to direct traffic towards (-1 discards the matched traffic)

To delete a filter:

```
# ethtool -U <ethX> delete <N>
```

Where <N> is the filter ID displayed when printing all the active filters, and may also have been specified using "loc <N>" when adding the filter.

Examples:

To add a filter that directs packet to queue 2:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 action 2 [loc 1]
```

To match TCP traffic sent from 192.168.0.1, port 5300, directed to 192.168.0.5, port 80, and then send it to queue 7:

```
# ethtool -U enp130s0 flow-type tcp4 src-ip 192.168.0.1 dst-ip 192.168.0.5 src-port
5300 dst-port 80 action 7
```

For each flow-type, the programmed filters must all have the same matching input set. For example, issuing the following two commands is acceptable:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 src-port 5300 action 7
```

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.5 src-port 55 action 10
```

Issuing the next two commands, however, is not acceptable, since the first specifies `src-ip` and the second specifies `dst-ip`:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 src-port 5300 action 7
```

```
# ethtool -U enp130s0 flow-type ip4 dst-ip 192.168.0.5 src-port 55 action 10
```

The second command will fail with an error. You may program multiple filters with the same fields, using different values, but, on one device, you may not program two `tcp4` filters with different matching fields.

Matching on a subportion of a field is not supported by the driver, thus partial mask fields are not supported.

Flex Byte Flow Director Filters

The driver also supports matching user-defined data within the packet payload. This flexible data is specified using the "user-def" field of the `ethtool` command in the following way:

31	28	24	20	16	15	12	8	4	0
offset into packet payload					2 bytes of flexible data				

For example,

```
... user-def 0x4FFFF ...
```

tells the filter to look 4 bytes into the payload and match that value against `0xFFFF`. The offset is based on the beginning of the payload, and not the beginning of the packet. Thus

```
flow-type tcp4 ... user-def 0x8BEAF ...
```

would match TCP/IPv4 packets which have the value `0xBEAF` 8 bytes into the TCP/IPv4 payload.

Note that ICMP headers are parsed as 4 bytes of header and 4 bytes of payload. Thus to match the first byte of the payload, you must actually add 4 bytes to the offset. Also note that `ip4` filters match both ICMP frames as well as raw (unknown) `ip4` frames, where the payload will be the L3 payload of the IP4 frame.

The maximum offset is 64. The hardware will only read up to 64 bytes of data from the payload. The offset must be even because the flexible data is 2 bytes long and must be aligned to byte 0 of the packet payload.

The user-defined flexible offset is also considered part of the input set and cannot be programmed separately for multiple filters of the same type. However, the flexible data is not part of the input set and multiple filters may use the same offset but match against different data.

Filters to Direct Traffic to a Specific VF

It is possible to create filters that direct traffic to a specific Virtual Function. For older versions of ethtool, this depends on the "action" parameter. Specify the action as a 64-bit value, where the lower 32 bits represent the queue number, while the next 8 bits represent the VF ID. Note that 0 is the PF, so the VF identifier is offset by 1. For example:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 action 0x800000002 [loc 1]
```

The action field specifies to direct traffic to Virtual Function 7 (8 minus 1) into queue 2 of that VF.

Newer versions of ethtool (version 4.11 and later) use "vf" and "queue" parameters instead of the "action" parameter. Note that using the new ethtool "vf" parameter does not require the value to be offset by 1. This command is equivalent to the above example:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 vf 7 queue 2 [loc 1]
```

Note that these filters will not break internal routing rules, and will not route traffic that otherwise would not have been sent to the specified VF.

Support for UDP RSS

This feature adds an ON/OFF switch for hashing over certain flow types. Only UDP can be turned on. The default setting is disabled.

Only support for enabling/disabling hashing on ports for UDP over IPv4 (UDP4) or IPv6 (UDP6) is supported.



NOTE: Fragmented packets may arrive out of order when RSS UDP support is configured.

Supported ethtool Commands and Options

```
-n --show-nfc
```

Retrieves the receive network flow classification configurations.

```
rx-flow-hash tcp4|udp4|ah4|esp4|sctp4|tcp6|udp6|ah6|esp6|sctp6
```

Retrieves the hash options for the specified network traffic type.

```
-N --config-nfc
```

Configures the receive network flow classification.

```
rx-flow-hash tcp4|udp4|ah4|esp4|sctp4|tcp6|udp6|ah6|esp6|sctp6 m|v|t|s|d|f|n|r...
```

Configures the hash options for the specified network traffic type.

udp4 UDP over IPv4

udp6 UDP over IPv6

f Hash on bytes 0 and 1 of the Layer 4 header of the rx packet.

n Hash on bytes 2 and 3 of the Layer 4 header of the rx packet.

Parameters FdirPballoc and AtrSampleRate impact Flow Director.

Configuring VLAN Tagging on SR-IOV Enabled Adapter Ports

To configure VLAN tagging for the ports on an SR-IOV enabled adapter, use the following command. The VLAN configuration should be done before the VF driver is loaded or the VM is booted. The VF is not aware of the VLAN tag being inserted on transmit and removed on received frames (sometimes called "port VLAN" mode).

```
# ip link set dev <PF netdev id> vf <id> vlan <vlan id>
```

For example, the following will configure PF eth0 and the first VF on VLAN 10:

```
# ip link set dev eth0 vf 0 vlan 10
```

Data Center Bridging (DCB)

Note: The kernel assumes that TC0 is available, and will disable Priority Flow Control (PFC) on the device if TC0 is not available. To fix this, ensure TC0 is enabled when setting up DCB on your switch.

DCB is a configuration Quality of Service implementation in hardware. It uses the VLAN priority tag (802.1p) to filter traffic. That means that there are 8 different priorities that traffic can be filtered into. It also enables priority flow control (802.1Qbb) which can limit or eliminate the number of dropped packets during network stress. Bandwidth can be allocated to each of these priorities, which is enforced at the hardware level (802.1Qaz).

Adapter firmware implements LLDP and DCBX protocol agents as per 802.1AB and 802.1Qaz respectively. There are potentially two DCBX modes on Linux, depending on the underlying PF device:

- Intel Ethernet Controller 500 Series adapters only support software DCBX mode. They do not support firmware DCBX.

DCB parameters can be established via a firmware LLDP/DCBX agent. Only one LLDP/DCBX agent can be active on a single interface at a time. When the firmware DCBX agent is active, software agents will not be able to receive or transmit LLDP frames.

When operating in firmware DCBX mode, the adapter is in an "always willing" state. DCB settings are applied on the adapter by transmitting a nonwilling configuration from the link partner. Typically this is a switch. For configuring DCBX parameters on a switch, please consult the switch manufacturer's documentation.



NOTE: Intel Ethernet Controller 500 Series adapters: You can configure DCB parameters using software LLDP/DCBX agents that interface with the Linux kernel's DCB Netlink API. We recommend using OpenLLDP as the DCBX agent. For more information, see the OpenLLDP man pages and <https://github.com/intel/openlldp>.

Malicious Driver Detection (MDD) for VFs

Some Intel Ethernet devices use Malicious Driver Detection (MDD) to detect malicious traffic from the VF and disable Tx/Rx queues or drop the offending packet until a VF driver reset occurs. You can view MDD messages in the PF's system log using the `dmesg` command.

- If the PF driver logs MDD events from the VF, confirm that the correct VF driver is installed.
- To restore functionality, you can manually reload the VF or VM.

MAC and VLAN Anti-Spoofing Feature for VFs

When a malicious driver on a Virtual Function (VF) interface attempts to send a spoofed packet, it is dropped by the hardware and not transmitted.

An interrupt is sent to the PF driver notifying it of the spoof attempt. When a spoofed packet is detected, the PF driver will send the following message to the system log (displayed by the `"dmesg"` command):

When a spoofed packet is detected the PF driver will send the following message to the system log (displayed by the `"dmesg"` command):

```
ixgbe <ethX>: ixgbe_spoof_check: n spoofed packets detected
```

where "X" is the PF interface number and "n" is number of spoofed packets.

This feature can be disabled for a specific VF:

```
# ip link set <pf dev> vf <vf id> spoofchk {off|on}
```

Setting MAC Address, VLAN and Rate Limit Using IProute2 Tool

You can set a MAC address of a Virtual Function (VF), a default VLAN and the rate limit using the IProute2 tool. Download the latest version of the iproute2 tool from Sourceforge if your version does not have all the features you require.

Wake on LAN (WoL) Support

Some adapters do not support Wake on LAN (WoL). To determine if your adapter supports WoL, run the following command:

```
# ethtool <ethX>
```

WoL is configured through the ethtool utility. ethtool is included with all versions of Red Hat after Red Hat 7.2. For other Linux distributions, download and install ethtool from the following website: <https://kernel.org/pub/software/network/ethtool/>.

For instructions on enabling WoL with ethtool, refer to the website listed above.

WoL will be enabled on the system during the next shutdown or reboot. For this driver version, in order to enable WoL, the driver must be loaded prior to shutting down or suspending the system.



NOTE: The Intel® Ethernet Converged Network Adapter X550-T1 and Intel® Ethernet Converged Network Adapter X550-T2 have a manageability/AUX power connector. These devices only support WoL if AUX power is supplied via this connector. Note that this is system and adapter specific. Some with this connector do not support WoL. Some systems do not provide the correct power connection. See your system documentation for details.

IEEE 1588 Precision Time Protocol (PTP) Hardware Clock (PHC)

Precision Time Protocol (PTP) is used to synchronize clocks in a computer network. PTP support varies among Intel devices that support this driver. Use 'ethtool -T <ethX>' to get a definitive list of PTP capabilities supported by the device.

Tunnel/Overlay Stateless Offloads

Supported tunnels and overlays include VXLAN, GENEVE, and others depending on hardware and software configuration. Stateless offloads are enabled by default.

To view the current state of all offloads:

```
# ethtool -k <ethX>
```

Virtual Function (VF) Tx Rate Limit

Use the ip command to configure the Tx rate limit for a VF from the PF interface.

For example, to set a Tx rate limit of 1000Mbps for VF 0:

```
# ip link set eth0 vf 0 rate 1000
```

Note that the limit is set per queue and not for the entire VF interface.

Interrupt Rate Limiting

This driver supports an adaptive interrupt throttle rate (ITR) mechanism that is tuned for general workloads. The user can customize the interrupt rate control for specific workloads, via ethtool, adjusting the number of microseconds between interrupts.

Syntax:

```
# ethtool -C <ethX> rx-usecs N
```

Values for N:

- 0 - no limit
- 1 - adaptive (default)
- 2-1022 - minimum microseconds between each interrupt

The range of 0-1022 microseconds provides an effective range of 978 to 500,000 interrupts per second. The underlying hardware supports granularity in 2us intervals at 1Gbps and 10Gbps and 20us at 100Mbps, so adjacent values may result in the same interrupt rate.

For lower CPU utilization:

- Lower Rx and Tx interrupts per queue using ethtool.
- Setting rx-usecs to 125 will limit interrupts to about 8,000 interrupts per second per queue:

```
# ethtool -C <ethX> rx-usecs 125
```

For reduced latency:

- Disable ITR by setting rx-usecs to 0 using ethtool:

```
# ethtool -C <ethX> rx-usecs 0
```

Known Issues

MAC Address of Virtual Function Changes Unexpectedly

If a Virtual Function's MAC address is not assigned in the host, then the VF (virtual function) driver will use a random MAC address. This random MAC address may change each time the VF driver is reloaded. You can assign a static MAC address in the host machine. This static MAC address will survive a VF driver reload.

Hardware Issues

For known hardware and troubleshooting issues, either refer to the "Release Notes" in your User Guide, or for more detailed information, go to <http://www.intel.com>.

In the search box enter your devices controller ID followed by "spec update". The specification update file has complete information on known hardware issues.

Software Issues



NOTE: After installing the driver, if your Intel® Ethernet Network Connection is not working, verify that you have installed the correct driver. Intel® Active Management Technology 2.0, 2.1, and 2.5 are not supported in conjunction with the Linux driver.

LRO and iSCSI Incompatibility

LRO is incompatible with iSCSI target or initiator traffic. A panic may occur when iSCSI traffic is received through the ixgbe driver with LRO enabled. To workaround this, the driver should be built and installed with:

```
# make CFLAGS_EXTRA=-DIXGBE_NO_LRO install
```

Multiple Interfaces on Same Ethernet Broadcast Network

Due to the default ARP behavior on Linux, it is not possible to have one system on two IP networks in the same Ethernet broadcast domain (non-partitioned switch) behave as expected. All Ethernet interfaces will respond to IP traffic for any IP address assigned to the system. This results in unbalanced receive traffic.

If you have multiple interfaces in a server, either turn on ARP filtering by entering:

```
# echo 1 > /proc/sys/net/ipv4/conf/all/arp_filter
```

This only works if your kernel's version is higher than 2.4.5.



NOTE: This setting is not saved across reboots. The configuration change can be made permanent by adding the following line to the file `/etc/sysctl.conf`:

```
net.ipv4.conf.all.arp_filter = 1
```

Another alternative is to install the interfaces in separate broadcast domains (either in different switches or in a switch partitioned to VLANs).

UDP Stress Test Dropped Packet Issue

Under small packet UDP stress with the ixgbe driver, the system may drop UDP packets due to socket buffers being full. Setting the driver Flow Control variables to the minimum may resolve the issue. You may also try increasing the kernel's default buffer sizes by changing the values in `/proc/sys/net/core/rmem_default` and `rmem_max`.

Rx Page Allocation Errors

'Page allocation failure. order:0' errors may occur under stress with kernels 2.6.25 and newer. This is caused by the way the Linux kernel reports this stressed condition.

DCB: Generic segmentation offload on causes bandwidth allocation issues

In order for DCB to work correctly, Generic Segmentation Offload (GSO), also known as software TSO, must be disabled using `ethtool`. Since the hardware supports TSO (hardware offload of segmentation), GSO will not be running by default. The GSO state can be queried with `ethtool` using `ethtool -k ethX`. When using 82598-based network connections, ixgbe driver only supports 16 queues on a platform with more than 16 cores.

Due to known hardware limitations, RSS can only filter in a maximum of 16 receive queues.

82599 and X540, and X550-based network connections support up to 64 queues.

Lower Than Expected Performance

Some PCIe x8 slots are actually configured as x4 slots. These slots have insufficient bandwidth for full line rate with dual port and quad port devices. In addition, if you put a PCIe v4.0 or v3.0-capable adapter into a PCIe v2.x slot, you cannot get full bandwidth. The driver detects this situation and writes one of the following messages in the system log:

"PCI-Express bandwidth available for this card is not sufficient for optimal performance. For optimal performance a x8 PCI-Express slot is required."

or

"PCI-Express bandwidth available for this device may be insufficient for optimal performance. Please move the device to a different PCI-e link with more lanes and/or higher transfer rate."

If this error occurs, moving your adapter to a true PCIe v3.0 x8 slot will resolve the issue.

ethtool may incorrectly display SFP+ fiber module as direct attached cable

Due to kernel limitations, port type can only be correctly displayed on kernel 2.6.33 or greater.

Under Redhat 5.4, system may crash when closing guest OS window after loading/unloading the Physical Function (PF) driver. Do not remove the ixgbe driver from Dom0 while Virtual Functions (VFs) are assigned to guests. VFs must first use the `xm "pci-detach"` command to hot-plug the VF device out of the VM it is assigned to or else shut down the VM.

Unloading Physical Function (PF) driver may cause kernel panic or system reboot when VM is running and VF is loaded on the VM. On pre-3.2 Linux kernels, unloading the Physical Function (PF) driver causes system reboots when the VM is running and VF is loaded on the VM. Do not unload the PF driver (ixgbe) while VFs are assigned to guests.

Running `ethtool -t ethX` command causes break between PF and test client

When there are active VFs, "ethtool -t" will only run the link test. The driver will also log in syslog that VFs should be shut down to run a full diagnostic test.

Unable to obtain DHCP lease on boot with RedHat

In configurations where the auto-negotiation process takes more than 5 seconds, the boot script may fail with the following message:

"ethX: failed. No link present. Check cable?"

This error may occur even though the presence of link can be confirmed using `ethtool ethx`. In this case, try setting "LINKDELAY=30" in `/etc/sysconfig/network-scripts/ifcfg-ethx`.

The same issue can occur during a network boot (via PXE) on RedHat distributions that use the dracut script:

"Warning: No carrier detected on interface <interface_name>"

In this case add "rd.net.timeout.carrier=30" at the kernel command line.



NOTE: Link time can vary. Adjust LINKDELAY value accordingly.

Host May Reboot after Removing PF when VF is Active in Guest

Using kernel versions earlier than 3.2, do not unload the PF driver with active VFs. Doing this will cause your VFs to stop working until you reload the PF driver and may cause a spontaneous reboot of your system.

Prior to unloading the PF driver, you must first ensure that all VFs are no longer active. Do this by shutting down all VMs and unloading the VF driver.

ixgbevf Linux* Driver for the Intel® 10 Gigabit Server Adapters

ixgbevf Overview

SR-IOV is supported by the ixgbevf driver, which should be loaded on both the host and VMs. This driver supports upstream kernel versions 2.6.30 (or higher) x86_64.

The ixgbevf driver supports 82599, X540, and X550 virtual function devices that can only be activated on kernels supporting SR-IOV. SR-IOV requires the correct platform and OS support.

The ixgbevf driver requires the ixgbe driver, version 2.0 or later. The ixgbevf driver supports virtual functions generated by the ixgbe driver with a max_vfs value of 1 or greater. For more information on the max_vfs parameter refer to the section on the [ixgbe](#) driver.

The guest OS loading the ixgbevf driver must support MSI-X interrupts.

This driver is only supported as a loadable module at this time. Intel is not supplying patches against the kernel source to allow for static linking of the driver. For questions related to hardware requirements, refer to the documentation supplied with your Intel 10GbE adapter. All hardware requirements listed apply to use with Linux.

ixgbevf Linux Base Driver Supported Adapters

The following Intel network adapters are compatible with the ixgbevf Linux driver in this release and can support up to 63 virtual functions per port.

- Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz
- Intel® Ethernet 10G 2P X540-t Adapter
- Intel® Ethernet 10G 2P X550-t Adapter
- Intel® Ethernet 10G 4P X550 rNDC
- Intel® Ethernet 10G 4P X550/I350 rNDC
- Intel® Ethernet 10G 4P X540/I350 rNDC
- Intel® Ethernet 10G 4P X520/I350 rNDC
- Intel® Ethernet 10G 2P X520-k bNDC
- Intel® Ethernet 10G 2P X520 Adapter
- Intel® Ethernet 10G X520 LOM

SR-IOV Capable Operating Systems

- Citrix XenServer 6.0 with Red Hat Enterprise Linux
- VMWare* ESXi* 6.x
- Red Hat* Enterprise Linux* (RHEL) 8.3
- Red Hat* Enterprise Linux* (RHEL) 8.2
- Red Hat* Enterprise Linux* (RHEL) 7.9
- Novell* SUSE* Linux Enterprise Server (SLES) 15 SP2

Building and Installation

To enable SR-IOV on your system:

1. Ensure both Virtualization and SR-IOV are enabled in the BIOS.
2. Install the Linux operating system. You can verify that the KVM driver is loaded by typing: `lsmod | grep -i kvm`
3. Load the Linux Base Driver using the modprobe command: `modprobe ixgbe option max_vfs=xx,yy`

xx and yy are the number of virtual functions you want to create. You must specify a number for each port with each parameter separated by a comma. For example, xx is the number of virtual functions for port 1; and yy, for port 2. You can create up to 63 functions per port.
4. Compile and install the ixgbevf driver for SR-IOV. This is loaded against the virtual functions created.

 **NOTE:** For VLANs, there is a limit of a total of 32 shared VLANs to 1 or more virtual functions.

There are three methods for installing the Linux driver:

- [Install from Source Code](#)
- [Install Using KMP RPM](#)
- [Install Using KMOD RPM](#)

Install from Source Code

To build a binary RPM* package of this driver, run 'rpmbuild -tb <filename.tar.gz>'. Replace <filename.tar.gz> with the specific filename of the driver.

NOTES:

- For the build to work properly it is important that the currently running kernel MATCH the version and configuration of the installed kernel source. If you have just recompiled your kernel, reboot the system.
- RPM functionality has only been tested in Red Hat distributions.

1. Download the base driver tar file to the directory of your choice. For example, use '/home/username/ixgbevf' or '/usr/local/src/ixgbevf'.
2. Untar/unzip the archive, where <x.x.x> is the version number for the driver tar:

```
# tar xzf ixgbevf-<x.x.x>.tar.gz
```

3. Change to the driver src directory, where <x.x.x> is the version number for the driver tar:

```
# cd ixgbevf-<x.x.x>/src/
```

4. Compile the driver module:

```
# make install
```

The binary will be installed as:

```
/lib/modules/<KERNEL_VERSION>/kernel/drivers/net/ixgbevf/ixgbevf.ko
```

The install locations listed above are the default locations. This might differ for various Linux distributions. For more information, see the ldistrib.txt file included in the driver tar.

5. Remove the old driver:

```
# rmmod ixgbevf
```

6. Install the module using the modprobe command:

```
# modprobe ixgbevf <parameter>=<value>
```

7. Update the system image:

```
dracut -f
```

8. Assign an IP address to and activate the Ethernet interface by entering the following, where <ethx> is the interface name:

```
# ifconfig <ethX> <IP_address> netmask <netmask> up
```

9. Verify that the interface works. Enter the following, where <IP_address> is the IP address for another machine on the same subnet as the interface that is being tested:

```
# ping <IP_address>
```

Install Using KMP RPM

The KMP RPMs update existing ixgbevf RPMs currently installed on the system. These updates are provided by SuSE in the SLES release. If an RPM does not currently exist on the system, the KMP will not install.

The RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
intel-<component name>-<component version>.<arch type>.rpm
```

For example, intel-ixgbevf-1.3.8.6-1.x86_64.rpm:

- ixgbevf is the component name
- 1.3.8.6-1 is the component version
- x86_64 is the architecture type

KMP RPMs are provided for supported Linux distributions. The naming convention for the included KMP RPMs is:

```
intel-<component name>-kmp-<kernel type>-<component version>_<kernel version>.<arch type>.rpm
```

For example, intel-ixgbevf-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm:

- ixgbevf is the component name
- default is the kernel type
- 1.3.8.6 is the component version
- 2.6.27.19_5-1 is the kernel version
- x86_64 is the architecture type

To install the KMP RPM, type the following two commands:

```
# rpm -i <rpm filename>
# rpm -i <kmp rpm filename>
```

For example, to install the ixgbevf KMP RPM package, type the following:

```
# rpm -i intel-ixgbevf-1.3.8.6-1.x86_64.rpm
# rpm -i intel-ixgbevf-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm
```

Install Using KMOD RPM

The KMOD RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
kmod-<driver name>-<version>-1.<arch type>.rpm
```

For example, kmod-ixgbevf-2.3.4-1.x86_64.rpm:

- ixgbevf is the driver name
- 2.3.4 is the version
- x86_64 is the architecture type

To install the KMOD RPM, go to the directory of the RPM and type the following command:

```
# rpm -i <rpm filename>
```

For example, to install the ixgbevf KMOD RPM package, type the following:

```
# rpm -i kmod-ixgbevf-2.3.4-1.x86_64.rpm
```

Command Line Parameters

If the driver is built as a module, the following optional parameters are used by entering them on the command line with the modprobe command using this syntax:

```
# modprobe ixgbevf [<option>=<VAL1>,<VAL2>,...]
```

For example:

```
# modprobe ixgbevf InterruptThrottleRate=16000,16000
```

The default value for each parameter is generally the recommended setting, unless otherwise noted.

The following table contains parameters and possible values for modprobe commands:

Parameter Name	Valid Range/Settings	Default	Description
InterruptThrottleRate	0, 1, 956 - 488,281	8000	0=off 1=dynamic <min_ITR>-<max_ITR> Use ethtool to control InterruptThrottleRate, as shown below: <pre># ethtool -C <ethX> rx-usecs N</pre> where N is the time in microseconds between each interrupt. Interrupt Throttle Rate controls the number of interrupts each interrupt vector can generate per second. Increasing ITR lowers latency at the cost of increased CPU utilization, though it may help throughput in some circumstances. • 0 = Setting InterruptThrottleRate to 0 turns off any interrupt moderation and may improve small packet latency. However, this is generally not suitable for bulk throughput traffic due to the increased CPU utilization of the higher interrupt rate. • 1 = Setting InterruptThrottleRate to Dynamic mode attempts to moderate interrupts per vector while maintaining very low latency. This can sometimes cause extra CPU utilization. If planning on deploying this driver in a latency sensitive environment, this parameter should be considered. • <min_ITR>-<max_ITR> = Setting InterruptThrottleRate to a value greater or equal to <min_ITR> will program the adapter to send at most that many interrupts per second, even if more packets have come in. This reduces interrupt load on the system and can lower CPU utilization under heavy load, but will increase latency as packets are not processed as quickly.  On 82599, and X540, and X550-based adapters, disabling InterruptThrottleRate will also result in the driver disabling HW RSC. • On 82598-based adapters, disabling InterruptThrottleRate will also result in disabling LRO (Large Receive Offloads).



NOTES:

- For more information about the InterruptThrottleRate parameter, see the application note at <http://www.in->

tel.com/design/network/applnotes/ap450.htm.

- A descriptor describes a data buffer and attributes related to the data buffer. This information is accessed by the hardware.

Additional Configurations

Configuring the Driver on Different Distributions

Configuring a network driver to load properly when the system is started is distribution dependent. Typically, the configuration process involves adding an alias line to `/etc/modules.conf` or `/etc/modprobe.conf` as well as editing other system start up scripts and/or configuration files. Many popular Linux distributions ship with tools to make these changes for you. To learn the proper way to configure a network device for your system, refer to your distribution documentation. If during this process you are asked for the driver or module name, the name for the Linux Base Driver for the device is `ixgbevf`.

For example, if you install the `ixgbevf` driver for two adapters (`eth0` and `eth1`) and want to set the interrupt mode to MSI-X and MSI, respectively, add the following to `modules.conf` or `/etc/modprobe.conf`:

```
# alias eth0 ixgbevf
# alias eth1 ixgbevf
# options ixgbevf IntMode=2,1
```

Viewing Link Messages

Link messages will not be displayed to the console if the distribution is restricting system messages. In order to see network driver link messages on your console, set `dmesg` to eight by entering the following:

```
# dmesg -n 8
```



NOTE: This setting is not saved across reboots.

ethtool

The driver utilizes the `ethtool` interface for driver configuration and diagnostics, as well as displaying statistical information. The latest `ethtool` version is required for this functionality. Download it at: <https://kernel.org/pub/software/network/ethtool/>

MACVLAN

This driver supports MACVLAN. Kernel support for MACVLAN can be tested by checking if the MACVLAN driver is loaded. You can run `'lsmod | grep macvlan'` to see if the MACVLAN driver is loaded or run `'modprobe macvlan'` to try to load the MACVLAN driver.



NOTE:

- In passthru mode, you can only set up one MACVLAN device. It will inherit the MAC address of the underlying PF (Physical Function) device.

NAPI

This driver supports NAPI (Rx polling mode). For more information on NAPI, see <https://wiki.linux-foundation.org/networking/napi>.

Known Issues

MAC Address of Virtual Function Changes Unexpectedly

If a Virtual Function's MAC address is not assigned in the host, then the VF (virtual function) driver will use a random MAC address. This random MAC address may change each time the VF driver is reloaded. You can assign a static MAC address in the host machine. This static MAC address will survive a VF driver reload.

Hardware Issues

For known hardware and troubleshooting issues, either refer to the "Release Notes" in your User Guide, or for more detailed information, go to <http://www.intel.com>.

In the search box enter your devices controller ID followed by "spec update". The specification update file has complete information on known hardware issues.

Software Issues



NOTE: After installing the driver, if your Intel® Ethernet Network Connection is not working, verify that you have installed the correct driver.

Compiling the Driver

When trying to compile the driver by running `make install`, the following error may occur: "Linux kernel source not configured - missing version.h"

To solve this issue, create the `version.h` file by going to the Linux source tree and entering:

```
# make include/linux/version.h
```

Multiple Interfaces on Same Ethernet Broadcast Network

Due to the default ARP behavior on Linux, it is not possible to have one system on two IP networks in the same Ethernet broadcast domain (non-partitioned switch) behave as expected. All Ethernet interfaces will respond to IP traffic for any IP address assigned to the system. This results in unbalanced receive traffic.

If you have multiple interfaces in a server, either turn on ARP filtering by entering:

```
# echo 1 > /proc/sys/net/ipv4/conf/all/arp_filter
```

This only works if your kernel's version is higher than 2.4.5.



NOTE: This setting is not saved across reboots. The configuration change can be made permanent by adding the following line to the file `/etc/sysctl.conf`:

```
net.ipv4.conf.all.arp_filter = 1
```

Another alternative is to install the interfaces in separate broadcast domains (either in different switches or in a switch partitioned to VLANs).

Rx Page Allocation Errors

'Page allocation failure. order:0' errors may occur under stress with kernels 2.6.25 and newer. This is caused by the way the Linux kernel reports this stressed condition.

Host May Reboot after Removing PF when VF is Active in Guest

Using kernel versions earlier than 3.2, do not unload the PF driver with active VFs. Doing this will cause your VFs to stop working until you reload the PF driver and may cause a spontaneous reboot of your system.

Prior to unloading the PF driver, you must first ensure that all VFs are no longer active. Do this by shutting down all VMs and unloading the VF driver.

i40e Linux Driver for the Intel Ethernet Controller 700 Series

i40e Overview

	NOTE: The kernel assumes that TC0 is available, and will disable Priority Flow Control (PFC) on the device if TC0 is not available. To fix this, ensure TC0 is enabled when setting up DCB on your switch.
	NOTE: If the physical function (PF) link is down, you can force link up (from the host PF) on any virtual functions (VF) bound to the PF. Note that this requires kernel support (Redhat kernel 3.10.0-327 or newer, upstream kernel 3.11.0 or newer, and associated iproute2 user space support). If the following command does not work, it may not be supported by your system. The following command forces link up on VF 0 bound to PF eth0: ip link set eth0 vf 0 state enable
	NOTE: Do not unload a port's driver if a Virtual Function (VF) with an active Virtual Machine (VM) is bound to it. Doing so will cause the port to appear to hang. Once the VM shuts down, or otherwise releases the VF, the command will complete.
	NOTE: In a virtualized environment, on Intel® Server Adapters that support SR-IOV, the virtual function (VF) may be subject to malicious behavior. Software-generated layer two frames, like IEEE 802.3x (link flow control), IEEE 802.1Qbb (priority based flow-control), and others of this type, are not expected and can throttle traffic between the host and the virtual switch, reducing performance. To resolve this issue, and to ensure isolation from unintended traffic streams, configure all SR-IOV enabled ports for VLAN tagging from the administrative interface on the PF. This configuration allows unexpected, and potentially malicious, frames to be dropped.

The i40e Linux* Base Driver for the Intel Ethernet Controller 700 Series family of adapters supports the 2.6.32 and newer kernels, and includes support for Linux supported x86_64 systems.

The following features are now available in supported kernels:

- VXLAN encapsulation
- Native VLANs
- Channel Bonding (teaming)
- Generic Receive Offload
- Data Center Bridging

Adapter teaming is implemented using the native Linux Channel bonding module. This is included in supported Linux kernels. Channel Bonding documentation can be found in the Linux kernel source: /Documentation/networking/bonding.txt

Use ethtool, lspci, or iproute2's ip command to obtain driver information. Instructions on updating ethtool can be found in the [Additional Configurations](#) section.

i40e Linux Base Driver Supported Devices

The following Intel network adapters are compatible with this driver:

- Intel® Ethernet 10G 4P X710-k bNDC
- Intel® Ethernet 10G 2P X710-k bNDC
- Intel® Ethernet 10G X710-k bNDC
- Intel® Ethernet Converged Network Adapter X710
- Intel® Ethernet Converged Network Adapter X710-T
- Intel® Ethernet 10G 4P X710/350 rNDC
- Intel® Ethernet 10G 4P X710 SFP+ rNDC
- Intel® Ethernet 10G X710 rNDC
- Intel® Ethernet Server Adapter X710-DA2 for OCP
- Intel® Ethernet 10G 2P X710 OCP
- Intel® Ethernet 10G 4P X710 OCP
- Intel® Ethernet 10G 2P X710-T2L-t OCP
- Intel® Ethernet 10G 4P X710-T4L-t OCP

- Intel® Ethernet 10G 2P X710-T2L-t Adapter
- Intel® Ethernet 10G 4P X710-T4L-t Adapter
- Intel® Ethernet 40G 2P XL710 QSFP+ rNDC
- Intel® Ethernet Converged Network Adapter XL710-Q2
- Intel® Ethernet 25G 2P XXV710 Adapter
- Intel® Ethernet 25G 2P XXV710 Mezz

Building and Installation

There are three methods for installing the Linux driver:

- [Install from Source Code](#)
- [Install Using KMP RPM](#)
- [Install Using KMOD RPM](#)

Install from Source Code

To build a binary RPM* package of this driver, run 'rpmbuild -tb <filename.tar.gz>'. Replace <filename.tar.gz> with the specific filename of the driver.



NOTES:

- For the build to work properly it is important that the currently running kernel MATCH the version and configuration of the installed kernel source. If you have just recompiled your kernel, reboot the system.
- RPM functionality has only been tested in Red Hat distributions.

1. Download the base driver tar file to the directory of your choice. For example, use '/home/username/i40e' or '/usr/local/src/i40e'.
2. Untar/unzip the archive, where <x.x.x> is the version number for the driver tar:

```
# tar xzf i40e-<x.x.x>.tar.gz
```

3. Change to the driver src directory, where <x.x.x> is the version number for the driver tar:

```
# cd i40e-<x.x.x>/src/
```

4. Compile the driver module:

```
# make install
```

The binary will be installed as:

```
/lib/modules/<KERNEL VERSION>/kernel/drivers/net/i40e/i40e.ko
```

The install locations listed above are the default locations. This might differ for various Linux distributions. For more information, see the ldistrib.txt file included in the driver tar.

5. Remove the old driver:

```
# rmmod i40e
```

6. Install the module using the modprobe command:

```
# modprobe i40e <parameter>=<value>
```



NOTE: For RHEL 7.5 or later, you must unload the i40iw driver before removing older i40e drivers.

7. Update the system image:

```
dracut -f
```

8. Assign an IP address to and activate the Ethernet interface by entering the following, where <ethx> is the interface name:

```
# ifconfig <ethX> <IP_address> netmask <netmask> up
```

9. Verify that the interface works. Enter the following, where <IP_address> is the IP address for another machine on the same subnet as the interface that is being tested:

```
# ping <IP_address>
```

Install Using KMP RPM

The KMP RPMs update existing i40e RPMs currently installed on the system. These updates are provided by SuSE in the SLES release. If an RPM does not currently exist on the system, the KMP will not install.

The RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
intel-<component name>-<component version>.<arch type>.rpm
```

For example, intel-i40e-1.3.8.6-1.x86_64.rpm:

- i40e is the component name
- 1.3.8.6-1 is the component version
- x86_64 is the architecture type

KMP RPMs are provided for supported Linux distributions. The naming convention for the included KMP RPMs is:

```
intel-<component name>-kmp-<kernel type>-<component version>_<kernel version>.<arch type>.rpm
```

For example, intel-i40e-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm:

- i40e is the component name
- default is the kernel type
- 1.3.8.6 is the component version
- 2.6.27.19_5-1 is the kernel version
- x86_64 is the architecture type

To install the KMP RPM, type the following two commands:

```
# rpm -i <rpm filename>
# rpm -i <kmp rpm filename>
```

For example, to install the i40e KMP RPM package, type the following:

```
# rpm -i intel-i40e-1.3.8.6-1.x86_64.rpm
# rpm -i intel-i40e-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm
```

Install Using KMOD RPM

The KMOD RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
kmod-<driver name>-<version>-1.<arch type>.rpm
```

For example, kmod-i40e-2.3.4-1.x86_64.rpm:

- i40e is the driver name
- 2.3.4 is the version
- x86_64 is the architecture type

To install the KMOD RPM, go to the directory of the RPM and type the following command:

```
# rpm -i <rpm filename>
```

For example, to install the i40e KMOD RPM package, type the following:

```
# rpm -i kmod-i40e-2.3.4-1.x86_64.rpm
```

Command Line Parameters

In general, ethtool and other OS specific commands are used to configure user changeable parameters after the driver is loaded. The i40e driver only supports the max_vfs kernel parameter on older kernels that do not have the standard sysfs interface. The only other module parameter is the debug parameter that can control the default logging verbosity of the driver.

If the driver is built as a module, the following optional parameters are used by entering them on the command line with the modprobe command using this syntax:

```
# modprobe i40e [<option>=<VAL1>]
```

For example:

```
# modprobe i40e max_vfs=7
```

The default value for each parameter is generally the recommended setting, unless otherwise noted.

The following table contains parameters and possible values for modprobe commands:

Parameter Name	Valid Range/Settings	Default	Description
max_vfs	1-32 (Intel Ethernet Controller X710 based devices) 1-64 (Intel Ethernet Controller XXV710/XL710 based devices)	0	This parameter adds support for SR-IOV. It causes the driver to spawn up to max_vfs worth of virtual functions. NOTE: This parameter is only used on kernel 3.7.x and below. On kernel 3.8.x and above, use sysfs to enable VFs. Also, for Red Hat distributions, this parameter is only used on version 6.6 and older. For version 6.7 and newer, use sysfs. For example, you can create 4 VFs as follows: <pre># echo 4 > /sys/class/net/<ethX>/device/sriov_numvfs</pre> To disable VFs, write 0 to the same file: <pre># echo 0 > /sys/class/net/<ethX>/device/sriov_numvfs</pre> The parameters for the driver are referenced by position. Thus, if you have a dual port adapter, or more than one adapter in your system, and want N virtual functions per port, you must specify a number for each port with each parameter separated by a comma. For example: <pre># modprobe i40e max_vfs=4</pre> This will spawn 4 VFs on the first port. <pre># modprobe i40e max_vfs=2,4</pre> This will spawn 2 VFs on the first port and 4 VFs on the second port.  NOTES: • Caution must be used in loading the driver with these parameters. Depending on your system configuration, number of slots, etc., it is impossible to predict in all cases where the positions would be on the command line. • Neither the device nor the driver control how VFs are mapped into config space. Bus layout will vary by operating system. On operating systems that support it, you can check sysfs to find the mapping. Some hardware configurations support fewer SR-IOV instances, as the whole Intel Ethernet Controller XL710 (all functions) is limited to 128 SR-IOV interfaces in total.

Parameter Name	Valid Range/Settings	Default	Description
			When SR-IOV mode is enabled, hardware VLAN filtering and VLAN tag stripping/insertion will remain enabled. Please remove the old VLAN filter before the new VLAN filter is added. For example: <pre># ip link set eth0 vf 0 vlan 100 // set vlan 100 for VF 0</pre> <pre># ip link set eth0 vf 0 vlan 0 // Delete vlan 100</pre> <pre># ip link set eth0 vf 0 vlan 200 // set a new vlan 200 for VF 0</pre>

Additional Configurations

ethtool

The driver utilizes the ethtool interface for driver configuration and diagnostics, as well as displaying statistical information. The latest ethtool version is required for this functionality. Download it at: <https://kernel.org/pub/software/network/ethtool/>

Viewing Link Messages

Link messages will not be displayed to the console if the distribution is restricting system messages. In order to see network driver link messages on your console, set dmesg to eight by entering the following:

```
# dmesg -n 8
```



NOTE: This setting is not saved across reboots.

Configuring the Driver on Different Distributions

Configuring a network driver to load properly when the system is started is distribution dependent. Typically, the configuration process involves adding an alias line to `/etc/modules.conf` or `/etc/modprobe.conf` as well as editing other system start up scripts and/or configuration files. Many popular Linux distributions ship with tools to make these changes for you. To learn the proper way to configure a network device for your system, refer to your distribution documentation. If during this process you are asked for the driver or module name, the name for the Linux Base Driver for the device is `i40e`.

Displaying VF Statistics on the PF

Use the following ethtool command to display the statistics for all VFs on the PF:

```
# ethtool -S <ethX>
```



NOTE: The output of this command is very large due to the large number of VF statistics and the maximum number of possible VFs.

The PF driver will display a subset of the VF's statistics, as provided by the VF driver, for all VFs that are configured. The PF will always print a statistics block for each of the possible VFs, and it will show a zero for all unconfigured VFs. VF stats are listed in a single block at the end of the PF statistics, using the following naming convention:

```
vf<XXX>.<statistic name>
```

Where:

- `<XXX>` is the VF number (for example, `vf008`).
- `<statistic name>` is the name of the statistic as supplied by the VF driver.

For example:

```
vf008.rx_bytes: 0
vf008.rx_unicast: 0
vf008.rx_multicast: 0
vf008.rx_broadcast: 0
vf008.rx_discards: 0
vf008.rx_unknown_protocol: 0
vf008.tx_bytes: 0
vf008.tx_unicast: 0
vf008.tx_multicast: 0
vf008.tx_broadcast: 0
vf008.tx_discards: 0
vf008.tx_errors: 0
```

Configuring VLAN Tagging on SR-IOV Enabled Adapter Ports

To configure VLAN tagging for the ports on an SR-IOV enabled adapter, use the following command. The VLAN configuration should be done before the VF driver is loaded or the VM is booted. The VF is not aware of the VLAN tag being inserted on transmit and removed on received frames (sometimes called "port VLAN" mode).

```
# ip link set dev <PF netdev id> vf <id> vlan <vlan id>
```

For example, the following will configure PF eth0 and the first VF on VLAN 10:

```
# ip link set dev eth0 vf 0 vlan 10
```

Setting the MAC Address for a VF

To change the MAC address for the specified VF:

```
# ip link set <ethX> vf 0 mac <address>
```

For example:

```
# ip link set <ethX> vf 0 mac 00:01:02:03:04:05
```

This setting lasts until the PF is reloaded.



NOTE: Assigning a MAC address for a VF from the host will disable any subsequent requests to change the MAC address from within the VM. This is a security feature. The VM is not aware of this restriction, so if this is attempted in the VM, it will trigger MDD events.

Trusted VFs and VF Promiscuous Mode

This feature allows you to designate a particular VF as trusted and allows that trusted VF to request selective promiscuous mode on the Physical Function (PF).

To set a VF as trusted or untrusted, enter the following command in the Hypervisor:

```
# ip link set dev eth0 vf 1 trust [on|off]
```

 **NOTE:** It's important to set the VF to trusted before setting promiscuous mode. If the VM is not trusted, the PF will ignore promiscuous mode requests from the VF. If the VM becomes trusted after the VF driver is loaded, you must make a new request to set the VF to promiscuous.

Once the VF is designated as trusted, use the following commands in the VM to set the VF to promiscuous mode.

- For promiscuous all: `# ip link set eth2 promisc on`
Where eth2 is a VF interface in the VM
- For promiscuous Multicast: `# ip link set eth2 allmulticast on`
Where eth2 is a VF interface in the VM

 **NOTE:** By default, the ethtool private flag `vf-true-promisc-support` is set to "off," meaning that promiscuous mode for the VF will be limited. To set the promiscuous mode for the VF to true promiscuous and allow the VF to see all ingress traffic, use the following command:

```
# ethtool --set-priv-flags p261p1 vf-true-promisc-support on
```

The `vf-true-promisc-support` private flag does not enable promiscuous mode; rather, it designates which type of promiscuous mode (limited or true) you will get when you enable promiscuous mode using the `ip link` commands above. Note that this is a global setting that affects the entire device. However, the `vf-true-promisc-support` private flag is only exposed to the first PF of the device. The PF remains in limited promiscuous mode (unless it is in MFP mode) regardless of the `vf-true-promisc-support` setting.

Next, add a VLAN interface on the VF interface.

```
# ip link add link eth2 name eth2.100 type vlan id 100
```

Note that the order in which you set the VF to promiscuous mode and add the VLAN interface does not matter (you can do either first). The result in this example is that the VF will get all traffic that is tagged with VLAN 100.

Virtual Function (VF) Tx Rate Limit

Use the `ip` command to configure the Tx rate limit for a VF from the PF interface.

For example, to set a Tx rate limit of 1000Mbps for VF 0:

```
# ip link set eth0 vf 0 rate 1000
```

Malicious Driver Detection (MDD) for VFs

Some Intel Ethernet devices use Malicious Driver Detection (MDD) to detect malicious traffic from the VF and disable Tx/Rx queues or drop the offending packet until a VF driver reset occurs. You can view MDD messages in the PF's system log using the `dmesg` command.

- If the PF driver logs MDD events from the VF, confirm that the correct VF driver is installed.
- To restore functionality, you can manually reload the VF or VM.

MAC and VLAN Anti-Spoofing Feature for VFs

When a malicious driver on a Virtual Function (VF) interface attempts to send a spoofed packet, it is dropped by the hardware and not transmitted.

When a spoofed packet is detected the PF driver will send the following message to the system log (displayed by the "`dmesg`" command):

This feature can be disabled for a specific VF:

```
# ip link set <pf dev> vf <vf id> spoofchk {off|on}
```

Intel® Ethernet Flow Director

Note: Intel Ethernet Flow Director parameters are only supported on kernel versions 2.6.30 or newer.

The Intel Ethernet Flow Director performs the following tasks:

- Directs receive packets according to their flows to different queues
- Enables tight control on routing a flow in the platform
- Matches flows and CPU cores for flow affinity
- Supports multiple parameters for flexible flow classification and load balancing (in SFP mode only)

An included script (`set_irq_affinity`) automates setting the IRQ to CPU affinity.

This driver supports the following flow types:

- IPv4
- TCPv4
- UDPv4

For a given flow type, it supports valid combinations of IP addresses (source or destination) and UDP/TCP ports (source and destination). For example, you can supply only a source IP address, a source IP address and a destination port, or any combination of one or more of these four parameters.

This driver allows you to filter traffic based on a user-defined flexible two-byte pattern and offset by using the `ethtool` `user-def` and `mask` fields. Only L3 and L4 flow types are supported for user-defined flexible filters. For a given flow type, you must clear all Intel Ethernet Flow Director filters before changing the input set (for that flow type).

To enable or disable the Intel Ethernet Flow Director:

```
# ethtool -K <ethX> ntuple <on|off>
```

When disabling `ntuple` filters, all the user programmed filters are flushed from the driver cache and hardware. All needed filters must be re-added when `ntuple` is re-enabled.

Application Targeted Routing (ATR) Perfect Filters

Intel Ethernet Flow Director ATR is enabled by default when the kernel is in multiple transmit queue mode. A rule is added when a TCP flow starts and is deleted when the flow ends. Because this would interfere with sideband TCP rules, the driver automatically disables ATR when a TCP rule is added via `ethtool` (sideband). ATR is automatically re-enabled when all TCP sideband rules are deleted or when sideband is disabled.

You can disable or enable ATR using the `ethtool` private flags interface. To view the current setting:

```
# ethtool --show-priv-flags <ethX>
```

To change the setting:

```
# ethtool --set-priv-flags <ethX> flow-director-atr [off|on]
```

Packets that match the ATR rules will increment the `port.fdir_atr_match` stat in `ethtool`. The current operational state of ATR is reflected by the stat `port.fdir_atr_status`.

Sideband Perfect Filters

Sideband Perfect Filters are used to direct traffic that matches specified characteristics. They are enabled through `ethtool`'s `ntuple` interface. To enable or disable these filters:

```
# ethtool -K <ethX> ntuple <off|on>
```

To display all of the active filters:

```
# ethtool -u <ethX>
```

To add a new filter:

```
# ethtool -U <ethX> flow-type <type> src-ip <ip> dst-ip <ip> src-port <port> dst-  
port <port> action <queue>
```

Where:

- <ethX> - the Ethernet device to program
- <type> - can be ip4, tcp4, udp4, or sctp4
- <ip> - the ip address to match on
- <port> - the port number to match on
- <queue> - the queue to direct traffic towards (-1 discards the matched traffic)

To delete a filter:

```
# ethtool -U <ethX> delete <N>
```

Where <N> is the filter ID displayed when printing all the active filters, and may also have been specified using "loc <N>" when adding the filter.

Examples:

To add a filter that directs packet to queue 2:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 action 2 [loc 1]
```

To set a filter using only the source and destination IP address:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 action 2
[loc 1]
```

To set a filter based on a user-defined pattern and offset:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 user-def
0x4FFFF action 2 [loc 1]
```

where the value of the user-def field contains the offset (4 bytes) and the pattern (0xffff).

To match TCP traffic sent from 192.168.0.1, port 5300, directed to 192.168.0.5, port 80, and then send it to queue 7:

```
# ethtool -U enp130s0 flow-type tcp4 src-ip 192.168.0.1 dst-ip 192.168.0.5 src-port
5300 dst-port 80 action 7
```

For each flow-type, the programmed filters must all have the same matching input set. For example, issuing the following two commands is acceptable:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 src-port 5300 action 7
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.5 src-port 55 action 10
```

Issuing the next two commands, however, is not acceptable, since the first specifies `src-ip` and the second specifies `dst-ip`:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 src-port 5300 action 7
# ethtool -U enp130s0 flow-type ip4 dst-ip 192.168.0.5 src-port 55 action 10
```

The second command will fail with an error. You may program multiple filters with the same fields, using different values, but, on one device, you may not program two tcp4 filters with different matching fields.

Matching on a subportion of a field is not supported by the driver, thus partial mask fields are not supported.

Flex Byte Flow Director Filters

The driver also supports matching user-defined data within the packet payload. This flexible data is specified using the "user-def" field of the ethtool command in the following way:

31	28	24	20	16	15	12	8	4	0
offset into packet payload					2 bytes of flexible data				

For example,

```
... user-def 0x4FFFF ...
```

tells the filter to look 4 bytes into the payload and match that value against 0xFFFF. The offset is based on the beginning of the payload, and not the beginning of the packet. Thus

```
flow-type tcp4 ... user-def 0x8BEAF ...
```

would match TCP/IPv4 packets which have the value 0xBEAF 8 bytes into the TCP/IPv4 payload.

Note that ICMP headers are parsed as 4 bytes of header and 4 bytes of payload. Thus to match the first byte of the payload, you must actually add 4 bytes to the offset. Also note that ip4 filters match both ICMP frames as well as raw (unknown) ip4 frames, where the payload will be the L3 payload of the IP4 frame.

The maximum offset is 64. The hardware will only read up to 64 bytes of data from the payload. The offset must be even because the flexible data is 2 bytes long and must be aligned to byte 0 of the packet payload.

The user-defined flexible offset is also considered part of the input set and cannot be programmed separately for multiple filters of the same type. However, the flexible data is not part of the input set and multiple filters may use the same offset but match against different data.

Filters to Direct Traffic to a Specific VF

It is possible to create filters that direct traffic to a specific Virtual Function. For older versions of ethtool, this depends on the "action" parameter. Specify the action as a 64-bit value, where the lower 32 bits represent the queue number, while the next 8 bits represent the VF ID. Note that 0 is the PF, so the VF identifier is offset by 1. For example:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 action 0x800000002 [loc 1]
```

The action field specifies to direct traffic to Virtual Function 7 (8 minus 1) into queue 2 of that VF.

Newer versions of ethtool (version 4.11 and later) use "vf" and "queue" parameters instead of the "action" parameter. Note that using the new ethtool "vf" parameter does not require the value to be offset by 1. This command is equivalent to the above example:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 vf 7 queue 2 [loc 1]
```

Note that these filters will not break internal routing rules, and will not route traffic that otherwise would not have been sent to the specified VF.

Cloud Filter Support

On a complex network that supports multiple types of traffic (such as for storage as well as cloud), cloud filter support allows you to send one type of traffic (for example, the storage traffic) to the Physical Function (PF) and another type (say, the cloud traffic) to a Virtual Function (VF). Because cloud networks are typically VXLAN/GENEVE-based, you can define a cloud filter to identify VXLAN/GENEVE packets and send them to a queue in the VF to be processed by the virtual machine (VM). Similarly, other cloud filters can be designed for various other traffic tunneling.



NOTE:

- Cloud filters are only supported when the underlying device is in Single Function per Port mode.
- The "action -1" option, which drops matching packets in regular Intel Ethernet Flow Director filters, is not available to drop packets when used with cloud filters.
- For IPv4 and ether flow-types, cloud filters cannot be used for TCP or UDP filters.
- Cloud filters can be used as a method for implementing queue splitting in the PF.

The following filters are supported:

- Cloud Filters
 - Inner MAC, Inner VLAN (for NVGRE, VXLAN or GENEVE packets)
 - Inner MAC, Inner VLAN, Tenant ID (for NVGRE, VXLAN or GENEVE packets)
 - Inner MAC, Tenant ID (NVGRE packet or VXLAN/GENEVE packets)

- Outer MAC L2 filter
- Inner MAC filter
- Outer MAC, Tenant ID, Inner MAC
- Application Destination IP
- Application Source-IP, Inner MAC
- ToQueue: Use MAC, VLAN to point to a queue
- L3 filters
 - Application Destination IP

Cloud filters are specified using ethtool's ntuple interface, but the driver uses user-def to determine whether to treat the filter as a cloud filter or a regular filter. To enable a cloud filter, set the highest bit of the user-def field, "user-def 0x8000000000000000" to enable the cloud features described below. This specifies to the driver to treat the filter specially and not treat it like the regular filters described above. Note that cloud filters also read the other bits in the user-def field separately so you cannot use the flexible data feature described above.

For regular Intel Ethernet Flow Director filters:

- No user-def specified or highest bit (bit 63) is 0. For example:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 dst-ip 192.168.0.109 action 6 loc
```

For L3 filters (non-tunneled packets):

- user-def 0x8000000000000000 (no Tenant ID/VNI specified in remaining bits of the user-def field)
- Only L3 parameters (src-IP, dst-IP) are considered
- For example:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.42.13 dst-ip 192.168.42.33 /
src-port 12344 dst-port 12344 user-def 0x8000000000000000 action /
0x200000000 loc 3
```

Redirect traffic coming from 192.168.42.13 port 12344 with destination 192.168.42.33 port 12344 into VF id 1, and call this "rule 3"

For cloud filters (tunneled packets):

- All other filters, including where Tenant ID/VNI is specified. The lower 32 bits of the user-def field can carry the tenant ID/VNI if required.
- The VF can be specified using the "action" field, just as regular filters described in the Flow Director Filter section above.
- Cloud filters can be defined with inner MAC, outer MAC, inner IP address, inner VLAN, and VNI as part of the cloud tuple. Cloud filters filter on destination (not source) MAC and IP. The destination and source MAC address fields in the ethtool command are overloaded as dst = outer, src = inner MAC address to facilitate tuple definition for a cloud filter.
- The 'loc' parameter specifies the rule number of the filter as being stored in the base driver.
- For example:

```
# ethtool -U enp130s0 flow-type ether dst 8b:9d:ed:6a:ce:43 \
src 1d:44:9d:54:da:de user-def 0x8000000000000022 loc 38 \
action 0x200000000
```

Redirect traffic on VXLAN using tunnel id 34 (hex 0x22) coming from outer MAC address 8b:9d:ed:6a:ce:43 and inner MAC address 1d:44:9d:54:da:de into VF id 1 and call this "rule 38".

Flow Control

Ethernet Flow Control (IEEE 802.3x) can be configured with ethtool to enable receiving and transmitting pause frames for this driver. When transmit is enabled, pause frames are generated when the receive packet buffer crosses a predefined threshold. When receive is enabled, the transmit unit will halt for the time delay specified when a pause frame is received.



NOTES:

- You must have a flow control capable link partner.
- This driver requires flow control on both the port and link partner. If flow control is disabled on one of the sides,

the port may appear to hang on heavy traffic.

Use `ethtool` to change the flow control settings.

To enable or disable Rx or Tx Flow Control:

```
# ethtool -A <ethX> rx <on|off> tx <on|off>
```



NOTE: This command only enables or disables Flow Control if auto-negotiation is disabled. If auto-negotiation is enabled, this command changes the parameters used for auto-negotiation with the link partner.

To enable or disable auto-negotiation:

```
# ethtool -s <ethX> autoneg <on|off>
```



NOTE: Flow Control auto-negotiation is part of link auto-negotiation. Depending on your device, you may not be able to change the auto-negotiation setting.

RSS Hash Flow

Allows you to set the hash bytes per flow type and any combination of one or more options for Receive Side Scaling (RSS) hash byte configuration.

```
# ethtool -N <ethX> rx-flow-hash <type> <option>
```

Where `<type>` is:

tcp4 signifying TCP over IPv4

udp4 signifying UDP over IPv4

tcp6 signifying TCP over IPv6

udp6 signifying UDP over IPv6

And `<option>` is one or more of:

s Hash on the IP source address of the Rx packet.

d Hash on the IP destination address of the Rx packet.

f Hash on bytes 0 and 1 of the Layer 4 header of the Rx packet.

n Hash on bytes 2 and 3 of the Layer 4 header of the Rx packet.

Application Device Queues (ADQ)

Application Device Queues (ADQ) allow you to dedicate one or more queues to a specific application. This can reduce latency for the specified application, and allow Tx traffic to be rate limited per application.

Requirements:

- Kernel version 4.19.58 or later
- The `sch_mqprio`, `act_mirred` and `cls_flower` modules must be loaded. For example:

```
# modprobe sch_mqprio
# modprobe act_mirred
# modprobe cls_flower
```

- The latest version of `iproute2`.

```
# cd iproute2
# ./configure
# make DESTDIR=/opt/iproute2 install
```

- NVM version 6.01 or later
- ADQ cannot be enabled when the following features are enabled: Data Center Bridging (DCB), Multiple Functions per Port (MFP), or Sideband Filters.
- If another driver (for example, DPDK) has set cloud filters, you cannot enable ADQ.

To create TCs on the interface:

 **NOTE:** Run all TC commands from the `../iproute2/tc/` directory.

1. Use the `tc` command to create traffic classes (TCs). You can create a maximum of 8 TCs per interface.

```
# tc qdisc add dev <ethX> root mqprio num_tc <tcs> map <priorities> queues <count1@offset1 ...> hw 1 mode channel shaper bw_rlimit min_rate <min_rate1 ...> max_rate <max_rate1 ...>
```

Where:

- `num_tc <tcs>`: The number of TCs to use.
- `map <priorities>`: The map of priorities to TCs. You can map up to 16 priorities to TCs.
- `queues <count1@offset1 ...>`: For each TC, `<num queues>@<offset>`. The max total number of queues for all TCs is the number of cores.
- `hw 1 mode channel`: 'channel' with 'hw' set to 1 is a new hardware offload mode in mqprio that makes full use of the mqprio options, the TCs, the queue configurations, and the QoS parameters.
- `shaper bw_rlimit`: For each TC, sets the minimum and maximum bandwidth rates. The totals must be equal to or less than the port speed. This parameter is optional and is required only to set up the Tx rates.
- `min_rate <min_rate1>`: Sets the minimum bandwidth rate limit for each TC.
- `max_rate <max_rate1 ...>`: Sets the maximum bandwidth rate limit for each TC. You can set a min and max rate together.

 **NOTE:** See the mqprio man page and the examples below for more information.

2. Verify the bandwidth limit using network monitoring tools such as `ifstat` or `sar -n DEV [interval] [number of samples]`

 **NOTE:** Setting up channels via `ethtool` (`ethtool -L`) is not supported when the TCs are configured using mqprio.

3. Enable hardware TC offload on the interface:

```
# ethtool -K <ethX> hw-tc-offload on
```

4. Apply TCs to ingress (Rx) flow of the interface:

```
# tc qdisc add dev <ethX> ingress
```

 **NOTE:**

- Tunnel filters are not supported in ADQ. If encapsulated packets do arrive in non-tunnel mode, filtering will be done on the inner headers. For example, for VXLAN traffic in non-tunnel mode, if PCTYPE is identified as a VXLAN encapsulated packet, then the outer headers are ignored. Therefore, inner headers are matched.
- If a TC filter on a PF matches traffic over a VF (on the PF), that traffic will be routed to the appropriate queue of the PF, and will not be passed on the VF. Such traffic will end up getting dropped higher up in the TCP/IP stack as it does not match PF address data.
- If traffic matches multiple TC filters that point to different TCs, that traffic will be duplicated and sent to all matching TC queues. The hardware switch mirrors the packet to a VSI list when multiple filters are matched.

Example:

 **NOTE:** See the `tc` and `tc-flower` man pages for more information on traffic control and TC flower filters.

To set up two TCs (`tc0` and `tc1`), with 16 queues each, priorities 0-3 for `tc0` and 4-7 for `tc1`, and max Tx rate set to 1Gbit for `tc0` and 3Gbit for `tc1`:

```
# tc qdisc add dev ens4f0 root mqprio num_tc 2 map 0 0 0 0 1 1 1 1 queues 16@0 16@16
hw 1 mode channel shaper bw_rlimit max_rate 1Gbit 3Gbit
```

Where:

- map 0 0 0 0 1 1 1 1: Sets priorities 0-3 to use tc0 and 4-7 to use tc1
- queues 16@0 16@16: Assigns 16 queues to tc0 at offset 0 and 16 queues to tc1 at offset 16

You can add multiple filters to the device, using the same recipe (and requires no additional recipe resources), either on the same interface or on different interfaces. Each filter uses the same fields for matching, but can have different match values.

```
tc filter add dev ethx protocol ip ingress prio 1 flower ip_proto tcp dst_port $app_port skip_sw
hw_tc 1
```

For example:

```
# tc filter add dev ethx protocol ip ingress prio 1 flower ip_proto tcp dst_port
5555 skip_sw hw_tc 1
```

EEE (Energy Efficient Ethernet)

A link between two EEE-compliant devices will result in periodic bursts of data followed by periods where the link is in an idle state. This Low Power Idle (LPI) state is supported at 2.5 Gbps and 5 Gbps link speeds.



NOTES:

- EEE support requires auto-negotiation.
- Both link partners must support EEE.
- EEE is not supported on all Intel® Ethernet Network devices or at all link speeds.

Example:

```
# ethtool --show-eee <ethX>
# ethtool --set-eee <ethX> [eee on|off]
```

Setting the link-down-on-close Private Flag

When the link-down-on-close private flag is set to "on", the port's link will go down when the interface is brought down using the `ifconfig <ethX> down` command.

Use `ethtool` to view and set link-down-on-close, as follows:

```
# ethtool --show-priv-flags <ethX>
# ethtool --set-priv-flags <ethX> link-down-on-close [on|off]
```

Jumbo Frames

Jumbo Frames support is enabled by changing the MTU to a value larger than the default of 1500 bytes. Use the `ifconfig` command to increase the MTU size. For example, enter the following where `<ethX>` is the interface number::

```
# ifconfig <ethX> mtu 9000 up
```

Alternatively, you can use the `ip` command as follows:

```
# ip link set mtu 9000 dev <ethX>
# ip link set up dev <ethX>
```

This setting is not saved across reboots. The setting change can be made permanent by adding 'MTU = 9000' to the following file:

- `/etc/sysconfig/network-scripts/ifcfg-<ethX>` for RHEL
- `/etc/sysconfig/network/<config_file>` for SLES

NOTES:

- The maximum MTU setting for Jumbo Frames is 9710 bytes. This value coincides with the maximum Jumbo Frames size of 9728 bytes.
- This driver will attempt to use multiple page sized buffers to receive each jumbo packet. This should help to avoid buffer starvation issues when allocating receive packets.
- Packet loss may have a greater impact on throughput when you use jumbo frames. If you observe a drop in performance after enabling jumbo frames, enabling flow control may mitigate the issue.

Speed and Duplex Configuration

In addressing speed and duplex configuration issues, you need to distinguish between copper-based adapters and fiber-based adapters.

In the default mode, an Intel® Ethernet Network Adapter using copper connections will attempt to auto-negotiate with its link partner to determine the best setting. If the adapter cannot establish link with the link partner using auto-negotiation, you may need to manually configure the adapter and link partner to identical settings to establish link and pass packets. This should only be needed when attempting to link with an older switch that does not support auto-negotiation or one that has been forced to a specific speed or duplex mode. Your link partner must match the setting you choose. 1 Gbps speeds and higher cannot be forced. Use the autonegotiation advertising setting to manually set devices for 1 Gbps and higher.

Speed, duplex, and autonegotiation advertising are configured through the ethtool utility. ethtool is included with all versions of Red Hat after Red Hat 7.2. To see the speed configurations your device supports, run the following:

```
# ethtool <ethX>
```



CAUTION: Only experienced network administrators should force speed and duplex or change autonegotiation advertising manually. The settings at the switch must always match the adapter settings. Adapter performance may suffer or your adapter may not operate if you configure the adapter differently from your switch.



NOTE: You cannot set the speed for devices based on the Intel® Ethernet Network Adapter XXV710 based devices.

NAPI

This driver supports NAPI (Rx polling mode). For more information on NAPI, see <https://wiki.linux-foundation.org/networking/napi>.

IEEE 1588 Precision Time Protocol (PTP) Hardware Clock (PHC)

Precision Time Protocol (PTP) is used to synchronize clocks in a computer network. PTP support varies among Intel devices that support this driver. Use 'ethtool -T <ethX>' to get a definitive list of PTP capabilities supported by the device.

Tunnel/Overlay Stateless Offloads

Supported tunnels and overlays include VXLAN, GENEVE, and others depending on hardware and software configuration. Stateless offloads are enabled by default.

To view the current state of all offloads:

```
# ethtool -k <ethX>
```

Data Center Bridging (DCB)

Note: The kernel assumes that TC0 is available, and will disable Priority Flow Control (PFC) on the device if TC0 is not available. To fix this, ensure TC0 is enabled when setting up DCB on your switch.

DCB is a configuration Quality of Service implementation in hardware. It uses the VLAN priority tag (802.1p) to filter traffic. That means that there are 8 different priorities that traffic can be filtered into. It also enables priority flow control (802.1Qbb) which can limit or eliminate the number of dropped packets during network stress. Bandwidth can be allocated to each of these priorities, which is enforced at the hardware level (802.1Qaz).

Adapter firmware implements LLDP and DCBX protocol agents as per 802.1AB and 802.1Qaz respectively. There are potentially two DCBX modes on Linux, depending on the underlying PF device:

- Intel Ethernet Controller 700 Series adapters only support firmware DCBX mode. They do not support software DCBX.

DCB parameters can be established via a firmware LLDP/DCBX agent. Only one LLDP/DCBX agent can be active on a single interface at a time. When the firmware DCBX agent is active, software agents will not be able to receive or transmit LLDP frames.

See the "FW-LLDP (Firmware Link Layer Discovery Protocol)" subsection for this driver for the ethtool commands to query the status of the firmware LLDP/DCBX agent.

When operating in firmware DCBX mode, the adapter is in an "always willing" state. DCB settings are applied on the adapter by transmitting a nonwilling configuration from the link partner. Typically this is a switch. For configuring DCBX parameters on a switch, please consult the switch manufacturer's documentation.



NOTE:

Intel Ethernet Controller 700 Series adapters:

- You can disable the firmware agent using ethtool private flags. For DCB to be operational, the firmware LLDP agent is required.
- Software configuration of DCBX parameters via dcbtool/lldptool is not supported.
- The driver implements the DCB netlink interface layer to allow user-space to communicate with the driver and query DCB configuration for the port.

FW-LLDP (Firmware Link Layer Discovery Protocol)

Use ethtool to change FW-LLDP settings. The FW-LLDP setting is per port and persists across boots.

To enable LLDP:

```
# ethtool --set-priv-flags <ethX> disable-fw-lldp off
```

To disable LLDP:

```
# ethtool --set-priv-flags <ethX> disable-fw-lldp on
```

To check the current LLDP setting:

```
# ethtool --show-priv-flags <ethX>
```



NOTE: You must enable the UEFI HII "LLDP Agent" attribute for this setting to take effect. If "LLDP AGENT" is set to disabled, you cannot enable it from the OS.

Forward Error Correction (FEC)

Allows you to set the Forward Error Correction (FEC) mode. FEC improves link stability, but increases latency. Many high quality optics, direct attach cables, and backplane channels provide a stable link without FEC.



NOTE: For devices to benefit from this feature, link partners must have FEC enabled.

On kernels older than 4.14, use the following private flags to disable FEC modes:

- rs-fec (0 to disable, 1 to enable)
- base-r-fec (0 to disable, 1 to enable)

On kernel 4.14 or later, use ethtool to get/set the following FEC modes:

- No FEC
- Auto FEC
- BASE-R FEC
- RS FEC

Dynamic Device Personalization

Dynamic Device Personalization (DDP) allows you to change the packet processing pipeline of a device by applying a profile package to the device at runtime. Profiles can be used to, for example, add support for new protocols, change existing protocols, or change default settings. DDP profiles can also be rolled back without rebooting the system.

Requirements:

- Intel Ethernet X710/XXV710/XL710 adapter
- Firmware 6.0 or newer
- RHEL 7.5 or later or Linux Kernel 4.0.1 or newer

To apply a profile, copy it first to the `intel/i40e/ddp` directory relative to your firmware root (usually `/lib/firmware` or `/lib/firmware/updates`).

For example:

```
/lib/firmware/intel/i40e/ddp
```

Then use the `ethtool -f` flag with region 100:

```
# ethtool -f <ethX> <profile name> 100
```

For example:

```
# ethtool -f eth0 gtp.pkgo 100
```

You can roll back to a previously loaded profile using `-` instead of profile name:

```
# ethtool -f <ethX> - 100
```

For example:

```
# ethtool -f eth0 - 100
```

For every rollback request one profile will be removed, from last to first (LIFO) order.



NOTE:

- DDP profiles are loaded only on the interface corresponding to first physical function of the device (PF0), but the configuration is applied to all ports of the adapter.
- DDP profiles are not persistent. A system reboot will reset the device to its default configuration.
- DDP profiles are NOT automatically unloaded when the driver is unbound/unloaded. Please note that subsequent driver reload may corrupt the profile configuration during its initialization and is NOT recommended.
- DDP profiles should be manually rolled-back before driver unload/unbind if the intention is to start with clean HW configuration.
- Exercise caution while loading DDP profiles. Attempting to load files other than DDP profiles provided by Intel may cause system instability, system crashes, or system hangs.

More details about Dynamic Device Personalization can be found on the Intel Developer Zone site:

<https://software.intel.com/en-us/articles/dynamic-device-personalization-for-intel-ethernet-700-series>

SR-IOV Hypervisor Management Interface

The `sysfs` file structure below supports the SR-IOV hypervisor management interface.

`/sys/class/net/<ethX>/device/sriov` (see 1 below)

```
+- [VF-id, 0 .. 255] (see 2 below)
```

```
| +- trunk
```

```
| +--- vlan_mirror
| +--- egress_mirror
| +--- ingress_mirror
| +--- mac_anti_spoof
| +--- vlan_anti_spoof
| +--- loopback
| +--- mac
| +--- mac_list
| +--- promisc
| +--- vlan_strip
| +--- stats
| +--- link_state
| +--- max_tx_rate
| +--- min_tx_rate
| +--- spoofcheck
| +--- trust
| +--- vlan
```

**NOTES:**

1. kobject started from “sriov” is not available from existing kernel sysfs, and it requires device driver to implement this interface.
2. maximum number of SR-IOV instances is 256. The actual number of instances created depends on the value set for /sys/bus/devices/<device pci address>/sriov_numvfs

SR-IOV hypervisor functions:

- trunk
 - Supports two operations:
 - add: adds one or more VLAN id into VF VLAN filtering.
 - rem: removes VLAN ids from the VF VLAN filtering list.
 - Example 1: add multiple VLAN tags, VLANs 2,4,5,10-20, by PF, p1p2, on a selected VF, 1, for filtering, with the sysfs support:


```
# echo add 2,4,5,10-20 > /sys/class/net/plp2/device/sriov/1/trunk
```
 - Example 2: remove VLANs 5, 11-13 from PF p1p2 VF 1 with sysfs:


```
# echo rem 5,11-13 > /sys/class/net/plp2/device/sriov/1/trunk
```
 - Note: for rem, if VLAN id is not on the VLAN filtering list, the VLAN id will be ignored.
- vlan_mirror
 - Supports both ingress and egress traffic mirroring.
 - Example 1: mirror traffic based upon VLANs 2,4,6,18-22 to VF 3 of PF p1p1.


```
# echo add 2,4,6,18-22 > /sys/class/net/plp1/device/sriov/3/vlan_mirror
```
 - Example 2: remove VLAN 4, 15-17 from traffic mirroring at destination VF 3.


```
# echo rem 15-17 > /sys/class/net/plp1/device/sriov/3/vlan_mirror
```
 - Example 3: remove all VLANs from mirroring at VF 3.


```
# echo rem 0 - 4095 > /sys/class/net/plp1/device/sriov/3/vlan_mirror
```
- egress_mirror
 - Supports egress traffic mirroring.
 - Example 1: add egress traffic mirroring on PF p1p2 VF 1 to VF 7.


```
# echo add 7 > /sys/class/net/plp2/device/sriov/1/egress_mirror
```

- Example 2: remove egress traffic mirroring on PF p1p2 VF 1 to VF 7.
echo rem 7 > /sys/class/net/plp2/device/sriov/1/egress_mirror
- ingress_mirror
 - Supports ingress traffic mirroring.
 - Example 1: mirror ingress traffic on PF p1p2 VF 1 to VF 7.
echo add 7 > /sys/class/net/plp2/device/sriov/1/ingress_mirror
 - Example 2: show current ingress mirroring configuration.
cat /sys/class/net/plp2/device/sriov/1/ingress_mirror
- mac_anti_spoof
 - Supports Enable/Disable MAC anti-spoof. Allows VFs to transmit packets with any SRC MAC, which is needed for some L2 applications as well as vNIC bonding within VMs if set to OFF.
 - Example 1: enable MAC anti-spoof for PF p2p1 VF 1.
echo ON /sys/class/net/plp2/device/sriov/1/mac_anti_spoof
 - Example 2: disable MAC anti-spoof for PF p2p1 VF 1.
echo OFF /sys/class/net/plp2/device/sriov/1/mac_anti_spoof
- vlan_anti_spoof
 - Supports Enable/Disable VLAN anti-spoof. Allows VFs to transmit packets only with VLAN tag specified in “trunk” settings, also will not allow to transmit “untagged” packets if set to ON. Violation have to increment tx_spoof stats counter.
 - Example 1: enable VLAN anti-spoof for PF p2p1 VF 1.
echo ON /sys/class/net/plp2/device/sriov/1/vlan_anti_spoof
 - Example 2: disable VLAN anti-spoof for PF p2p1 VF 1.
echo OFF /sys/class/net/plp2/device/sriov/1/vlan_anti_spoof
- loopback
 - Supports Enable/Disable VEB/VEPA (Local loopback).
 - Example 1: allow traffic switching between VFs on the same PF.
echo ON > /sys/class/net/plp2/device/sriov/loopback
 - Example 2: send Hairpin traffic to the switch to which the PF is connected.
echo OFF > /sys/class/net/plp2/device/sriov/loopback
 - Example 3: show loopback configuration.
cat /sys/class/net/plp2/device/sriov/loopback
- mac
 - Supports setting default MAC address. If MAC address is set by this command, the PF will not allow VF to change it using an MBOX request.
 - Example 1: set default MAC address to VF 1.
echo "00:11:22:33:44:55" > /sys/class/net/plp2/device/sriov/1/mac
 - Example 2: show default MAC address.
cat /sys/class/net/plp2/device/sriov/1/mac
- mac_list
 - Supports adding additional MACs to the VF. The default MAC is taken from "ip link set p1p2 vf 1 mac 00:11:22:33:44:55" if configured. If not, a random address is assigned to the VF by the NIC. If the MAC is configured using the IP LINK command, the VF cannot change it via MBOX/AdminQ requests.
 - Example 1: add mac 00:11:22:33:44:55 and 00:66:55:44:33:22 to PF p1p2 VF 1.
echo add "00:11:22:33:44:55,00:66:55:44:33:22" > /sys/class/net/plp2/device/sriov/1/mac_list
 - Example 2: delete mac 00:11:22:33:44:55 from above VF device.
echo rem 00:11:22:33:44:55 > /sys/class/net/plp2/device/sriov/1/mac_list
 - Example 3: display a VF MAC address list.
cat /sys/class/net/plp2/device/sriov/1/mac_list
- promisc
 - Supports setting/unsetting VF device unicast promiscuous mode and multicast promiscuous mode.
 - Example 1: set MCAST promiscuous on PF p1p2 VF 1.
echo add mcast > /sys/class/net/plp2/device/sriov/1/promisc
 - Example 2: set UCAST promiscuous on PF p1p2 VF 1.
echo add ucast > /sys/class/net/plp2/device/sriov/1/promisc

- Example 3: unset MCAST promiscuous on PF p1p2 VF 1.
echo rem mcast > /sys/class/net/plp2/device/sriov/1/promisc
- Example 4: show current promiscuous mode configuration.
cat /sys/class/net/plp2/device/sriov/1/promisc
- vlan_strip
 - Supports enabling/disabling VF device outer VLAN stripping
 - Example 1: enable VLAN strip on VF 3.
echo ON > /sys/class/net/plp1/device/sriov/3/vlan_strip
 - Example 2: disable VLAN stripping VF 3.
echo OFF > /sys/class/net/plp1/device/sriov/3/vlan_strip
- stats
 - Supports getting VF statistics
 - Example 1: display stats of VF 1.
cat /sys/class/net/plp2/device/sriov/1/stats
- link_state
 - Sets/displays link status.
 - Example 1: display link status on link speed.
cat /sys/class/net/plp2/device/sriov/1/link_state
 - Example 2: set VF 1 to track status of PF link.
echo auto > /sys/class/net/plp2/device/sriov/1/link_state
 - Example 3: disable VF 1.
echo disable > /sys/class/net/plp2/device/sriov/1/link_state

Performance Optimization

The driver defaults are meant to fit a wide variety of workloads. If further optimization is required, we recommend experimenting with the following settings.

Small Frame Sizes

For better performance when processing small (64B) frame sizes:

1. Try enabling Hyper threading in the BIOS in order to increase the number of logical cores in the system.
2. Increase the number of queues available to the adapter:
ethtool -L

IRQ to Adapter Queue Alignment

Pin the adapter's IRQs to specific cores by disabling the irqbalance service and using the included set_irq_affinity script. Please see the script's help text for further options.

The following settings will distribute the IRQs across all the cores evenly:

```
# scripts/set_irq_affinity -X all <interfacel> , [ <interface2>, ... ]
```

The following settings will distribute the IRQs across all the cores that are local to the adapter (same NUMA node):

```
# scripts/set_irq_affinity -X local <interfacel> , [ <interface2>, ... ]
```

For very CPU-intensive workloads, we recommend pinning the IRQs to all cores.

Rx Descriptor Ring Size

To reduce the number of Rx packet discards, increase the number of Rx descriptors for each Rx ring using ethtool.

Check if the interface is dropping Rx packets due to buffers being full (rx_dropped.nic means there is no PCIe bandwidth):

```
# ethtool -S <interface> | grep "rx_dropped"
```

Increase the number of Rx descriptors for each Rx ring using ethtool. This may help reduce Rx packet drops at the expense of system resources:

```
# ethtool -G <interface> rx N
```

Where N is the desired number of rings

Interrupt Rate Limiting

This driver supports an adaptive interrupt rate mechanism that is tuned for general workloads. The user can customize the interrupt rate control for specific workloads, via ethtool, adjusting the number of microseconds between interrupts.

To set the interrupt rate manually, you must disable adaptive mode:

```
# ethtool -C <interface> adaptive-rx off adaptive-tx off
```

For lower CPU utilization:

1. Disable adaptive ITR and lower Rx and Tx interrupts. The examples below affect every queue of the specified interface.
2. Setting rx-usecs and tx-usecs to 125 will limit interrupts to about 8000 interrupts per second per queue:

```
# ethtool -C <interface> adaptive-rx off adaptive-tx off rx-usecs 125 tx-usecs 125
```

For reduced latency:

Disable adaptive ITR and ITR by setting rx-usecs and tx-usecs to 0 using ethtool:

```
# ethtool -C <interface> adaptive-rx off adaptive-tx off rx-usecs 0 tx-usecs 0
```

Known Issues

Unexpected i40iw error messages in dmesg

If you install newer i40e drivers over the drivers in your Linux distro, you may see i40iw error messages in dmesg. This is because the i40iw drivers must be updated at the same time as the i40e drivers. These messages may be ignored if you do not use iWARP in your configuration.

Linux bonding fails with Virtual Functions bound to an Intel® Ethernet Controller 700 series based device

If you bind Virtual Functions (VFs) to an Intel® Ethernet Controller 700 series based device, the VF slaves may fail when they become the active slave. If the MAC address of the VF is set by the PF (Physical Function) of the device, when you add a slave, or change the active-backup slave, Linux bonding tries to sync the backup slave's MAC address to the same MAC address as the active slave. Linux bonding will fail at this point. This issue will not occur if the VF's MAC address is not set by the PF.

X710/XXV710 devices fail to enable MAX VFs when NPAR and SR-IOV are enabled

X710/XXV710 devices fail to enable Max VFs (64) when NPAR and SR-IOV are enabled. An error from i40e is logged that says "add vsi failed for VF N, aq_err 16". To workaround the issue, enable less than 64 virtual functions (VFs).

ip link show command shows incorrect VF MAC if VF MAC was set from VF side

Executing the command "ip link show" only shows MAC addresses if they are set by the PF. Otherwise, it shows all zeros.

This is expected behavior. The PF driver is passing zeroes to the VF driver that the VF driver can generate its own random MAC address and report it to the guest OS. Without this feature, some guest operating systems will incorrectly assign the VF a new interface name each time they reboot.

IPv6/UDP checksum offload does not work on some older kernels

Some distributions with older kernels do not properly enable IPv6/UDP checksum offload. To use IPv6 checksum offload, it may be necessary to upgrade to a newer kernel.

depmod warning messages about unknown symbol during installation

During driver installation, you may see depmod warning messages referring to unknown symbols `i40e_register_client` and `i40e_unregister_client`. These messages are informational only; no user action is required. The installation should complete successfully.

Error: <ifname> selects TX queue XX but real number of TX queues is YY

When configuring the number of queues under heavy traffic load, you may see an error message stating "<ifname> selects TX queue XX, but real number of TX queues is YY". This message is informational only and does not affect functionality.

Fixing Performance Issues When Using IOMMU in Virtualized Environments

The IOMMU feature of the processor prevents I/O devices from accessing memory outside the boundaries set by the OS. It also allows devices to be directly assigned to a Virtual Machine. However, IOMMU may affect performance, both in latency (each DMA access by the device must be translated by the IOMMU) and in CPU utilization (each buffer assigned to every device must be mapped in the IOMMU).

If you experience significant performance issues with IOMMU, try using it in "passthrough" mode by adding the following to the kernel boot command line:

```
intel_iommu=on iommu=pt
```



NOTE: This mode enables remapping for assigning devices to VMs, providing near-native I/O performance, but does not provide the additional memory protection.

Transmit hangs leading to no traffic

Disabling flow control while the device is under stress may cause tx hangs and eventually lead to the device no longer passing traffic. You must reboot the system to resolve this issue.

Incomplete messages in the system log

The NVMMUpdate utility may write several incomplete messages in the system log.

These messages take the form:

```
in the driver Pci Ex config function byte index 114
in the driver Pci Ex config function byte index 115
```

These messages can be ignored.

Bad checksum counter incorrectly increments when using VxLAN

When passing non-UDP traffic over a VxLAN interface, the `port.rx_csum_bad` counter increments for the packets.

Statistic counters reset when promiscuous mode is changed

Changing promiscuous mode triggers a reset of the physical function driver. This will reset the statistic counters.

Virtual machine does not get link

If the virtual machine has more than one virtual port assigned to it, and those virtual ports are bound to different physical ports, you may not get link on all of the virtual ports. The following command may work around the issue:

```
# ethtool -r <ethX>
```

Where <ethX> is the PF interface in the host, for example: p5p1. You may need to run the command more than once to get link on all virtual ports.

MAC Address of Virtual Function Changes Unexpectedly

If a Virtual Function's MAC address is not assigned in the host, then the VF (virtual function) driver will use a random MAC address. This random MAC address may change each time the VF driver is reloaded. You can assign a static MAC address in the host machine. This static MAC address will survive a VF driver reload.

Changing the number of Rx or Tx queues with ethtool -L may cause a kernel panic

Changing the number of Rx or Tx queues with ethtool -L while traffic is flowing and the interface is up may cause a kernel panic. Bring the interface down first to avoid the issue. For example:

```
# ip link set <ethX> down
# ethtool -L <ethX> combined 4
```

Adding an Intel Ethernet Flow Director Sideband rule fails incorrectly

If you try to add an Intel Ethernet Flow Director rule when no more sideband rule space is available, the driver logs an error that the rule could not be added, but ethtool returns success. You can remove rules to free up space. In addition, remove the rule that failed. This will evict it from the driver's cache.

Intel Ethernet Flow Director Sideband Logic adds duplicate filter

The Intel Ethernet Flow Director Sideband Logic adds a duplicate filter in the software filter list if the location is not specified or the specified location differs from the previous location but has the same filter criteria. In this case, the second of the two filters that appear is the valid one in hardware and it decides the filter action.

Multiple Interfaces on Same Ethernet Broadcast Network

Due to the default ARP behavior on Linux, it is not possible to have one system on two IP networks in the same Ethernet broadcast domain (non-partitioned switch) behave as expected. All Ethernet interfaces will respond to IP traffic for any IP address assigned to the system. This results in unbalanced receive traffic.

If you have multiple interfaces in a server, either turn on ARP filtering by entering:

```
# echo 1 > /proc/sys/net/ipv4/conf/all/arp_filter
```

This only works if your kernel's version is higher than 2.4.5.



NOTE: This setting is not saved across reboots. The configuration change can be made permanent by adding the following line to the file /etc/sysctl.conf:

```
net.ipv4.conf.all.arp_filter = 1
```

Another alternative is to install the interfaces in separate broadcast domains (either in different switches or in a switch partitioned to VLANs).

UDP Stress Test Dropped Packet Issue

Under small packet UDP stress with the i40e driver, the system may drop UDP packets due to socket buffers being full. Setting the driver Flow Control variables to the minimum may resolve the issue. You may also try increasing the kernel's default buffer sizes by changing the values in /proc/sys/net/core/rmem_default and rmem_max

Unplugging Network Cable While ethtool -p is Running

In kernel versions 2.6.32 and newer, unplugging the network cable while ethtool -p is running will cause the system to become unresponsive to keyboard commands, except for control-alt-delete. Restarting the system appears to be the only remedy.

Rx Page Allocation Errors

'Page allocation failure. order:0' errors may occur under stress with kernels 2.6.25 and newer. This is caused by the way the Linux kernel reports this stressed condition.

Lower Than Expected Performance

Some PCIe x8 slots are actually configured as x4 slots. These slots have insufficient bandwidth for full line rate with dual port and quad port devices. In addition, if you put a PCIe v4.0 or v3.0-capable adapter into a PCIe v2.x slot, you cannot get full bandwidth. The driver detects this situation and writes one of the following messages in the system log:

"PCI-Express bandwidth available for this card is not sufficient for optimal performance. For optimal performance a x8 PCI-Express slot is required."

or

"PCI-Express bandwidth available for this device may be insufficient for optimal performance. Please move the device to a different PCI-e link with more lanes and/or higher transfer rate."

If this error occurs, moving your adapter to a true PCIe v3.0 x8 slot will resolve the issue.

ethtool may incorrectly display SFP+ fiber module as direct attached cable

Due to kernel limitations, port type can only be correctly displayed on kernel 2.6.33 or greater.

Running ethtool -t ethX command causes break between PF and test client

When there are active VFs, "ethtool -t" performs a full diagnostic. In the process, it resets itself and all attached VFs. The VF drivers encounter a disruption, but are able to recover.

Unable to obtain DHCP lease on boot with RedHat

In configurations where the auto-negotiation process takes more than 5 seconds, the boot script may fail with the following message:

"ethX: failed. No link present. Check cable?"

This error may occur even though the presence of link can be confirmed using ethtool ethx. In this case, try setting "LINKDELAY=30" in /etc/sysconfig/network-scripts/ifdfg-ethx.

The same issue can occur during a network boot (via PXE) on RedHat distributions that use the dracut script:

"Warning: No carrier detected on interface <interface_name>"

In this case add "rd.net.timeout.carrier=30" at the kernel command line.



NOTE: Link time can vary. Adjust LINKDELAY value accordingly.

Alternatively, NetworkManager can be used to configure the interfaces, which avoids the set timeout. For configuration instructions of NetworkManager refer to the documentation provided by your distribution.

Loading i40e driver in 3.2.x and newer kernels displays kernel tainted message

Due to recent kernel changes, loading an out of tree driver causes the kernel to be tainted.

SR-IOV virtual functions have identical MAC addresses in RHEL8

When you create multiple SR-IOV virtual functions on Red Hat Enterprise Linux 8, the VFs may have identical MAC addresses. Only one VF will pass traffic, and all traffic on other VFs with identical MAC addresses will fail. This is related to the "MACAddressPolicy=persistent" setting in `/usr/lib/systemd/network/99-default.link`.

To resolve this issue, edit the `/usr/lib/systemd/network/99-default.link` file and change the `MACAddressPolicy` line to "`MACAddressPolicy=none`". For more information, see the `systemd` documentation.

'VF X failed opcode 24' error message in dmesg on host

With a Microsoft Windows Server 2019 guest machine running on a Linux host, you may see 'VF <vf_number> failed opcode 24' error messages in `dmesg` on the host. This error is benign and does not affect traffic. Installing the latest `iavf` driver in the guest will resolve the issue.

Windows guest OSs on a Linux host may not pass traffic across VLANs

The VF is not aware of the VLAN configuration if you use Load Balancing and Failover (LBFO) to configure VLANs in a Windows guest. VLANs configured using LBFO on a VF driver may result in failure to pass traffic.

ice Linux Driver for the Intel Ethernet Controller 800 Series

ice Overview

	NOTE: Devices based on the Intel® Ethernet Controller 800 Series may exhibit poor receive performance and dropped packets. The following steps may improve the situation: 1. In your system's BIOS/UEFI settings, select the "Performance" profile. 2. On RHEL 7.x/8.x, use the tuned power management tool to set the "latency-performance" profile. 3. In other operating systems and environments, use the equivalent tool to set the equivalent profile.
	NOTE: Do not unload a port's driver if a Virtual Function (VF) with an active Virtual Machine (VM) is bound to it. Doing so will cause the port to appear to hang. Once the VM shuts down, or otherwise releases the VF, the command will complete.
	NOTE: In a virtualized environment, on Intel® Server Adapters that support SR-IOV, the virtual function (VF) may be subject to malicious behavior. Software-generated layer two frames, like IEEE 802.3x (link flow control), IEEE 802.1Qbb (priority based flow-control), and others of this type, are not expected and can throttle traffic between the host and the virtual switch, reducing performance. To resolve this issue, and to ensure isolation from unintended traffic streams, configure all SR-IOV enabled ports for VLAN tagging from the administrative interface on the PF. This configuration allows unexpected, and potentially malicious, frames to be dropped.

The ice Linux* Base Driver for the Intel Ethernet Controller 800 Series family of adapters supports the 3.10.0 and newer kernels.

Driver information can be obtained using `ethtool`, `lspci`, and `ifconfig`. Instructions on updating `ethtool` can be found in the section [Additional Configurations](#) later in this document. This driver is only supported as a loadable module at this time. Intel is not supplying patches against the kernel source to allow for static linking of the drivers.

ice Linux Base Driver Supported Devices

- Intel® Ethernet 25G 2P E810-XXV OCP
- Intel® Ethernet 25G 2P E810-XXV Adapter

Building and Installation

There are three methods for installing the Linux driver:

- [Install from Source Code](#)
- [Install Using KMP RPM](#)
- [Install Using KMOD RPM](#)

Install from Source Code

To build a binary RPM* package of this driver, run `rpmbuild -tb <filename.tar.gz>`. Replace `<filename.tar.gz>` with the specific filename of the driver.



NOTES:

- For the build to work properly it is important that the currently running kernel MATCH the version and configuration of the installed kernel source. If you have just recompiled your kernel, reboot the system.
 - RPM functionality has only been tested in Red Hat distributions.
1. Download the base driver tar file to the directory of your choice. For example, use `"/home/username/ice"` or `"/usr/local/src/ice"`.
 2. Untar/unzip the archive, where `<x.x.x>` is the version number for the driver tar:


```
# tar xzf ice-<x.x.x>.tar.gz
```
 3. Change to the driver src directory, where `<x.x.x>` is the version number for the driver tar:


```
# cd ice-<x.x.x>/src/
```

4. Compile the driver module:

```
# make install
```

The binary will be installed as:

```
/lib/modules/<KERNEL_VERSION>/kernel/drivers/net/ice/ice.ko
```

The install locations listed above are the default locations. This might differ for various Linux distributions. For more information, see the `ldistrib.txt` file included in the driver tar.

5. Remove the old driver:

```
# rmmod ice
```

6. Install the module using the `modprobe` command:

```
# modprobe ice <parameter>=<value>
```

7. Update the system image:

```
dracut -f
```

8. Assign an IP address to and activate the Ethernet interface by entering the following, where `<ethx>` is the interface name:

```
# ifconfig <ethX> <IP_address> netmask <netmask> up
```

9. Verify that the interface works. Enter the following, where `<IP_address>` is the IP address for another machine on the same subnet as the interface that is being tested:

```
# ping <IP_address>
```

Install Using KMP RPM

The KMP RPMs update existing ice RPMs currently installed on the system. These updates are provided by SuSE in the SLES release. If an RPM does not currently exist on the system, the KMP will not install.

The RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
intel-<component name>-<component version>.<arch type>.rpm
```

For example, `intel-ice-1.3.8.6-1.x86_64.rpm`:

- ice is the component name
- 1.3.8.6-1 is the component version
- x86_64 is the architecture type

KMP RPMs are provided for supported Linux distributions. The naming convention for the included KMP RPMs is:

```
intel-<component name>-kmp-<kernel type>-<component version>_<kernel version>.<arch type>.rpm
```

For example, `intel-ice-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm`:

- ice is the component name
- default is the kernel type
- 1.3.8.6 is the component version
- 2.6.27.19_5-1 is the kernel version
- x86_64 is the architecture type

To install the KMP RPM, type the following two commands:

```
# rpm -i <rpm filename>
# rpm -i <kmp rpm filename>
```

For example, to install the ice KMP RPM package, type the following:

```
# rpm -i intel-ice-1.3.8.6-1.x86_64.rpm
# rpm -i intel-ice-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm
```

Install Using KMOD RPM

The KMOD RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
kmod-<driver name>-<version>-1.<arch type>.rpm
```

For example, kmod-ice-2.3.4-1.x86_64.rpm:

- ice is the driver name
- 2.3.4 is the version
- x86_64 is the architecture type

To install the KMOD RPM, go to the directory of the RPM and type the following command:

```
# rpm -i <rpm filename>
```

For example, to install the ice KMOD RPM package, type the following:

```
# rpm -i kmod-ice-2.3.4-1.x86_64.rpm
```

Command Line Parameters

The only command line parameter the ice driver supports is the debug parameter that can control the default logging verbosity of the driver. (Note: dyndbg also provides dynamic debug information.)

In general, use ethtool and other OS-specific commands to configure user-changeable parameters after the driver is loaded.

Additional Configurations

ethtool

The driver utilizes the ethtool interface for driver configuration and diagnostics, as well as displaying statistical information. The latest ethtool version is required for this functionality. Download it at: <https://kernel.org/pub/software/network/ethtool/>

The rx_bytes value of ethtool does not match the rx_bytes value of Netdev, due to the 4-byte CRC being stripped by the device. The difference between the two rx_bytes values will be 4 x the number of Rx packets. For example, if Rx packets are 10 and Netdev (software statistics) displays rx_bytes as "X", then ethtool (hardware statistics) will display rx_bytes as "X+40" (4 bytes CRC x 10 packets).

Viewing Link Messages

Link messages will not be displayed to the console if the distribution is restricting system messages. In order to see network driver link messages on your console, set dmesg to eight by entering the following:

```
# dmesg -n 8
```



NOTE: This setting is not saved across reboots.

Dynamic Device Personalization

Dynamic Device Personalization (DDP) allows you to change the packet processing pipeline of a device by applying a profile package to the device at runtime. Profiles can be used to, for example, add support for new protocols, change existing protocols, or change default settings. DDP profiles can also be rolled back without rebooting the system.

The driver automatically installs the default DDP package file during driver installation.



NOTE: It's important to do 'make install' during initial driver installation so that the driver loads the DDP package automatically.

The DDP package loads during device initialization. The driver looks for `intel/ice/ddp/ice.pkg` in your firmware root (typically `/lib/firmware/` or `/lib/firmware/updates/`) and checks that it contains a valid DDP package file.

If the driver is unable to load the DDP package, the device will enter Safe Mode. Safe Mode disables advanced and performance features and supports only basic traffic and minimal functionality, such as updating the NVM or downloading a new driver or DDP package. Safe Mode only applies to the affected physical function and does not impact any other PFs. For more details, see "Dynamic Device Personalization (DDP)" on page 13 and "Safe Mode" on page 178 in this user guide.

**NOTE:**

- If you encounter issues with the DDP package file, you may need to download an updated driver or DDP package file. See the log messages for more information.
- The `ice.pkg` file is a symbolic link to the default DDP package file installed by the Linux-firmware software package or the ice out-of-tree driver installation.
- You cannot update the DDP package if any PF drivers are already loaded. To overwrite a package, unload all PFs and then reload the driver with the new package.
- Only the first loaded PF per device can download a package for that device.
- You can install specific DDP package files for different physical devices in the same system. See the Linux driver README in your installation for instructions on how to install a specific DDP package file.

RDMA (Remote Direct Memory Access)

Remote Direct Memory Access, or RDMA, allows a network device to transfer data directly to and from application memory on another system, increasing throughput and lowering latency in certain networking environments.

The ice driver supports the following RDMA protocols:

- iWARP (Internet Wide Area RDMA Protocol)
- RoCEv2 (RDMA over Converged Ethernet)

The major difference is that iWARP performs RDMA over TCP, while RoCEv2 uses UDP.

For detailed installation and configuration information, see the README file in the RDMA driver tarball.

**NOTES:**

- Devices based on the Intel® Ethernet Controller 800 Series do not support RDMA when operating in multiport mode with more than 4 ports.
- You cannot use RDMA or SR-IOV when link aggregation (LAG)/bonding is active, and vice versa. To enforce this, on kernels 4.5 and above, the ice driver checks for this mutual exclusion. On kernels older than 4.5, the ice driver cannot check for this exclusion and is unaware of bonding events.

Application Device Queues (ADQ)

Application Device Queues (ADQ) allow you to dedicate one or more queues to a specific application. This can reduce latency for the specified application, and allow Tx traffic to be rate limited per application.

For requirements and configuration information, refer to the [Intel® Ethernet Controller E810 Application Device Queues \(ADQ\) Configuration Guide \(https://cdrdv2.intel.com/v1/dl/getContent/609008\)](https://cdrdv2.intel.com/v1/dl/getContent/609008)

Intel® Ethernet Flow Director

The Intel Ethernet Flow Director performs the following tasks:

- Directs receive packets according to their flows to different queues
- Enables tight control on routing a flow in the platform
- Matches flows and CPU cores for flow affinity

An included script (`set_irq_affinity`) automates setting the IRQ to CPU affinity.

This driver supports the following flow types:

- IPv4
- TCPv4
- UDPv4
- IPv6
- TCPv6
- UDPv6

For a given flow type, it supports valid combinations of IP addresses (source or destination) and UDP/TCP ports (source and destination). For example, you can supply only a source IP address, a source IP address and a destination port, or any combination of one or more of these four parameters.

This driver allows you to filter traffic based on a user-defined flexible two-byte pattern and offset by using the ethtool user-def and mask fields. Only L3 and L4 flow types are supported for user-defined flexible filters. For a given flow type, you must clear all Intel Ethernet Flow Director filters before changing the input set (for that flow type).



NOTE: Flow Director filters impact only LAN traffic. RDMA filtering occurs before Flow Director, so Flow Director filters will not impact RDMA.

The following table summarizes supported Intel Ethernet Flow Director features across Intel® Ethernet controllers.

Feature	500 Series	700 Series	800 Series
VF Flow Director	Supported	Routing to VF not supported	Not supported
IP Address Range Filter	Supported	Not supported	Not supported
IPv6 Support	Not supported	Not supported	Supported
Configurable Input Set	Configured per port	Configured globally	Configured per port
ATR	Supported	Supported	Not supported
Flex Byte Filter	Starts at beginning of packet	Starts at beginning of payload	Starts at beginning of packet
Tunneled Packets	Filter matches outer header	Filter matches Inner header	Filter matches outer header

To enable or disable the Intel Ethernet Flow Director:

```
# ethtool -K <ethX> ntuple <on|off>
```

When disabling ntuple filters, all the user programmed filters are flushed from the driver cache and hardware. All needed filters must be re-added when ntuple is re-enabled.

Flow Director Filters

Flow Director filters are used to direct traffic that matches specified characteristics. They are enabled through ethtool's ntuple interface. To enable or disable these filters:

```
# ethtool -K <ethX> ntuple <off|on>
```

To display all of the active filters:

```
# ethtool -u <ethX>
```

To add a new filter:

```
# ethtool -U <ethX> flow-type <type> src-ip <ip> dst-ip <ip> src-port <port> dst-port <port> action <queue>
```

Where:

- <ethX> - the Ethernet device to program
- <type> - can be ip4, tcp4, udp4, sctp4, ip6, tcp6, udp6, or sctp6
- <ip> - the ip address to match on
- <port> - the port number to match on
- <queue> - the queue to direct traffic towards (-1 discards the matched traffic)

To delete a filter:

```
# ethtool -U <ethX> delete <N>
```

Where <N> is the filter ID displayed when printing all the active filters, and may also have been specified using "loc <N>" when adding the filter.

Examples:

To add a filter that directs packet to queue 2:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 src-port
2000 dst-port 2001 action 2 [loc 1]
```

To set a filter using only the source and destination IP address:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 action 2
[loc 1]
```

To set a filter based on a user-defined pattern and offset:

```
# ethtool -N <ethX> flow-type tcp4 src-ip 192.168.10.1 dst-ip 192.168.10.2 user-def
0x4FFFF action 2 [loc 1]
```

where the value of the user-def field contains the offset (4 bytes) and the pattern (0xffff).

To match TCP traffic sent from 192.168.0.1, port 5300, directed to 192.168.0.5, port 80, and then send it to queue 7:

```
# ethtool -U enp130s0 flow-type tcp4 src-ip 192.168.0.1 dst-ip 192.168.0.5 src-port
5300 dst-port 80 action 7
```

For each flow-type, the programmed filters must all have the same matching input set. For example, issuing the following two commands is acceptable:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 src-port 5300 action 7
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.5 src-port 55 action 10
```

Issuing the next two commands, however, is not acceptable, since the first specifies `src-ip` and the second specifies `dst-ip`:

```
# ethtool -U enp130s0 flow-type ip4 src-ip 192.168.0.1 src-port 5300 action 7
# ethtool -U enp130s0 flow-type ip4 dst-ip 192.168.0.5 src-port 55 action 10
```

The second command will fail with an error. You may program multiple filters with the same fields, using different values, but, on one device, you may not program two tcp4 filters with different matching fields.

Matching on a subportion of a field is not supported by the driver, thus partial mask fields are not supported.

Flex Byte Flow Director Filters

The driver also supports matching user-defined data within the packet payload. This flexible data is specified using the "user-def" field of the ethtool command in the following way:

31	28	24	20	16	15	12	8	4	0
offset into packet payload					2 bytes of flexible data				

For example,

```
... user-def 0x4FFFF ...
```

tells the filter to look 4 bytes into the payload and match that value against 0xFFFF. The offset is based on the beginning of the payload, and not the beginning of the packet. Thus

```
flow-type tcp4 ... user-def 0x8BEAF ...
```

would match TCP/IPv4 packets which have the value 0xBEAF 8 bytes into the TCP/IPv4 payload.

Note that ICMP headers are parsed as 4 bytes of header and 4 bytes of payload. Thus to match the first byte of the payload, you must actually add 4 bytes to the offset. Also note that ip4 filters match both ICMP frames as well as raw (unknown) ip4 frames, where the payload will be the L3 payload of the IP4 frame.

The maximum offset is 64. The hardware will only read up to 64 bytes of data from the payload. The offset must be even because the flexible data is 2 bytes long and must be aligned to byte 0 of the packet payload.

The user-defined flexible offset is also considered part of the input set and cannot be programmed separately for multiple filters of the same type. However, the flexible data is not part of the input set and multiple filters may use the same offset but match against different data.

RSS Hash Flow

Allows you to set the hash bytes per flow type and any combination of one or more options for Receive Side Scaling (RSS) hash byte configuration.

```
# ethtool -N <ethX> rx-flow-hash <type> <option>
```

Where <type> is:

```
tcp4 signifying TCP over IPv4
udp4 signifying UDP over IPv4
tcp6 signifying TCP over IPv6
udp6 signifying UDP over IPv6
```

And <option> is one or more of:

```
s Hash on the IP source address of the Rx packet.
d Hash on the IP destination address of the Rx packet.
f Hash on bytes 0 and 1 of the Layer 4 header of the Rx packet.
n Hash on bytes 2 and 3 of the Layer 4 header of the Rx packet.
```

Accelerated Receive Flow Steering (aRFS)

Devices based on the Intel® Ethernet Controller 800 Series support Accelerated Receive Flow Steering (aRFS) on the PF. aRFS is a load-balancing mechanism that allows you to direct packets to the same CPU where an application is running or consuming the packets in that flow.



NOTES:

- aRFS requires that ntuple filtering is enabled via ethtool.
- aRFS support is limited to the following packet types:
 - TCP over IPv4 and IPv6
 - UDP over IPv4 and IPv6
 - Nonfragmented packets
- aRFS only supports Flow Director filters, which consist of the source/destination IP addresses and source/destination ports.
- aRFS and ethtool's ntuple interface both use the device's Flow Director. aRFS and ntuple features can coexist,

but you may encounter unexpected results if there's a conflict between aRFS and ntuple requests.

To set up aRFS:

1. Enable the Intel Ethernet Flow Director and ntuple filters using ethtool.

```
# ethtool -K <ethX> ntuple on
```

2. Set up the number of entries in the global flow table. For example:

```
# NUM_RPS_ENTRIES=16384
# echo $NUM_RPS_ENTRIES > /proc/sys/net/core/rps_sock_flow_entries
```

3. Set up the number of entries in the per-queue flow table. For example:

```
# NUM_RX_QUEUES=64
# for file in /sys/class/net/$IFACE/queues/rx-*/rps_flow_cnt; do
# echo $((($NUM_RPS_ENTRIES/$NUM_RX_QUEUES)) > $file;
# done
```

4. Disable the IRQ balance daemon (this is only a temporary stop of the service until the next reboot).

```
# systemctl stop irqbalance
```

5. Configure the interrupt affinity.

```
# set_irq_affinity <ethX>
```

To disable aRFS using ethtool:

```
# ethtool -K <ethX> ntuple off
```

NOTE: This command will disable ntuple filters and clear any aRFS filters in software and hardware.

Example Use Case:

1. Set the server application on the desired CPU (e.g., CPU 4).

```
# taskset -c 4 netserver
```

2. Use netperf to route traffic from the client to CPU 4 on the server with aRFS configured. This example uses TCP over IPv4.

```
# netperf -H <Host IPv4 Address> -t TCP_STREAM
```

Enabling Virtual Functions (VFs)

Use sysfs to enable virtual functions (VF).

For example, you can create 4 VFs as follows:

```
# echo 4 > /sys/class/net/<ethX>/device/sriov_numvfs
```

To disable VFs, write 0 to the same file:

```
# echo 0 > /sys/class/net/<ethX>/device/sriov_numvfs
```

The maximum number of VFs for the ice driver is 256 total (all ports). To check how many VFs each PF supports, use the following command:

```
# cat /sys/class/net/<ethX>/device/sriov_totalvfs
```



NOTE: You cannot use RDMA or SR-IOV when link aggregation (LAG)/bonding is active, and vice versa. To enforce this, on kernels 4.5 and above, the ice driver checks for this mutual exclusion. On kernels older than 4.5, the ice driver cannot check for this exclusion and is unaware of bonding events.

Configuring VLAN Tagging on SR-IOV Enabled Adapter Ports

To configure VLAN tagging for the ports on an SR-IOV enabled adapter, use the following command. The VLAN configuration should be done before the VF driver is loaded or the VM is booted. The VF is not aware of the VLAN tag being inserted on transmit and removed on received frames (sometimes called "port VLAN" mode).

```
# ip link set dev <PF netdev id> vf <id> vlan <vlan id>
```

For example, the following will configure PF eth0 and the first VF on VLAN 10:

```
# ip link set dev eth0 vf 0 vlan 10
```

Setting the MAC Address for a VF

To change the MAC address for the specified VF:

```
# ip link set <ethX> vf 0 mac <address>
```

For example:

```
# ip link set <ethX> vf 0 mac 00:01:02:03:04:05
```

This setting lasts until the PF is reloaded.



NOTE: Assigning a MAC address for a VF from the host will disable any subsequent requests to change the MAC address from within the VM. This is a security feature. The VM is not aware of this restriction, so if this is attempted in the VM, it will trigger MDD events.

Trusted VFs and VF Promiscuous Mode

This feature allows you to designate a particular VF as trusted and allows that trusted VF to request selective promiscuous mode on the Physical Function (PF).

To set a VF as trusted or untrusted, enter the following command in the Hypervisor:

```
# ip link set dev eth0 vf 1 trust [on|off]
```



NOTE: It's important to set the VF to trusted before setting promiscuous mode. If the VM is not trusted, the PF will ignore promiscuous mode requests from the VF. If the VM becomes trusted after the VF driver is loaded, you must make a new request to set the VF to promiscuous.

Once the VF is designated as trusted, use the following commands in the VM to set the VF to promiscuous mode.

- For promiscuous all: # ip link set eth2 promisc on
Where eth2 is a VF interface in the VM
- For promiscuous Multicast: # ip link set eth2 allmulticast on
Where eth2 is a VF interface in the VM



NOTE: By default, the ethtool private flag vf-true-promisc-support is set to "off," meaning that promiscuous mode for the VF will be limited. To set the promiscuous mode for the VF to true promiscuous and allow the VF to see all ingress traffic, use the following command:

```
# ethtool --set-priv-flags p261p1 vf-true-promisc-support on
```

The vf-true-promisc-support private flag does not enable promiscuous mode; rather, it designates which type of promiscuous mode (limited or true) you will get when you enable promiscuous mode using the ip link commands above. Note that this is a global setting that affects the entire device. However, the vf-true-promisc-support private flag is only exposed to the first PF of the device. The PF remains in limited promiscuous mode regardless of the vf-true-promisc-support setting.

Next, add a VLAN interface on the VF interface.

```
# ip link add link eth2 name eth2.100 type vlan id 100
```

Note that the order in which you set the VF to promiscuous mode and add the VLAN interface does not matter (you can do either first). The result in this example is that the VF will get all traffic that is tagged with VLAN 100.

Virtual Function (VF) Tx Rate Limit

Use the `ip` command to configure the maximum or minimum Tx rate limit for a VF from the PF interface.

For example, to set a maximum Tx rate limit of 8000Mbps for VF 0:

```
# ip link set eth0 vf 0 max_tx_rate 8000
```

For example, to set a minimum Tx rate limit of 1000Mbps for VF 0:

```
# ip link set eth0 vf 0 min_tx_rate 1000
```

Malicious Driver Detection (MDD) for VFs

Some Intel Ethernet devices use Malicious Driver Detection (MDD) to detect malicious traffic from the VF and disable Tx/Rx queues or drop the offending packet until a VF driver reset occurs. You can view MDD messages in the PF's system log using the `dmesg` command.

- If the PF driver logs MDD events from the VF, confirm that the correct VF driver is installed.
- To restore functionality, you can manually reload the VF or VM or enable automatic VF resets.
- When automatic VF resets are enabled, the PF driver will immediately reset the VF and reenables queues when it detects MDD events on the receive path.
- If automatic VF resets are disabled, the PF will not automatically reset the VF when it detects MDD events.

To enable or disable automatic VF resets, use the following command:

```
# ethtool --set-priv-flags <ethX> mdd-auto-reset-vf on|off
```

MAC and VLAN Anti-Spoofing Feature for VFs

When a malicious driver on a Virtual Function (VF) interface attempts to send a spoofed packet, it is dropped by the hardware and not transmitted.

When a spoofed packet is detected the PF driver will send the following message to the system log (displayed by the "`dmesg`" command):

This feature can be disabled for a specific VF:

```
# ip link set <pf dev> vf <vf id> spoofchk {off|on}
```

Jumbo Frames

Jumbo Frames support is enabled by changing the MTU to a value larger than the default of 1500 bytes. Use the `ifconfig` command to increase the MTU size. For example, enter the following where `<ethX>` is the interface number::

```
# ifconfig <ethX> mtu 9000 up
```

Alternatively, you can use the `ip` command as follows:

```
# ip link set mtu 9000 dev <ethX>
```

```
# ip link set up dev <ethX>
```

This setting is not saved across reboots. The setting change can be made permanent by adding 'MTU = 9000' to the following file:

- `/etc/sysconfig/network-scripts/ifcfg-<ethX>` for RHEL
- `/etc/sysconfig/network/<config_file>` for SLES

 NOTES:

- The maximum MTU setting for Jumbo Frames is 9702 bytes. This value coincides with the maximum Jumbo Frames size of 9728 bytes.
- This driver will attempt to use multiple page sized buffers to receive each jumbo packet. This should help to avoid buffer starvation issues when allocating receive packets.
- Packet loss may have a greater impact on throughput when you use jumbo frames. If you observe a drop in performance after enabling jumbo frames, enabling flow control may mitigate the issue.

Speed and Duplex Configuration

In addressing speed and duplex configuration issues, you need to distinguish between copper-based adapters and fiber-based adapters.

In the default mode, an Intel® Ethernet Network Adapter using copper connections will attempt to auto-negotiate with its link partner to determine the best setting. If the adapter cannot establish link with the link partner using auto-negotiation, you may need to manually configure the adapter and link partner to identical settings to establish link and pass packets. This should only be needed when attempting to link with an older switch that does not support auto-negotiation or one that has been forced to a specific speed or duplex mode. Your link partner must match the setting you choose. 1 Gbps speeds and higher cannot be forced. Use the autonegotiation advertising setting to manually set devices for 1 Gbps and higher.

Speed, duplex, and autonegotiation advertising are configured through the ethtool utility. ethtool is included with all versions of Red Hat after Red Hat 7.2. To see the speed configurations your device supports, run the following:

```
# ethtool <ethX>
```



CAUTION: Only experienced network administrators should force speed and duplex or change autonegotiation advertising manually. The settings at the switch must always match the adapter settings. Adapter performance may suffer or your adapter may not operate if you configure the adapter differently from your switch.

Data Center Bridging (DCB)

Note: The kernel assumes that TC0 is available, and will disable Priority Flow Control (PFC) on the device if TC0 is not available. To fix this, ensure TC0 is enabled when setting up DCB on your switch.

DCB is a configuration Quality of Service implementation in hardware. It uses the VLAN priority tag (802.1p) to filter traffic. That means that there are 8 different priorities that traffic can be filtered into. It also enables priority flow control (802.1Qbb) which can limit or eliminate the number of dropped packets during network stress. Bandwidth can be allocated to each of these priorities, which is enforced at the hardware level (802.1Qaz).

Adapter firmware implements LLDP and DCBX protocol agents as per 802.1AB and 802.1Qaz respectively. There are potentially two DCBX modes on Linux, depending on the underlying PF device:

- Intel Ethernet Controller 800 Series adapters support both firmware DCBX and software DCBX. If FW-LLDP is enabled, DCBX will run in firmware.

DCB parameters can be established via a firmware LLDP/DCBX agent or a software LLDP/DCBX agent. Only one LLDP/DCBX agent can be active on a single interface at a time. When the firmware DCBX agent is active, software agents will not be able to receive or transmit LLDP frames.

See the "FW-LLDP (Firmware Link Layer Discovery Protocol)" subsection for this driver for the ethtool commands to query the status of the firmware LLDP/DCBX agent.

When operating in firmware DCBX mode, the adapter is in an "always willing" state. DCB settings are applied on the adapter by transmitting a nonwilling configuration from the link partner. Typically this is a switch. For configuring DCBX parameters on a switch, please consult the switch manufacturer's documentation.

**NOTE:**

Intel Ethernet Controller 800 Series adapters:

- This driver supports DCB when the firmware agent is on or off by supporting software DCBX agents.
- When the firmware LLDP agent is disabled, you can configure DCB parameters using software LLDP/DCBX

agents that interface with the Linux kernel's DCB Netlink API. We recommend using OpenLLDP as the DCBX agent when running in software mode. For more information, see the OpenLLDP man pages and <https://github.com/intel/openlldp>.

- iSCSI with DCB is not supported.

FW-LLDP (Firmware Link Layer Discovery Protocol)

Use ethtool to change FW-LLDP settings. The FW-LLDP setting is per port and persists across boots.

To enable LLDP:

```
# ethtool --set-priv-flags <ethX> fw-lldp-agent on
```

To disable LLDP:

```
# ethtool --set-priv-flags <ethX> fw-lldp-agent off
```

To check the current LLDP setting:

```
# ethtool --show-priv-flags <ethX>
```



NOTE: You must enable the UEFI HII "LLDP Agent" attribute for this setting to take effect. If "LLDP AGENT" is set to disabled, you cannot enable it from the OS.

Flow Control

Ethernet Flow Control (IEEE 802.3x) can be configured with ethtool to enable receiving and transmitting pause frames for this driver. When transmit is enabled, pause frames are generated when the receive packet buffer crosses a predefined threshold. When receive is enabled, the transmit unit will halt for the time delay specified when a pause frame is received.



NOTES:

- You must have a flow control capable link partner.
- This driver requires flow control on both the port and link partner. If flow control is disabled on one of the sides, the port may appear to hang on heavy traffic.

Use ethtool to change the flow control settings.

To enable or disable Rx or Tx Flow Control:

```
# ethtool -A <ethX> rx <on|off> tx <on|off>
```



NOTE: This command only enables or disables Flow Control if auto-negotiation is disabled. If auto-negotiation is enabled, this command changes the parameters used for auto-negotiation with the link partner.

To enable or disable auto-negotiation:

```
# ethtool -s <ethX> autoneg <on|off>
```



NOTE: Flow Control auto-negotiation is part of link auto-negotiation. Depending on your device, you may not be able to change the auto-negotiation setting.

You may encounter issues with link-level flow control (LFC) after disabling DCB. The LFC status may show as enabled but traffic is not paused. To resolve this issue, disable and reenable LFC using ethtool:

```
# ethtool -A <ethX> rx off tx off
```

```
# ethtool -A <ethX> rx on tx on
```

NAPI

This driver supports NAPI (Rx polling mode). For more information on NAPI, see <https://wiki.linux-foundation.org/networking/napi>.

MACVLAN

This driver supports MACVLAN. Kernel support for MACVLAN can be tested by checking if the MACVLAN driver is loaded. You can run 'lsmod | grep macvlan' to see if the MACVLAN driver is loaded or run 'modprobe macvlan' to try to load the MACVLAN driver.



NOTE:

- In passthru mode, you can only set up one MACVLAN device. It will inherit the MAC address of the underlying PF (Physical Function) device.
- You cannot enable MACVLAN offloading and ADQ at the same time.

ice devices support L2 Forwarding Offload. This will offload the processing required for L2 Forwarding from the system processors to the ice device.

Perform the following steps to enable L2 Forwarding Offload:

1. Enable L2 Forwarding offload:
`ethtool -K <ethX> l2-fwd-offload on`
2. Create the MACVLAN netdevs and bind them to the PF.
3. Bring up/enable the MACVLAN netdevs.

IEEE 802.1ad (QinQ) Support

The IEEE 802.1ad standard, informally known as QinQ, allows for multiple VLAN IDs within a single Ethernet frame. VLAN IDs are sometimes referred to as "tags," and multiple VLAN IDs are thus referred to as a "tag stack." Tag stacks allow L2 tunneling and the ability to segregate traffic within a particular VLAN ID, among other uses.



NOTES:

- 802.1ad (QinQ) is supported in 3.19 and later kernels.
- 802.1ad (QinQ) and RDMA are not compatible.
- Receive checksum offloads and VLAN acceleration are not supported for 802.1ad (QinQ) packets.
- 0x88A8 traffic will not be received unless VLAN stripping is disabled with the following command:
`ethtool -K <ethX> rxvlan off`
- The VF can only transmit 0x88A8/0x8100 (i.e., 802.1ad/802.1Q) traffic if:
 1. VF is not assigned a port VLAN.
 2. spoofchk is disabled from the PF. If you enable spoofchk, the VF will not transmit 0x88A8/0x8100 traffic.

The following are examples of how to configure 802.1ad (QinQ):

```
# ip link add link eth0 eth0.24 type vlan proto 802.1ad id 24
# ip link add link eth0.24 eth0.24.371 type vlan proto 802.1Q id 371
```

Where "24" and "371" are example VLAN IDs.

IEEE 1588 Precision Time Protocol (PTP) Hardware Clock (PHC)

Precision Time Protocol (PTP) is used to synchronize clocks in a computer network. PTP support varies among Intel devices that support this driver. Use 'ethtool -T <ethX>' to get a definitive list of PTP capabilities supported by the device.

Tunnel/Overlay Stateless Offloads

Supported tunnels and overlays include VXLAN, GENEVE, and others depending on hardware and software configuration. Stateless offloads are enabled by default.

To view the current state of all offloads:

```
# ethtool -k <ethX>
```

UDP Segmentation Offload

Allows the adapter to offload transmit segmentation of UDP packets with payloads up to 64K into valid Ethernet frames. Because the adapter hardware is able to complete data segmentation much faster than operating system software, this feature may improve transmission performance.

In addition, the adapter may use fewer CPU resources.



NOTES:

- UDP transmit segmentation offload requires Linux kernel 4.18 or later.
- The application sending UDP packets must support UDP segmentation offload.

To enable/disable UDP Segmentation Offload, issue the following command:

```
# ethtool -K <ethX> tx-udp-segmentation [off|on]
```

Performance Optimization

The driver defaults are meant to fit a wide variety of workloads. If further optimization is required, we recommend experimenting with the following settings.

IRQ to Adapter Queue Alignment

Pin the adapter's IRQs to specific cores by disabling the irqbalance service and using the included `set_irq_affinity` script. Please see the script's help text for further options.

The following settings will distribute the IRQs across all the cores evenly:

```
# scripts/set_irq_affinity -X all <interface1> , [ <interface2>, ... ]
```

The following settings will distribute the IRQs across all the cores that are local to the adapter (same NUMA node):

```
# scripts/set_irq_affinity -X local <interface1> , [ <interface2>, ... ]
```

For very CPU-intensive workloads, we recommend pinning the IRQs to all cores.

Rx Descriptor Ring Size

To reduce the number of Rx packet discards, increase the number of Rx descriptors for each Rx ring using `ethtool`.

Check if the interface is dropping Rx packets due to buffers being full (`rx_dropped.nic` means there is no PCIe bandwidth):

```
# ethtool -S <interface> | grep "rx_dropped"
```

Increase the number of Rx descriptors for each Rx ring using `ethtool`. This may help reduce Rx packet drops at the expense of system resources:

```
# ethtool -G <interface> rx N
```

Where N is the desired number of rings

Interrupt Rate Limiting

This driver supports an adaptive interrupt rate mechanism that is tuned for general workloads. The user can customize the interrupt rate control for specific workloads, via ethtool, adjusting the number of microseconds between interrupts.

To set the interrupt rate manually, you must disable adaptive mode:

```
# ethtool -C <interface> adaptive-rx off adaptive-tx off
```

For lower CPU utilization:

1. Disable adaptive ITR and lower Rx and Tx interrupts. The examples below affect every queue of the specified interface.
2. Setting rx-usecs and tx-usecs to 125 will limit interrupts to about 8000 interrupts per second per queue:

```
# ethtool -C <interface> adaptive-rx off adaptive-tx off rx-usecs 125 tx-usecs 125
```

For reduced latency:

Disable adaptive ITR and ITR by setting rx-usecs and tx-usecs to 0 using ethtool:

```
# ethtool -C <interface> adaptive-rx off adaptive-tx off rx-usecs 0 tx-usecs 0
```

Virtualized Environments

In addition to the other suggestions in this section, the following may be helpful to optimize performance in VMs.

Using the appropriate mechanism (vcpupin) in the VM, pin the CPUs to individual LCPUs, making sure to use a set of CPUs included in the device's local_cpulist: /sys/class/net/<ethX>/device/local_cpulist.

Configure as many Rx/Tx queues in the VM as available. (See the iavf driver documentation for the number of queues supported.) For example:

```
# ethtool -L <virt_interface> rx <max> tx <max>
```

Known Issues

Dynamic Debug

If you encounter unexpected issues during driver load, some of the most useful information for developers to receive in a bug report can include driver logging. This logging uses a kernel feature called Dynamic Debug, which is generally enabled in most kernel configurations (CONFIG_DYNAMIC_DEBUG=y).

To load the driver with dynamic debug enabled, run modprobe with the dyndbg parameter:

```
# modprobe ice dyndbg=+p
```

The driver will then load and print debugging information into the kernel log (dmesg) and is usually logged into the system log viewable by journalctl or in /var/log/messages. Saving this information to a file and attaching it to any bug report can help shorten the reproduction and debugging time for a developer.

To enable dynamic debug during runtime operation of the driver, use this command:

```
# echo "module ice +p" > /sys/kernel/debug/dynamic_debug/control
```

For more details, see the Dynamic Debug documentation included in the Linux kernel instructions.

'ethtool -S' does not display Tx/Rx packet statistics

Issuing the command 'ethtool -S' does not display Tx/Rx packet statistics. This is by convention. Use other tools (e.g. ifconfig, ip) that display standard netdev statistics such as Tx/Rx packet statistics.

Fiber optics and auto-negotiation

Modules based on 100GBASE-SR4, active optical cable (AOC), and active copper cable (ACC) do not support auto-negotiation per the IEEE specification. To obtain link with these modules, you must turn off auto-negotiation on the link partner's switch ports.

'ethtool -a' autonegotiate result may vary between drivers

For kernel versions 4.6 or higher, 'ethtool -a' will show the advertised and negotiated autoneg settings. For kernel versions below 4.6, ethtool will only report the negotiated link status.

The issue is cosmetic and does not affect functionality. Installing the latest ice driver and upgrading your kernel to version 4.6 or higher will resolve the issue.

AF_XDP fails to allocate buffers

On kernels older than 5.3, you may see an undesirable CPU load during packet processing if you enable AF_XDP in native mode and the Rx ring size is larger than the UMEM fill queue. This is due to a known issue in the kernel and was fixed in 5.3. To address the issue, upgrade your kernel to 5.3 or newer.

SCTP checksum offloads aren't indicated on Geneve tunnel

For SCTP traffic over a Geneve tunnel, the SCTP checksum isn't offloaded to the device, even when tx-checksum-sctp is on. This is due to a limitation in the Linux kernel. However, for Rx traffic, the SCTP checksum is verified if rx-checksumming is on. For both Tx and Rx traffic, you can offload the outer UDP checksum to the device.

Incorrect link speed reported on older VF drivers

Linux distributions with older iavf or i40evf drivers (including Red Hat Enterprise Linux 8) may show an incorrect link speed on VF interfaces. This issue is cosmetic and does not affect VF functionality. To resolve the issue, download the latest iavf driver.

Older VF drivers on Intel Ethernet Controller 800 Series based adapters

Some Windows* VF drivers from Release 22.9 or older may encounter errors when loaded on a PF based on the Intel Ethernet Controller 800 Series on Linux KVM. You may see errors and the VF may not load. This issue does not occur starting with the following Windows VF drivers:

- v40e64, v40e65: Version 1.5.65.0 and newer

To resolve this issue, download and install the latest iavf driver.

'VF X failed opcode 24' error message in dmesg on host

With a Microsoft Windows Server 2019 guest machine running on a Linux host, you may see 'VF <vf_number> failed opcode 24' error messages in dmesg on the host. This error is benign and does not affect traffic. Installing the latest iavf driver in the guest will resolve the issue.

Windows guest OSs on a Linux host may not pass traffic across VLANs

The VF is not aware of the VLAN configuration if you use Load Balancing and Failover (LBFO) to configure VLANs in a Windows guest. VLANs configured using LBFO on a VF driver may result in failure to pass traffic.

MDD events in dmesg when creating maximum number of VLANs on the VF

When you create the maximum number of VLANs on the VF, you may see MDD events in dmesg on the host. This is due to the asynchronous design of the iavf driver. It always reports success to any VLAN requests, but the requests may fail later. The guest OS could try to send traffic on a VLAN that is not configured on the VF, which will cause a Malicious Driver Detection (MDD) event in dmesg on the host.

This issue is cosmetic. You do not need to reload the PF driver.

'ip address' or 'ip link' command displays an error on a single-port NIC with 245+ VFs

When you use the 'ip address' or 'ip link' command on a Linux host configured with 245 or more VFs on a single-port adapter, you may encounter a "Buffer too small for object" error. This is due to a known issue in the iproute2 tools. Please use ifconfig instead of iproute2. You can install ifconfig via the net-tools-deprecated package.

iavf Linux Driver

iavf Overview

The i40evf driver was renamed to the iavf (Intel Adaptive Virtual Function) driver. This was done to reduce the impact of future Intel Ethernet Controllers. The iavf driver allows you to upgrade your hardware without needing to upgrade the virtual function driver in each of the VMs running on top of the hardware.

iavf Linux Base Driver Supported Devices

The following Intel network adapters are compatible with this driver:

- Intel® Ethernet 10G 4P X710-k bNDC
- Intel® Ethernet 10G 2P X710-k bNDC
- Intel® Ethernet 10G X710-k bNDC
- Intel® Ethernet Converged Network Adapter X710
- Intel® Ethernet Converged Network Adapter X710-T
- Intel® Ethernet 10G 4P X710/I350 rNDC
- Intel® Ethernet 10G 4P X710 SFP+ rNDC
- Intel® Ethernet 10G X710 rNDC
- Intel® Ethernet Server Adapter X710-DA2 for OCP
- Intel® Ethernet 10G 2P X710 OCP
- Intel® Ethernet 10G 4P X710 OCP
- Intel® Ethernet 10G 2P X710-T2L-t OCP
- Intel® Ethernet 10G 4P X710-T4L-t OCP
- Intel® Ethernet 10G 2P X710-T2L-t Adapter
- Intel® Ethernet 10G 4P X710-T4L-t Adapter
- Intel® Ethernet 25G 2P XXV710 Adapter
- Intel® Ethernet 25G 2P XXV710 Mezz
- Intel® Ethernet 25G 2P E810-XXV OCP
- Intel® Ethernet 25G 2P E810-XXV Adapter
- Intel® Ethernet 40G 2P XL710 QSFP+ rNDC
- Intel® Ethernet Converged Network Adapter XL710-Q2

Building and Installation

To enable SR-IOV on your system:

1. Ensure both Virtualization and SR-IOV are enabled in the BIOS.
2. Install the Linux operating system. You can verify that the KVM driver is loaded by typing: `lsmod | grep -i kvm`
3. Load the Linux Base Driver using the modprobe command: `modprobe i40e option max_vfs=xx,yy`

xx and yy are the number of virtual functions you want to create. You must specify a number for each port with each parameter separated by a comma. For example, xx is the number of virtual functions for port 1; and yy, for port 2. You can create up to 63 functions per port.
4. Compile and install the iavf driver for SR-IOV. This is loaded against the virtual functions created.



NOTE: For VLANs, there is a limit of a total of 32 shared VLANs to 1 or more virtual functions.

There are three methods for installing the Linux driver:

- [Install from Source Code](#)
- [Install Using KMP RPM](#)
- [Install Using KMOD RPM](#)

Install from Source Code

To build a binary RPM* package of this driver, run 'rpmbuild -tb <filename.tar.gz>'. Replace <filename.tar.gz> with the specific filename of the driver.



NOTES:

- For the build to work properly it is important that the currently running kernel MATCH the version and configuration of the installed kernel source. If you have just recompiled your kernel, reboot the system.
- RPM functionality has only been tested in Red Hat distributions.

1. Download the base driver tar file to the directory of your choice. For example, use '/home/username/iavf' or '/usr/local/src/iavf'.
2. Untar/unzip the archive, where <x.x.x> is the version number for the driver tar:

```
# tar xzf iavf-<x.x.x>.tar.gz
```

3. Change to the driver src directory, where <x.x.x> is the version number for the driver tar:

```
# cd iavf-<x.x.x>/src/
```

4. Compile the driver module:

```
# make install
```

The binary will be installed as:

```
/lib/modules/<KERNEL_VERSION>/kernel/drivers/net/iavf/iavf.ko
```

The install locations listed above are the default locations. This might differ for various Linux distributions. For more information, see the lldistrib.txt file included in the driver tar.

5. Remove the old driver:

```
# rmmod iavf
```

6. Install the module using the modprobe command:

```
# modprobe iavf <parameter>=<value>
```

7. Update the system image:

```
dracut -f
```

8. Assign an IP address to and activate the Ethernet interface by entering the following, where <ethx> is the interface name:

```
# ifconfig <ethX> <IP_address> netmask <netmask> up
```

9. Verify that the interface works. Enter the following, where <IP_address> is the IP address for another machine on the same subnet as the interface that is being tested:

```
# ping <IP_address>
```

Install Using KMP RPM

The KMP RPMs update existing iavf RPMs currently installed on the system. These updates are provided by SuSE in the SLES release. If an RPM does not currently exist on the system, the KMP will not install.

The RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
intel-<component name>-<component version>.<arch type>.rpm
```

For example, intel-iavf-1.3.8.6-1.x86_64.rpm:

- iavf is the component name
- 1.3.8.6-1 is the component version
- x86_64 is the architecture type

KMP RPMs are provided for supported Linux distributions. The naming convention for the included KMP RPMs is:

```
intel-<component name>-kmp-<kernel type>-<component version>_<kernel version>.<arch
type>.rpm
```

For example, intel-iavf-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm:

- ia vf is the component name
- default is the kernel type
- 1.3.8.6 is the component version
- 2.6.27.19_5-1 is the kernel version
- x86_64 is the architecture type

To install the KMP RPM, type the following two commands:

```
# rpm -i <rpm filename>
# rpm -i <kmp rpm filename>
```

For example, to install the ia vf KMP RPM package, type the following:

```
# rpm -i intel-iavf-1.3.8.6-1.x86_64.rpm
# rpm -i intel-iavf-kmp-default-1.3.8.6_2.6.27.19_5-1.x86_64.rpm
```

Install Using KMOD RPM

The KMOD RPMs are provided for supported Linux distributions. The naming convention for the included RPMs is:

```
kmod-<driver name>-<version>-1.<arch type>.rpm
```

For example, kmod-iavf-2.3.4-1.x86_64.rpm:

- ia vf is the driver name
- 2.3.4 is the version
- x86_64 is the architecture type

To install the KMOD RPM, go to the directory of the RPM and type the following command:

```
# rpm -i <rpm filename>
```

For example, to install the ia vf KMOD RPM package, type the following:

```
# rpm -i kmod-iavf-2.3.4-1.x86_64.rpm
```

Command Line Parameters

The ia vf driver does not support any command line parameters.

Additional Configurations

Viewing Link Messages

Link messages will not be displayed to the console if the distribution is restricting system messages. In order to see network driver link messages on your console, set dmesg to eight by entering the following:

```
# dmesg -n 8
```



NOTE: This setting is not saved across reboots.

ethtool

The driver utilizes the ethtool interface for driver configuration and diagnostics, as well as displaying statistical information. The latest ethtool version is required for this functionality. Download it at: <https://kernel.org/pub/software/network/ethtool/>

Setting VLAN Tag Stripping

If you have applications that require Virtual Functions (VFs) to receive packets with VLAN tags, you can disable VLAN tag stripping for the VF. The Physical Function (PF) processes requests issued from the VF to enable or disable VLAN tag stripping. Note that if the PF has assigned a VLAN to a VF, then requests from that VF to set VLAN tag stripping will be ignored.

To enable/disable VLAN tag stripping for a VF, issue the following command from inside the VM in which you are running the VF:

```
# ethtool -K <ethX> rxvlan on/off
```

or alternatively:

```
# ethtool --offload <ethX> rxvlan on/off
```

IEEE 802.1ad (QinQ) Support

The IEEE 802.1ad standard, informally known as QinQ, allows for multiple VLAN IDs within a single Ethernet frame. VLAN IDs are sometimes referred to as "tags," and multiple VLAN IDs are thus referred to as a "tag stack." Tag stacks allow L2 tunneling and the ability to segregate traffic within a particular VLAN ID, among other uses.



NOTES:

- 802.1ad (QinQ) is supported in 3.19 and later kernels.
- Receive checksum offloads and VLAN acceleration are not supported for 802.1ad (QinQ) packets.

The following are examples of how to configure 802.1ad (QinQ):

```
# ip link add link eth0 eth0.24 type vlan proto 802.1ad id 24
```

```
# ip link add link eth0.24 eth0.24.371 type vlan proto 802.1Q id 371
```

Where "24" and "371" are example VLAN IDs.

Application Device Queues (ADQ)

Application Device Queues (ADQ) allow you to dedicate one or more queues to a specific application. This can reduce latency for the specified application, and allow Tx traffic to be rate limited per application.

Requirements:

- Kernel version 4.19.58 or later
- Depending on the underlying PF device, ADQ cannot be enabled when the following features are enabled: Data Center Bridging (DCB), Multiple Functions per Port (MFP), or Sideband Filters.
- If another driver (for example, DPDK) has set cloud filters, you cannot enable ADQ.

To create TCs on the interface:



NOTE: Run all TC commands from the `../iproute2/tc/` directory.

1. Use the `tc` command to create traffic classes (TCs). You can create a maximum of 8 TCs per interface.

```
# tc qdisc add dev <ethX> root mqprio num_tc <tcs> map <priorities> queues <count1@offset1 ...> hw 1 mode channel shaper bw_rlimit min_rate <min_rate1 ...> max_rate <max_rate1 ...>
```

Where:

- `num_tc <tcs>`: The number of TCs to use.
- `map <priorities>`: The map of priorities to TCs. You can map up to 16 priorities to TCs.
- `queues <count1@offset1 ...>`: For each TC, `<num queues>@<offset>`. The max total number of queues for all TCs is the number of cores.

- `hw 1 mode channel`: 'channel' with 'hw' set to 1 is a new hardware offload mode in `mqprio` that makes full use of the `mqprio` options, the TCs, the queue configurations, and the QoS parameters.
- `shaper bw_rlimit`: For each TC, sets the minimum and maximum bandwidth rates. The totals must be equal to or less than the port speed. This parameter is optional and is required only to set up the Tx rates.
- `min_rate <min_rate1>`: Sets the minimum bandwidth rate limit for each TC.
- `max_rate <max_rate1 ...>`: Sets the maximum bandwidth rate limit for each TC. You can set a min and max rate together.



NOTE: See the `mqprio` man page and the examples below for more information.

2. Verify the bandwidth limit using network monitoring tools such as `ifstat` or `sar -n DEV [interval] [number of samples]`



NOTE: Setting up channels via `ethtool` (`ethtool -L`) is not supported when the TCs are configured using `mqprio`.

3. Enable hardware TC offload on the interface:

```
# ethtool -K <ethX> hw-tc-offload on
```

4. Apply TCs to ingress (Rx) flow of the interface:

```
# tc qdisc add dev <ethX> ingress
```



NOTE:

- Tunnel filters are not supported in ADQ. If encapsulated packets do arrive in non-tunnel mode, filtering will be done on the inner headers. For example, for VXLAN traffic in non-tunnel mode, if PCTYPE is identified as a VXLAN encapsulated packet, then the outer headers are ignored. Therefore, inner headers are matched.
- If a TC filter on a PF matches traffic over a VF (on the PF), that traffic will be routed to the appropriate queue of the PF, and will not be passed on the VF. Such traffic will end up getting dropped higher up in the TCP/IP stack as it does not match PF address data.
- If traffic matches multiple TC filters that point to different TCs, that traffic will be duplicated and sent to all matching TC queues. The hardware switch mirrors the packet to a VSI list when multiple filters are matched.

Example:



NOTE: See the `tc` and `tc-flower` man pages for more information on traffic control and TC flower filters.

To set up two TCs (`tc0` and `tc1`), with 16 queues each, priorities 0-3 for `tc0` and 4-7 for `tc1`, and max Tx rate set to 1Gbit for `tc0` and 3Gbit for `tc1`:

```
# tc qdisc add dev ens4f0 root mqprio num_tc 2 map 0 0 0 0 1 1 1 1 queues 16@0 16@16
hw 1 mode channel shaper bw_rlimit max_rate 1Gbit 3Gbit
```

Where:

- `map 0 0 0 0 1 1 1 1`: Sets priorities 0-3 to use `tc0` and 4-7 to use `tc1`
- `queues 16@0 16@16`: Assigns 16 queues to `tc0` at offset 0 and 16 queues to `tc1` at offset 16

You can add multiple filters to the device, using the same recipe (and requires no additional recipe resources), either on the same interface or on different interfaces. Each filter uses the same fields for matching, but can have different match values.

```
tc filter add dev ethx protocol ip ingress prio 1 flower ip_proto tcp dst_port $app_port skip_sw
hw_tc 1
```

For example:

```
# tc filter add dev ethx protocol ip ingress prio 1 flower ip_proto tcp dst_port
5555 skip_sw hw_tc 1
```

SR-IOV Hypervisor Management Interface

The `sysfs` file structure below supports the SR-IOV hypervisor management interface.

`/sys/class/net/<ethX>/device/sriov` (see 1 below)

```

+--- [VF-id, 0 .. 255] (see 2 below)
| +--- trunk
| +--- vlan_mirror
| +--- egress_mirror
| +--- ingress_mirror
| +--- mac_anti_spoof
| +--- vlan_anti_spoof
| +--- loopback
| +--- mac
| +--- mac_list
| +--- promisc
| +--- vlan_strip
| +--- stats
| +--- link_state
| +--- max_tx_rate
| +--- min_tx_rate
| +--- spoofcheck
| +--- trust
| +--- vlan

```

**NOTES:**

1. kobject started from "sriov" is not available from existing kernel sysfs, and it requires device driver to implement this interface.
2. maximum number of SR-IOV instances is 256. The actual number of instances created depends on the value set for `/sys/bus/pci/devices/<device pci address>/sriov_numvfs`

SR-IOV hypervisor functions:

- trunk
 - Supports two operations:
 - add: adds one or more VLAN id into VF VLAN filtering.
 - rem: removes VLAN ids from the VF VLAN filtering list.
 - Example 1: add multiple VLAN tags, VLANs 2,4,5,10-20, by PF, p1p2, on a selected VF, 1, for filtering, with the sysfs support:


```
# echo add 2,4,5,10-20 > /sys/class/net/plp2/device/sriov/1/trunk
```
 - Example 2: remove VLANs 5, 11-13 from PF p1p2 VF 1 with sysfs:


```
# echo rem 5,11-13 > /sys/class/net/plp2/device/sriov/1/trunk
```
 - Note: for rem, if VLAN id is not on the VLAN filtering list, the VLAN id will be ignored.
- vlan_mirror
 - Supports both ingress and egress traffic mirroring.
 - Example 1: mirror traffic based upon VLANs 2,4,6,18-22 to VF 3 of PF p1p1.


```
# echo add 2,4,6,18-22 > /sys/class/net/plp1/device/sriov/3/vlan_mirror
```
 - Example 2: remove VLAN 4, 15-17 from traffic mirroring at destination VF 3.


```
# echo rem 15-17 > /sys/class/net/plp1/device/sriov/3/vlan_mirror
```
 - Example 3: remove all VLANs from mirroring at VF 3.


```
# echo rem 0 - 4095 > /sys/class/net/plp1/device/sriov/3/vlan_mirror
```

- **egress_mirror**
 - Supports egress traffic mirroring.
 - Example 1: add egress traffic mirroring on PF p1p2 VF 1 to VF 7.
echo add 7 > /sys/class/net/plp2/device/sriov/1/egress_mirror
 - Example 2: remove egress traffic mirroring on PF p1p2 VF 1 to VF 7.
echo rem 7 > /sys/class/net/plp2/device/sriov/1/egress_mirror
- **ingress_mirror**
 - Supports ingress traffic mirroring.
 - Example 1: mirror ingress traffic on PF p1p2 VF 1 to VF 7.
echo add 7 > /sys/class/net/plp2/device/sriov/1/ingress_mirror
 - Example 2: show current ingress mirroring configuration.
cat /sys/class/net/plp2/device/sriov/1/ingress_mirror
- **mac_anti_spoof**
 - Supports Enable/Disable MAC anti-spoof. Allows VFs to transmit packets with any SRC MAC, which is needed for some L2 applications as well as vNIC bonding within VMs if set to OFF.
 - Example 1: enable MAC anti-spoof for PF p2p1 VF 1.
echo ON /sys/class/net/plp2/device/sriov/1/mac_anti_spoof
 - Example 2: disable MAC anti-spoof for PF p2p1 VF 1.
echo OFF /sys/class/net/plp2/device/sriov/1/mac_anti_spoof
- **vlan_anti_spoof**
 - Supports Enable/Disable VLAN anti-spoof. Allows VFs to transmit packets only with VLAN tag specified in “trunk” settings, also will not allow to transmit “untagged” packets if set to ON. Violation have to increment tx_spoof stats counter.
 - Example 1: enable VLAN anti-spoof for PF p2p1 VF 1.
echo ON /sys/class/net/plp2/device/sriov/1/vlan_anti_spoof
 - Example 2: disable VLAN anti-spoof for PF p2p1 VF 1.
echo OFF /sys/class/net/plp2/device/sriov/1/vlan_anti_spoof
- **loopback**
 - Supports Enable/Disable VEB/MEPA (Local loopback).
 - Example 1: allow traffic switching between VFs on the same PF.
echo ON > /sys/class/net/plp2/device/sriov/loopback
 - Example 2: send Hairpin traffic to the switch to which the PF is connected.
echo OFF > /sys/class/net/plp2/device/sriov/loopback
 - Example 3: show loopback configuration.
cat /sys/class/net/plp2/device/sriov/loopback
- **mac**
 - Supports setting default MAC address. If MAC address is set by this command, the PF will not allow VF to change it using an MBOX request.
 - Example 1: set default MAC address to VF 1.
echo "00:11:22:33:44:55" > /sys/class/net/plp2/device/sriov/1/mac
 - Example 2: show default MAC address.
cat /sys/class/net/plp2/device/sriov/1/mac
- **mac_list**
 - Supports adding additional MACs to the VF. The default MAC is taken from "ip link set p1p2 vf 1 mac 00:11:22:33:44:55" if configured. If not, a random address is assigned to the VF by the NIC. If the MAC is configured using the IP LINK command, the VF cannot change it via MBOX/AdminQ requests.
 - Example 1: add mac 00:11:22:33:44:55 and 00:66:55:44:33:22 to PF p1p2 VF 1.
echo add "00:11:22:33:44:55,00:66:55:44:33:22" > /sys/class/net/plp2/device/sriov/1/mac_list
 - Example 2: delete mac 00:11:22:33:44:55 from above VF device.
echo rem 00:11:22:33:44:55 > /sys/class/net/plp2/device/sriov/1/mac_list
 - Example 3: display a VF MAC address list.
cat /sys/class/net/plp2/device/sriov/1/mac_list

- promisc
 - Supports setting/unsetting VF device unicast promiscuous mode and multicast promiscuous mode.
 - Example 1: set MCAST promiscuous on PF p1p2 VF 1.
`# echo add mcast > /sys/class/net/plp2/device/sriov/1/promisc`
 - Example 2: set UCAST promiscuous on PF p1p2 VF 1.
`# echo add ucast > /sys/class/net/plp2/device/sriov/1/promisc`
 - Example 3: unset MCAST promiscuous on PF p1p2 VF 1.
`# echo rem mcast > /sys/class/net/plp2/device/sriov/1/promisc`
 - Example 4: show current promiscuous mode configuration.
`# cat /sys/class/net/plp2/device/sriov/1/promisc`
- vlan_strip
 - Supports enabling/disabling VF device outer VLAN stripping
 - Example 1: enable VLAN strip on VF 3.
`# echo ON > /sys/class/net/plp1/device/sriov/3/vlan_strip`
 - Example 2: disable VLAN striping VF 3.
`# echo OFF > /sys/class/net/plp1/device/sriov/3/vlan_strip`
- stats
 - Supports getting VF statistics
 - Example 1: display stats of VF 1.
`# cat /sys/class/net/plp2/device/sriov/1/stats`
- link_state
 - Sets/displays link status.
 - Example 1: display link status on link speed.
`# cat /sys/class/net/plp2/device/sriov/1/link_state`
 - Example 2: set VF 1 to track status of PF link.
`# echo auto > /sys/class/net/plp2/device/sriov/1/link_state`
 - Example 3: disable VF 1.
`# echo disable > /sys/class/net/plp2/device/sriov/1/link_state`

Known Issues

Software Issues



NOTE: After installing the driver, if your Intel® Ethernet Network Connection is not working, verify that you have installed the correct driver.

Linux bonding fails with Virtual Functions bound to an Intel® Ethernet Controller 700 series based device

If you bind Virtual Functions (VFs) to an Intel® Ethernet Controller 700 series based device, the VF slaves may fail when they become the active slave. If the MAC address of the VF is set by the PF (Physical Function) of the device, when you add a slave, or change the active-backup slave, Linux bonding tries to sync the backup slave's MAC address to the same MAC address as the active slave. Linux bonding will fail at this point. This issue will not occur if the VF's MAC address is not set by the PF.

Traffic Is Not Being Passed Between VM and Client

You may not be able to pass traffic between a client system and a Virtual Machine (VM) running on a separate host if the Virtual Function (VF, or Virtual NIC) is not in trusted mode and spoof checking is enabled on the VF. Note that this situation can occur in any combination of client, host, and guest operating system. See the readme for the PF driver for information on spoof checking and how to set the VF to trusted mode.

Do not unload port driver if VF with active VM is bound to it



NOTE: Do not unload a port's driver if a Virtual Function (VF) with an active Virtual Machine (VM) is bound to it. Doing so will cause the port to appear to hang. Once the VM shuts down, or otherwise releases the VF, the command will complete.

Using four traffic classes fails

Do not try to reserve more than three traffic classes in the iavf driver. Doing so will fail to set any traffic classes and will cause the driver to write errors to stdout. Use a maximum of three queues to avoid this issue.

Multiple log error messages on iavf driver removal

If you have several VFs and you remove the iavf driver, several instances of the following log errors are written to the log:

```
Unable to send opcode 2 to PF, err I40E_ERR_QUEUE_EMPTY, aq_err ok
Unable to send the message to VF 2 aq_err 12
ARQ Overflow Error detected
```

Virtual machine does not get link

If the virtual machine has more than one virtual port assigned to it, and those virtual ports are bound to different physical ports, you may not get link on all of the virtual ports. The following command may work around the issue:

```
# ethtool -r <ethX>
```

Where <ethX> is the PF interface in the host, for example: p5p1. You may need to run the command more than once to get link on all virtual ports.

MAC Address of Virtual Function Changes Unexpectedly

If a Virtual Function's MAC address is not assigned in the host, then the VF (virtual function) driver will use a random MAC address. This random MAC address may change each time the VF driver is reloaded. You can assign a static MAC address in the host machine. This static MAC address will survive a VF driver reload.

Compiling the Driver

When trying to compile the driver by running `make install`, the following error may occur: "Linux kernel source not configured - missing version.h"

To solve this issue, create the `version.h` file by going to the Linux source tree and entering:

```
# make include/linux/version.h
```

Multiple Interfaces on Same Ethernet Broadcast Network

Due to the default ARP behavior on Linux, it is not possible to have one system on two IP networks in the same Ethernet broadcast domain (non-partitioned switch) behave as expected. All Ethernet interfaces will respond to IP traffic for any IP address assigned to the system. This results in unbalanced receive traffic.

If you have multiple interfaces in a server, either turn on ARP filtering by entering:

```
# echo 1 > /proc/sys/net/ipv4/conf/all/arp_filter
```

This only works if your kernel's version is higher than 2.4.5.



NOTE: This setting is not saved across reboots. The configuration change can be made permanent by adding the following line to the file `/etc/sysctl.conf`:

```
net.ipv4.conf.all.arp_filter = 1
```

Another alternative is to install the interfaces in separate broadcast domains (either in different switches or in a switch partitioned to VLANs).

Rx Page Allocation Errors

'Page allocation failure. order:0' errors may occur under stress with kernels 2.6.25 and newer. This is caused by the way the Linux kernel reports this stressed condition.

Host May Reboot after Removing PF when VF is Active in Guest

Using kernel versions earlier than 3.2, do not unload the PF driver with active VFs. Doing this will cause your VFs to stop working until you reload the PF driver and may cause a spontaneous reboot of your system.

Prior to unloading the PF driver, you must first ensure that all VFs are no longer active. Do this by shutting down all VMs and unloading the VF driver.

SR-IOV virtual functions have identical MAC addresses in RHEL8

When you create multiple SR-IOV virtual functions on Red Hat Enterprise Linux 8, the VFs may have identical MAC addresses. Only one VF will pass traffic, and all traffic on other VFs with identical MAC addresses will fail. This is related to the "MACAddressPolicy=persistent" setting in `/usr/lib/systemd/network/99-default.link`.

To resolve this issue, edit the `/usr/lib/systemd/network/99-default.link` file and change the `MACAddressPolicy` line to `"MACAddressPolicy=none"`. For more information, see the `systemd` documentation.

VMWare ESX Drivers and Support

Driver types

Intel provides the following types of drivers for VMware ESX:

- Native mode drivers are the default driver for the VMware ESX environment. They are interrupt driven and developed using VMware's native mode API.
- Enhanced Network Stack (ENS) drivers are intended for use in VMware NSX-T deployments. These drivers are polling mode drivers.
- Unified drivers support both interrupt and poll mode operation. Depending on the deployment model, the driver automatically utilizes the appropriate mode.

This release contains Native, ENS, and Unified drivers as follows:

Driver	Device Family	ESXi 6.7 U3	ESXi 7.0 U1
icen	Intel® Ethernet Controller 800 series devices	Unified	Unified
i40en	Intel® Ethernet Controller 700 series devices	Native, ENS	Native, ENS
ixgben	Intel® Ethernet Controller X550 series devices	Native, ENS	Native, ENS
	Intel® Ethernet Controller X540 series devices		
	Intel® Ethernet Controller X520 series devices		
igbn	Intel® Ethernet Controller X540 series devices	Native	Native

You can check your hardware compatibility at <http://www.vmware.com/resources/compatibility/search.php>

Installation and Configuration

For detailed installation information, please refer to the the readme and release notes files included with the driver.

For information on VMware ENS configuration, please refer to VMware documentation under the topics of Enhanced Network Stack (ENS) or Enhanced Data Path. To use an Intel ENS capable driver, you must install VMware NSX-T on the system hosting the adapters. VMware's NSX manager must then be used to create logical switches and attach ENS capable drivers to the adapter ports.



NOTE: To install the ENS driver, you must first install the native driver. You must leave the native driver installed for the ENS driver to function.

Remote Boot

Remote Boot allows you to boot a system using only an Ethernet adapter. You connect to a server that contains an operating system image and use that to boot your local system.

Flash Images

"Flash" is a generic term for nonvolatile RAM (NVRAM), firmware, and option ROM (OROM). Depending on the device, it can be on the NIC or on the system board.



NOTE: You cannot update the flash of a device in the "Pending Reboot" state. Reboot your system before attempting to update the device's flash.

Updating the Flash from Linux

The BootUtil command line utility can update the flash on an Intel Ethernet network adapter. Run BootUtil with the following command line options to update the flash on all supported Intel network adapters. For example, enter the following command line:

```
bootutil64e -up=efi -all
```

BootUtil can only be used to program add-in Intel network adapters. LOM (LAN On Motherboard) network connections cannot be programmed with the UEFI network driver option ROM.

See the bootutil.txt file for details on using BootUtil.

Installing the UEFI Network Driver Option ROM from the UEFI Shell

The BootUtil command line utility can install the UEFI network driver on an Intel network adapter's option ROM. The UEFI network driver will load automatically during system UEFI boot when installed into the option ROM. For example, run BootUtil with the following command line options to install the UEFI network driver on all supported Intel network adapters:

```
FS0:\>bootutil64e -up=efi -all
```

BootUtil can only be used to program add-in Intel Ethernet network adapters. LOM (LAN On Motherboard) network connections cannot be programmed with the UEFI network driver option ROM.

See the bootutil.txt file for details on using BootUtil.

Enable Remote Boot

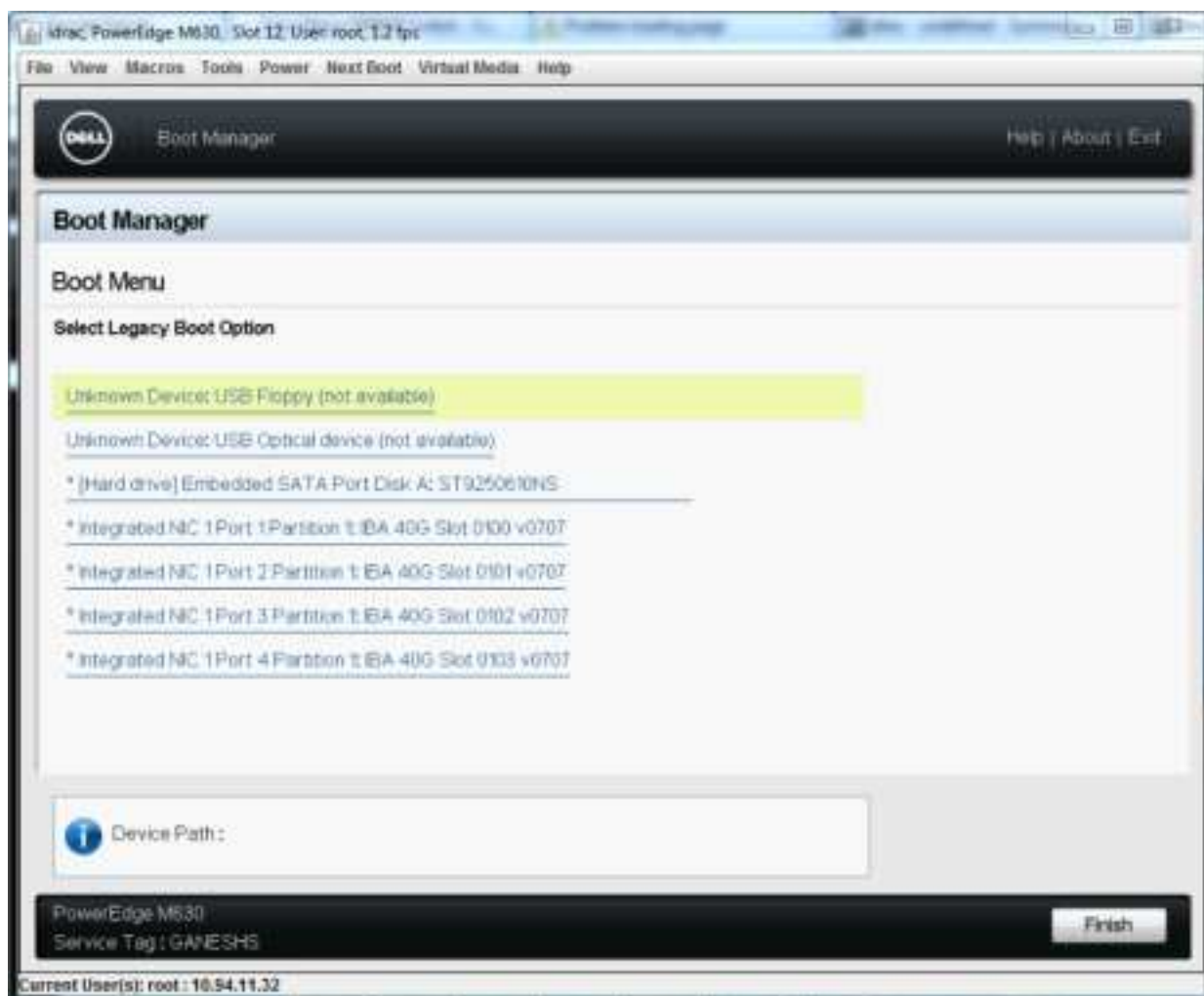
If you have an Intel Desktop Adapter installed in your client computer, the flash ROM device is already available in your adapter, and no further installation steps are necessary. For Intel Server Adapters, the flash ROM can be enabled using the BootUtil utility. For example, from the command line type:

```
BOOTUTIL -E  
BOOTUTIL -NIC=1 -FLASHENABLE
```

The first line will enumerate the ports available in your system. Choose a port. Then type the second line, selecting the port you wish to enable. For more details, see the bootutil.txt file.

Intel Adapters in the Boot Menu

The Boot Menu section of the Boot Manager will report the PXE-enabled ports on an Intel X710-based adapter as being 40G ports, as illustrated in the following figure. The ports on an X710 adapter are, in fact, 10G ports.



In the Boot Manager Boot Menu, Intel adapters are identified as follows:

- X710-controlled adapters: "IBA 40G"
- Other 10G adapters: "IBA XE"
- 1G adapters: "IBA 1G"

Intel® Boot Agent Configuration

Boot Agent Client Configuration

The Boot Agent is enabled and configured from HII.



CAUTION: If spanning tree protocol is enabled on a switch port through which a port is trying to use PXE, the delay before the port starts forwarding can cause a DHCP timeout. Either disable spanning tree or turn on the feature that allows the port to begin forwarding of packets immediately ('port fast learning' for Cisco switches), rather than wait until the spanning tree discovery is complete.

Intel Boot Agent Target/Server Setup

Overview

For the Intel® Boot Agent software to perform its intended job, there must be a server set up on the same network as the client computer. That server must recognize and respond to the PXE or BOOTP boot protocols that are used by the Intel Boot Agent software.

 **NOTE:** When the Intel Boot Agent software is installed as an upgrade for an earlier version boot ROM, the associated server-side software may not be compatible with the updated Intel Boot Agent. Contact your system administrator to determine if any server updates are necessary.

Linux* Server Setup

Consult your Linux* vendor for information about setting up the Linux Server.

Windows* Deployment Services

Nothing is needed beyond the standard driver files supplied on the media. Microsoft* owns the process and associated instructions for Windows Deployment Services. For more information on Windows Deployment Services perform a search of Microsoft articles at: <http://technet.microsoft.com/en-us/library/default.aspx>

Intel® Boot Agent Messages

Message Text	Cause
Invalid PMM function number.	PMM is not installed or is not working correctly. Try updating the BIOS.
PMM allocation error.	PMM could not or did not allocate the requested amount of memory for driver usage.
Option ROM initialization error. 64-bit PCI BAR addresses not supported, AX=	This may be caused by the system BIOS assigning a 64-bit BAR (Base Address Register) to the network port. Running the BootUtil utility with the -64d command line option may resolve this issue. To work around the issue on Intel® Ethernet X710 or XL710 based adapters, disable NPar and NParEP. Alternatively, put the system into UEFI boot mode.
PXE-E00: This system does not have enough free conventional memory. The Intel Boot Agent cannot continue.	System does not have enough free memory to run PXE image. The Intel Boot Agent was unable to find enough free base memory (below 640K) to install the PXE client software. The system cannot boot via PXE in its current configuration. The error returns control to the BIOS and the system does not attempt to remote boot. If this error persists, try updating your system's BIOS to the most-recent version. Contact your system administrator or your computer vendor's customer support to resolve the problem.
PXE-E01: PCI Vendor and Device IDs do not match!	Image vendor and device ID do not match those located on the card. Make sure the correct flash image is installed on the adapter.
PXE-E04: Error reading PCI configuration space. The Intel Boot Agent cannot con-	PCI configuration space could not be read. Machine is probably not PCI compliant. The Intel Boot Agent was unable to read one or more of the adapter's PCI configuration registers. The adapter may be mis-configured, or the wrong Intel Boot Agent image may be installed on the adapter. The Intel Boot Agent will return control to the BIOS and not attempt to remote boot. Try to update the flash image. If this does not solve the problem, contact your system administrator or Intel Cus-

tinue.	tomer Support.
PXE-E05: The LAN adapter's configuration is corrupted or has not been initialized. The Intel Boot Agent cannot continue.	The adapter's EEPROM is corrupted. The Intel Boot Agent determined that the adapter EEPROM checksum is incorrect. The agent will return control to the BIOS and not attempt to remote boot. Try to update the flash image. If this does not solve the problem, contact your system administrator or Intel Customer Support .
PXE-E06: Option ROM requires DDIM support.	The system BIOS does not support DDIM. The BIOS does not support the mapping of the PCI expansion ROMs into upper memory as required by the PCI specification. The Intel Boot Agent cannot function in this system. The Intel Boot Agent returns control to the BIOS and does not attempt to remote boot. You may be able to resolve the problem by updating the BIOS on your system. If updating your system's BIOS does not fix the problem, contact your system administrator or your computer vendor's customer support to resolve the problem.
PXE-E07: PCI BIOS calls not supported.	BIOS-level PCI services not available. Machine is probably not PCI compliant.
PXE-E09: Unexpected UNDI loader error. Status == xx	The UNDI loader returned an unknown error status. xx is the status returned.
PXE-E0C: Firmware recovery mode detected. Initialization failed.	The adapter is in firmware recovery mode. Refer to the "Firmware Recovery Mode" section of this document for details.
PXE-E20: BIOS extended memory copy error.	BIOS could not move the image into extended memory.
PXE-E20: BIOS extended memory copy error. AH == xx	Error occurred while trying to copy the image into extended memory. xx is the BIOS failure code.
PXE-E51: No DHCP or BOOTP offers received.	The Intel Boot Agent did not receive any DHCP or BOOTP responses to its initial request. Please make sure that your DHCP server (and/or proxyDHCP server, if one is in use) is properly configured and has sufficient IP addresses available for lease. If you are using BOOTP instead, make sure that the BOOTP service is running and is properly configured.
PXE-E53: No boot filename received.	The Intel Boot Agent received a DHCP or BOOTP offer, but has not received a valid filename to download. If you are using PXE, please check your PXE and BINL configuration. If using BOOTP, be sure that the service is running and that the specific path and filename are correct.
PXE-E61: Media test failure.	The adapter does not detect link. Please make sure that the cable is good and is attached to a working hub or switch. The link light visible from the back of the adapter should be lit.
PXE-EC1: Base-code ROM ID structure was not found.	No base code could be located. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image.

PXE-EC3: BC ROM ID structure is invalid.	Base code could not be installed. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image.
PXE-EC4: UNDI ID structure was not found.	UNDI ROM ID structure signature is incorrect. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image.
PXE-EC5: UNDI ROM ID structure is invalid.	The structure length is incorrect. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image.
PXE-EC6: UNDI driver image is invalid.	The UNDI driver image signature was invalid. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image.
PXE-EC8: !PXE structure was not found in UNDI driver code segment.	The Intel Boot Agent could not locate the needed !PXE structure resource. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image. This may also be caused by the system BIOS assigning a 64-bit BAR (Base Address Register) to the network port. Running the BootUtil utility with the -64d command line option may resolve this issue.
PXE-EC9: PXENV+ structure was not found in UNDI driver code segment.	The Intel Boot Agent could not locate the needed PXENV+ structure. An incorrect flash image is installed or the image has become corrupted. Try to update the flash image.
PXE-M0F: Exiting Intel Boot Agent.	Ending execution of the ROM image.
This option has been locked and cannot be changed.	You attempted to change a configuration setting that has been locked by your system administrator. This message can appear either from within Intel® PROSet's Boot Options tab when operating under Windows* or from the Configuration Setup Menu when operating in a stand-alone environment. If you think you should be able to change the configuration setting, consult your system administrator.
PXE-M0E: Retrying network boot; press ESC to cancel.	The Intel Boot Agent did not successfully complete a network boot due to a network error (such as not receiving a DHCP offer). The Intel Boot Agent will continue to attempt to boot from the network until successful or until canceled by the user. This feature is disabled by default. For information on how to enable this feature, contact Intel Customer Support .

Intel Boot Agent Troubleshooting Procedures

Common Issues

The following list of problems and associated solutions covers a representative set of problems that you might encounter while using the Intel Boot Agent.

After booting, my computer experiences problems

After the Intel® Boot Agent product has finished its sole task (remote booting), it no longer has any effect on the client computer operation. Thus, any issues that arise after the boot process is complete are most likely not related to the Intel Boot Agent product.

If you are having problems with the local (client) or network operating system, contact the operating system manufacturer for assistance. If you are having problems with some application program, contact the application manufacturer for assistance. If you are having problems with any of your computer's hardware or with the BIOS, contact your computer system manufacturer for assistance.

Cannot change boot order

If you are accustomed to redefining your computer's boot order using the motherboard BIOS setup program, the default settings of the Intel Boot Agent setup program can override that setup. To change the boot sequence, you must first override the Intel Boot Agent setup program defaults. A configuration setup menu appears allowing you to set configuration values for the Intel Boot Agent. To change your computer's boot order setting, see [Configuring the Boot Agent in a Pre-boot PXE Environment](#).

My computer does not complete POST

If your computer fails to boot with an adapter installed, but *does* boot when you remove the adapter, try moving the adapter to another computer and using BootUtil to disable the Flash ROM.

If this does not work, the problem may be occurring before the Intel Boot Agent software even begins operating. In this case, there may be a BIOS problem with your computer. Contact your computer manufacturer's customer support group for help in correcting your problem.

There are configuration/operation problems with the boot process

If your PXE client receives a DHCP address, but then fails to boot, you know the PXE client is working correctly. Check your network or PXE server configuration to troubleshoot the problem. Contact [Intel Customer Support](#) if you need further assistance.

PXE option ROM does not follow the PXE specification with respect to the final "discover" cycle

In order to avoid long wait periods, the option ROM no longer includes the final 32-second discover cycle. (If there was no response in the prior 16-second cycle, it is almost certain that there will be none in the final 32-second cycle.)

Known Issues

Incorrect port numbering in the Boot Options Menu

On certain platforms, some device entries in the legacy PXE option ROM Boot Option Menu are pre-pended with identical port number information (first part of the string that comes from BIOS). This is not an option ROM issue. The first device option ROM to be initialized on a platform exposes all boot options for the device. This is misinterpreted by BIOS. The second part of the string from the option ROM indicates the correct slot (port) numbers.

iSCSI Boot Configuration

iSCSI Initiator Setup

Configuring Intel® Ethernet iSCSI Boot on a Microsoft* Windows* Client Initiator

Requirements

1. Make sure the iSCSI initiator system starts the iSCSI Boot firmware. The firmware should be configured properly, be able to connect to iSCSI target, and detect the boot disk.
2. You will need Microsoft* iSCSI Software Initiator with integrated software boot support. This boot version of the initiator is available [here](#).
3. To enable crash dump support, follow the steps in [Crash Dump Support](#).

Configuring Intel® Ethernet iSCSI Boot on a Linux* Client Initiator

1. Install the Open-iSCSI initiator utilities.

```
#yum -y install iscsi-initiator-utils
```

2. Refer to the README file found at <https://github.com/mikechristie/open-iscsi>.
3. Configure your iSCSI array to allow access.
 - a. Examine `/etc/iscsi/initiatorname.iscsi` for the Linux host initiator name.
 - b. Update your volume manager with this host initiator name.
4. Set iscsi to start on boot.

```
#chkconfig iscsd on
#chkconfig iscsi on
```

5. Start iSCSI service (192.168.x.x is the IP Address of your target).

```
#iscsiadm -n discovery -t s -p 192.168.x.x
```

Observe the target names returned by iscsi discovery.

6. Log onto the target (-m XXX -T is XXX -l XXX -).

```
iscsiadm -m node -T iqn.2123-01.com:yada:yada: -p 192.168.2.124 -l
```

iSCSI Boot POST Setup

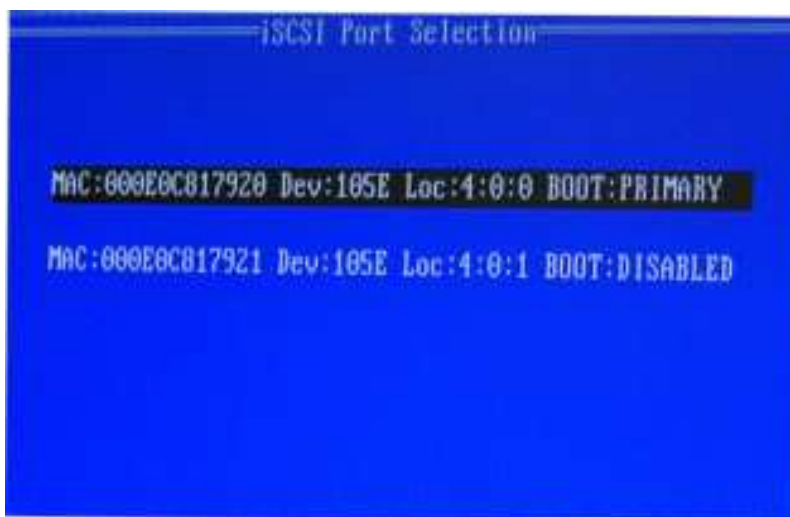
Intel® Ethernet iSCSI Boot features a setup menu which allows two network ports in one system to be enabled as iSCSI Boot devices. To configure Intel® iSCSI Boot, power-on or reset the system and press the Ctrl-D key when the message "Press <Ctrl-D> to run setup..." is displayed. After pressing the Ctrl-D key, you will be taken to the Intel® iSCSI Boot Port Selection Setup Menu.



NOTE: When booting an operating system from a local disk, Intel® Ethernet iSCSI Boot should be disabled for all network ports.

Intel® Ethernet iSCSI Boot Port Selection Menu

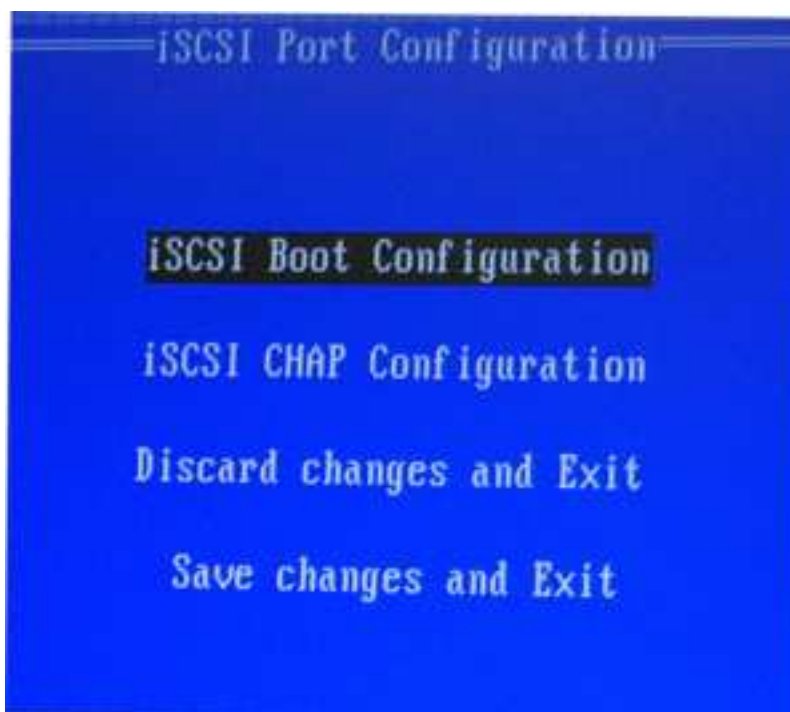
The first screen of the Intel® iSCSI Boot Setup Menu displays a list of Intel® iSCSI Boot-capable adapters. For each adapter port the associated PCI device ID, PCI bus/device/function location, and a field indicating Intel® Ethernet iSCSI Boot status is displayed. Up to 10 iSCSI Boot-capable ports are displayed within the Port Selection Menu. If there are more Intel® iSCSI Boot-capable adapters, these are not listed in the setup menu.



The usage of this menu is described below:

- One network port in the system can be selected as the primary boot port by pressing the 'P' key when highlighted. The primary boot port will be the first port used by Intel® Ethernet iSCSI Boot to connect to the iSCSI target. Only one port may be selected as a primary boot port.
- One network port in the system can be selected as the secondary boot port by pressing the 'S' key when highlighted. The secondary boot port will only be used to connect to the iSCSI target disk if the primary boot port fails to establish a connection. Only one port may be selected as a secondary boot port.
- Pressing the 'D' key with a network port highlighted will disable Intel® Ethernet iSCSI Boot on that port.
- Pressing the 'B' key with a network port highlighted will blink an LED on that port.
- Press the Esc key to leave the screen.

Intel® Ethernet iSCSI Boot Port Specific Setup Menu



The port specific iSCSI setup menu has four options:

- **Intel® iSCSI Boot Configuration** - Selecting this option will take you to the iSCSI Boot Configuration Setup Menu. The [iSCSI Boot Configuration Menu](#) is described in detail in the section below and will allow you to configure the iSCSI parameters for the selected network port.
- **CHAP Configuration** - Selecting this option will take you to the CHAP configuration screen. The [CHAP Configuration Menu](#) is described in detail in the section below.
- **Discard Changes and Exit** - Selecting this option will discard all changes made in the iSCSI Boot Configuration and CHAP Configuration setup screens, and return back to the iSCSI Boot Port Selection Menu.
- **Save Changes and Exit** - Selecting this option will save all changes made in the iSCSI Boot Configuration and CHAP Configuration setup screens. After selecting this option, you will return to the iSCSI Boot Port Selection Menu.

Intel® iSCSI Boot Configuration Menu

The Intel® iSCSI Boot Configuration Menu allows you to configure the iSCSI Boot and Internet Protocol (IP) parameters for a specific port. The iSCSI settings can be configured manually or retrieved dynamically from a DHCP server.

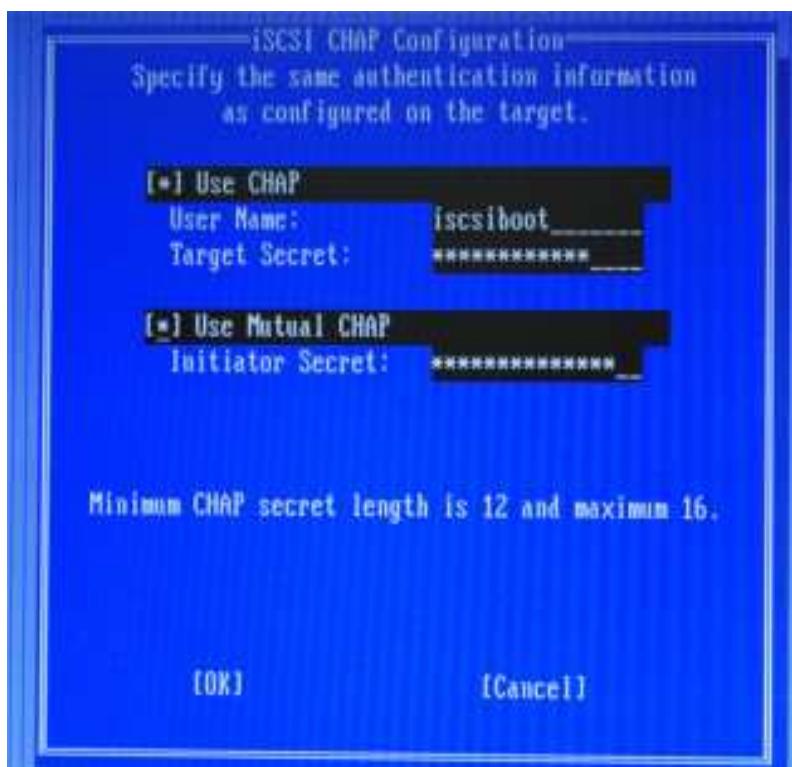


Listed below are the options in the Intel® iSCSI Boot Configuration Menu:

- **Use Dynamic IP Configuration (DHCP)** - Selecting this checkbox will cause iSCSI Boot to attempt to get the client IP address, subnet mask, and gateway IP address from a DHCP server. If this checkbox is enabled, these fields will not be visible.
- **Initiator Name** - Enter the iSCSI initiator name to be used by Intel® iSCSI Boot when connecting to an iSCSI target. The value entered in this field is global and used by all iSCSI Boot-enabled ports in the system. This field may be left blank if the "Use DHCP For Target Configuration" checkbox is enabled. For information on how to retrieve the iSCSI initiator name dynamically from a DHCP server see the section [DHCP Server Configuration](#).
- **Initiator IP** - Enter the client IP address to be used for this port as static IP configuration in this field. This IP address will be used by the port during the entire iSCSI session. This option is visible if DHCP is not enabled.
- **Subnet Mask** - Enter the IP subnet-mask in this field. This should be the IP subnet mask used on the network which the selected port will be connecting to for iSCSI. This option is visible if DHCP is not enabled.
- **Gateway IP** - Enter the IP address of the network gateway in this field. This field is necessary if the iSCSI target is located on a different sub-network than the selected Intel® iSCSI Boot port. This option is visible if DHCP is not enabled.
- **Use DHCP for iSCSI Target Information** - Selecting this checkbox will cause Intel® iSCSI Boot to attempt to gather the iSCSI target's IP address, IP port number, iSCSI target name, and SCSI LUN ID from a DHCP server on the network. For information on how to configure the iSCSI target parameters using DHCP see the section [DHCP Server Configuration](#). When this checkbox is enabled, these fields will not be visible.
- **Target Name** - Enter the IQN name of the iSCSI target in this field. This option is visible if DHCP for iSCSI target is not enabled.
- **Target IP** - Enter the target IP address of the iSCSI target in this field. This option is visible if DHCP for iSCSI target is not enabled.
- **Target Port** - TCP Port Number.
- **Boot LUN** - Enter the LUN ID of the boot disk on the iSCSI target in this field. This option is visible if DHCP for iSCSI target is not enabled.

iSCSI CHAP Configuration

Intel® iSCSI Boot supports Mutual CHAP MD5 authentication with an iSCSI target. Intel® iSCSI Boot uses the "MD5 Message Digest Algorithm" developed by RSA Data Security, Inc.



The iSCSI CHAP Configuration menu has the following options to enable CHAP authentication:

- **Use CHAP** - Selecting this checkbox will enable CHAP authentication for this port. CHAP allows the target to authenticate the initiator. After enabling CHAP authentication, a user name and target password must be entered.
- **User Name** - Enter the CHAP user name in this field. This must be the same as the CHAP user name configured on the iSCSI target.
- **Target Secret** - Enter the CHAP password in this field. This must be the same as the CHAP password configured on the iSCSI target and must be between 12 and 16 characters in length. This password can not be the same as the **Initiator Secret**.
- **Use Mutual CHAP** - Selecting this checkbox will enable Mutual CHAP authentication for this port. Mutual CHAP allows the initiator to authenticate the target. After enabling Mutual CHAP authentication, an initiator password must be entered. Mutual CHAP can only be selected if Use CHAP is selected.
- **Initiator Secret** - Enter the Mutual CHAP password in this field. This password must also be configured on the iSCSI target and must be between 12 and 16 characters in length. This password can not be the same as the **Target Secret**.

The CHAP Authentication feature of this product requires the following acknowledgments:

This product includes cryptographic software written by Eric Young (eay@cryptsoft.com). This product includes software written by Tim Hudson (tjh@cryptsoft.com).

This product includes software developed by the OpenSSL Project for use in the OpenSSL Toolkit. (<http://www.openssl.org/>).

Intel® PROSet

Many of the functions of the Intel® iSCSI Boot Port Selection Setup Menu can also be configured or revised from Windows Device Manager or the Intel® PROSet Adapter Configuration Utility. Open the adapter's property sheet and select the **Data Options** tab. You must install the latest Intel Ethernet Adapter drivers and software to access this.

iSCSI Boot Target Configuration

For specific information on configuring your iSCSI target system and disk volume, refer to instructions provided by your system or operating system vendor. Listed below are the basic steps necessary to setup Intel® Ethernet iSCSI Boot to work with most iSCSI target systems. The specific steps will vary from one vendor to another.



NOTES:

- To support iSCSI Boot, the target needs to support multiple sessions from the same initiator. Both the iSCSI Boot firmware initiator and the OS High Initiator need to establish an iSCSI session at the same time. Both these initiators use the same Initiator Name and IP Address to connect and access the OS disk but these two initiators will establish different iSCSI sessions. In order for the target to support iSCSI Boot, the target must be capable of supporting multiple sessions and client logins.
- If you set your client to BIOS boot mode, and your target boot server is a Linux target, your client may fail to boot with the error "Fail to find a suitable stage 1 device." Either set your client to UEFI boot mode or add the following to the iSCSI target server kernel command line: `rd.net.timeout.carrier=15`

1. Configure a disk volume on your iSCSI target system. Note the LUN ID of this volume for use when configuring in Intel® Ethernet iSCSI Boot firmware setup.

2. Note the iSCSI Qualified Name (IQN) of the iSCSI target, which will likely look like:

```
iqn.1986-03.com.intel:target1
```

This value is used as the iSCSI target name when you configuring your initiator system's Intel® Ethernet iSCSI Boot firmware.

3. Configure the iSCSI target system to accept the iSCSI connection from the iSCSI initiator. This usually requires listing the initiator's IQN name or MAC address for permitting the initiator to access to the disk volume. See the [Firmware Setup](#) section for information on how to set the iSCSI initiator name.
4. One-way authentication protocol can optionally be enabled for secure communications. Challenge-Handshake Authentication Protocol (CHAP) is enabled by configuring username/password on iSCSI target system. For setting up CHAP on the iSCSI initiator, refer to the section [Firmware Setup](#) for information.

Bootting from Targets Larger than 2TB

You can connect and boot from a target LUN that is larger than 2 Terabytes with the following restrictions:

- The block size on the target must be 512 bytes
- The following operating systems are supported:
 - VMware* ESXi 6.0, or later
 - Red Hat* Enterprise Linux* 6.3, or later
 - SUSE* Enterprise Linux 11SP2, or later
 - Microsoft* Windows Server* 2012 R2, or later
- You may be able to access data only within the first 2 TB.



NOTE: The Crash Dump driver does not support target LUNs larger than 2TB.

DHCP Server Configuration

If you are using Dynamic Host Configuration Protocol (DHCP), the DHCP server needs to be configured to provide the iSCSI Boot configurations to the iSCSI initiator. You must set up the DHCP server to specify Root Path option 17 and Host Name option 12 to respond iSCSI target information back to the iSCSI initiator. DHCP option 3, Router List may be necessary, depending on the network configuration.

DHCP Root Path Option 17:

The iSCSI root path option configuration string uses the following format:

```
iscsi:<server name or IP address>:<protocol>:<port>:<LUN>:<targetname>
```

- **Server name:** DHCP server name or valid IPv4 address literal.
Example: 192.168.0.20.
- **Protocol:** Transportation protocol used by iSCSI. Default is tcp (6).
No other protocols are currently supported.
- **Port:** Port number of the iSCSI. A default value of 3260 will be used if this field is left blank.
- **LUN:** LUN ID configured on iSCSI target system. Default is zero.
- **Target name:** iSCSI target name to uniquely identify an iSCSI target in IQN format.
Example: iqn.1986-03.com.intel:target1

DHCP Host Name Option 12:

Configure option 12 with the hostname of the iSCSI initiator.

DHCP Option 3, Router List:

Configure option 3 with the gateway or Router IP address, if the iSCSI initiator and iSCSI target are on different subnets.

Creating a Bootable Image for an iSCSI Target

There are two ways to create a bootable image on an iSCSI target:

- Install directly to a hard drive in an iSCSI storage array (Remote Install).
- Install to a local disk drive and then transfer this disk drive or OS image to an iSCSI Target (Local Install).

Microsoft* Windows*

Microsoft* Windows Server* natively supports OS installation to an iSCSI target without a local disk and also natively supports OS iSCSI boot. See Microsoft's installation instructions and Windows Deployment Services documentation for details.

SUSE* Linux Enterprise Server

For the easiest experience installing Linux onto an iSCSI target, you should use SLES10 or greater. SLES10 provides native support for iSCSI Booting and installing. This means that there are no additional steps outside of the installer that are necessary to install to an iSCSI target using an Intel Ethernet Server Adapter. Please refer to your SLES documentation for instructions on how to install to an iSCSI LUN.

Red Hat Enterprise Linux

For the easiest experience installing Linux onto an iSCSI target, you should use RHEL 5.1 or greater. RHEL 5.1 provides native support for iSCSI Booting and installing. This means that there are no additional steps outside of the installer that are necessary to install to an iSCSI target using an Intel Ethernet Server Adapter. Please refer to your RHEL documentation for instructions on how to install to an iSCSI LUN.

Microsoft Windows Server iSCSI Crash Dump Support

Crash dump file generation is supported for iSCSI-booted Windows Server x64 by the Intel iSCSI Crash Dump Driver. To ensure a full memory dump is created:

1. Set the page file size equal to or greater than the amount of RAM installed on your system is necessary for a full memory dump.
2. Ensure that the amount of free space on your hard disk is able to handle the amount of RAM installed on your system.

To setup crash dump support follow these steps:

1. Setup Windows iSCSI Boot.
2. If you have not already done so, install the latest Intel Ethernet Adapter drivers and Intel PROSet.
3. Open Intel PROSet for Windows Device Manager or Intel® PROSet Adapter Configuration Utility and select the Boot Options Tab.
4. From Settings select iSCSI Boot Crash Dump and the Value Enabled and click OK.

iSCSI Troubleshooting

The table below lists problems that can possibly occur when using Intel® Ethernet iSCSI Boot. For each problem a possible cause and resolution are provided.

Problem	Resolution
Intel® Ethernet iSCSI Boot does not load on system start-up and the sign-on banner is not displayed.	• While the system logon screen may display for a longer time during system start-up, Intel Ethernet iSCSI Boot may not be displayed during POST. It may be necessary to disable a system BIOS feature in order to display messages from Intel iSCSI Remote Boot. From the system BIOS Menu, disable any quiet boot or quick boot options. Also disable any BIOS splash screens. These options may be suppressing output from Intel iSCSI Remote Boot. • Intel Ethernet iSCSI Remote Boot has not been installed on the adapter or the adapter's flash ROM is disabled. Update the network adapter using the latest version of BootUtil as described in the Flash Images section of this document. If BootUtil reports the flash ROM is disabled, use the "BOOTUTIL -flashenable" command to enable the flash ROM and update the adapter. • The system BIOS may be suppressing output from Intel Ethernet iSCSI Boot. • Sufficient system BIOS memory may not be available to load Intel Ethernet iSCSI Boot. Attempt to disable unused disk controllers and devices in the system BIOS setup menu. SCSI controllers, RAID controller, PXE enabled network connections, and shadowing of system BIOS all reduce the memory area available to Intel Ethernet iSCSI Boot. Disable these devices and reboot the system to see if Intel iSCSI Boot is able to initialize. If disabling the devices in the system BIOS menu does not resolve the problem then attempt to remove unused disk devices or disk controllers from the system. Some system manufacturers allow unused devices to be disabled by jumper settings.
After installing Intel Ethernet iSCSI Boot, the system will not boot to a local disk or network boot device. The system becomes unresponsive after Intel Ethernet iSCSI Boot displays the sign-on banner or after connecting to the iSCSI target.	• A critical system error has occurred during iSCSI Remote Boot initialization. Power on the system and press the 's' key or 'ESC' key before Intel iSCSI Remote Boot initializes. This will bypass the Intel Ethernet iSCSI Boot initialization process and allow the system to boot to a local drive. Use the BootUtil utility to update to the latest version of Intel Ethernet iSCSI Remote Boot. • Updating the system BIOS may also resolve the issue.
"Intel® iSCSI Remote Boot" does not show up as a boot device in the system BIOS boot device menu.	• The system BIOS may not support Intel Ethernet iSCSI Boot. Update the system BIOS with the most recent version available from the system vendor. • A conflict may exist with another installed device. Attempt to disable unused disk and network controllers. Some SCSI and RAID controllers are known to cause compatibility problems with Intel iSCSI Remote Boot.

Error message displayed: "Failed to detect link"	Intel Ethernet iSCSI Boot was unable to detect link on the network port. Check the link detection light on the back of the network connection. The link light should illuminate green when link is established with the link partner. If the link light is illuminated but the error message still displays then attempt to run the Intel link and cable diagnostics tests using diags64e.efi for UEFI or Intel PROSet for Windows.
Error message displayed: "DHCP Server not found!"	iSCSI was configured to retrieve an IP address from DHCP but no DHCP server responded to the DHCP discovery request. This issue can have multiple causes: - The DHCP server may have used up all available IP address reservations. - The client iSCSI system may require static IP address assignment on the connected network. - There may not be a DHCP server present on the network. - Spanning Tree Protocol (STP) on the network switch may be preventing the Intel iSCSI Remote Boot port from contacting the DHCP server. Refer to your network switch documentation on how to disable Spanning Tree Protocol.
Error message displayed: "PnP Check Structure is invalid!"	Intel Ethernet iSCSI Boot was not able to detect a valid PnP PCI BIOS. If this message displays Intel Ethernet iSCSI Boot cannot run on the system in question. A fully PnP compliant PCI BIOS is required to run Intel iSCSI Remote Boot.
Error message displayed: "Invalid iSCSI connection information"	The iSCSI configuration information received from DHCP or statically configured in the setup menu is incomplete and an attempt to login to the iSCSI target system could not be made. Verify that the iSCSI initiator name, iSCSI target name, target IP address, and target port number are configured properly in the iSCSI setup menu (for static configuration) or on the DHCP server (for dynamic BOOTP configuration).
Error message displayed: "Unsupported SCSI disk block size!"	The iSCSI target system is configured to use a disk block size that is not supported by Intel Ethernet iSCSI Boot. Configure the iSCSI target system to use a disk block size of 512 bytes.
Error message displayed: "ERROR: Could not establish TCP/IP connection with iSCSI target system."	Intel Ethernet iSCSI Boot was unable to establish a TCP/IP connection with the iSCSI target system. Verify that the initiator and target IP address, subnet mask, port and gateway settings are configured properly. Verify the settings on the DHCP server if applicable. Check that the iSCSI target system is connected to a network accessible to the Intel iSCSI Remote Boot initiator. Verify that the connection is not being blocked by a firewall.
Error message displayed: "ERROR: CHAP authentication with target failed."	The CHAP user name or secret does not match the CHAP configuration on the iSCSI target system. Verify the CHAP configuration on the Intel iSCSI Remote Boot port matches the iSCSI target system CHAP configuration. Disable CHAP in the iSCSI Remote Boot setup menu if it is not enabled on the target.
Error message displayed: "ERROR: Login request rejected by iSCSI target system."	A login request was sent to the iSCSI target system but the login request was rejected. Verify the iSCSI initiator name, target name, LUN number, and CHAP authentication settings match the settings on the iSCSI target system. Verify that the target is configured to allow the Intel iSCSI Remote Boot initiator access to a LUN.

When installing Linux to NetApp Filer, after a successful target disk discovery, error messages may be seen similar to those listed below. Isctsi-sfnet:hostx: Connect failed with rc -113: No route to host Isctsi-sfnet:hostx: establish_session failed. Could not connect to target	• If these error messages are seen, unused iscsi interfaces on NetApp Filer should be disabled. • Continuous=no should be added to the iscsi.conf file
Error message displayed. "ERROR: iSCSI target not found."	• A TCP/IP connection was successfully made to the target IP address, however an iSCSI target with the specified iSCSI target name could not be found on the target system. Verify that the configured iSCSI target name and initiator name match the settings on the iSCSI target.
Error message displayed. "ERROR: iSCSI target can not accept any more connections."	• The iSCSI target cannot accept any new connections. This error could be caused by a configured limit on the iSCSI target or a limitation of resources (no disks available).
Error message displayed. "ERROR: iSCSI target has reported an error."	• An error has occurred on the iSCSI target. Inspect the iSCSI target to determine the source of the error and ensure it is configured properly.
Error message displayed. ERROR: There is an IP address conflict with another system on the network.	• A system on the network was found using the same IP address as the iSCSI Option ROM client. • If using a static IP address assignment, attempt to change the IP address to something which is not being used by another client on the network. • If using an IP address assigned by a DHCP server, make sure there are no clients on the network which are using an IP address which conflicts with the IP address range used by the DHCP server.

iSCSI Known Issues

A device is not listed in Lifecycle Controller - Network Settings menu

When an Intel® Ethernet iSCSI Boot device is connected to an iSCSI LUN in Legacy BIOS boot mode, the device will not be displayed in the Lifecycle Controller - Network Settings menu.

QoS filter is not created and iSCSI traffic is tagged with priority 0

On devices based on the Intel® Ethernet Controller X710, with NPAR enabled, if you connected the device to a switch that has DCBx enabled, the QoS filter is not created for the NPAR partition and the iSCSI traffic is tagged with priority 0 instead of priority 4. Restarting your system will resolve the issue.

A device cannot be uninstalled if it is configured as an iSCSI primary or secondary port.

Disabling the iSCSI primary port also disables the secondary port. To boot from the secondary port, change it to be the primary port.

iSCSI Remote Boot: Connecting back-to-back to a target with a Broadcom LOM

Connecting an iSCSI boot host to a target through a Broadcom LOM may occasionally cause the connection to fail. Use a switch between the host and target to avoid this.

iSCSI Remote Boot Firmware may show 0.0.0.0 in DHCP server IP address field

In a Linux base DHCP server, the iSCSI Remote Boot firmware shows 0.0.0.0 in the DHCP server IP address field. The iSCSI Remote Boot firmware looks at the DHCP server IP address from the Next-Server field in the DHCP response packet. However, the Linux base DHCP server may not set the field by default. Add "Next-Server <IP Address>," in dhcpd.conf to show the correct DHCP server IP address.

Microsoft Windows iSCSI Boot Issues**Microsoft Initiator does not boot without link on boot port:**

After setting up the system for Intel® Ethernet iSCSI Boot with two ports connected to a target and successfully booting the system, if you later try to boot the system with only the secondary boot port connected to the target, Microsoft Initiator will continuously reboot the system.

To work around this limitation follow these steps:

1. Using Registry Editor, expand the following registry key:

```
\System\CurrentControlSet\Services\Tcpip\Parameters
```

2. Create a DWORD value called DisableDHCPMediaSense and set the value to 0.

Support for Platforms Booted by UEFI iSCSI Native Initiator

Starting with version 2.2.0.0, the iSCSI crash dump driver gained the ability to support platforms booted using the native UEFI iSCSI initiator over supported Intel Network Adapters. This support is available on Microsoft Windows Server operating systems.

Since network adapters on UEFI platforms may not provide legacy iSCSI option ROM, the boot options tab in DMIX may not provide the setting to enable the iSCSI crash dump driver. If this is the case, the following registry entry has to be created:

```
HKLM\SYSTEM\CurrentControlSet\Control\Class\{4D36E97B-E325-11CE-BFC1-08002BE10318}\<InstanceID>\Parameters
DumpMiniport      REG_SZ      iscsdump.sys
```

Moving iSCSI adapter to a different slot:

In a Windows* installation, if you move the iSCSI adapter to a PCI slot other than the one that it occupied when the drivers and MS iSCSI Remote Boot Initiator were installed, then a System Error may occur during the middle of the Windows Splash Screen. This issue goes away if you return the adapter to its original PCI slot. We recommend not moving the adapter used for iSCSI boot installation. This is a known OS issue.

If you have to move the adapter to another slot, then perform the following:

1. Boot the operating system and remove the old adapter
2. Install a new adapter into another slot
3. Setup the new adapter for iSCSI Boot
4. Perform iSCSI boot to the OS via the original adapter
5. Make the new adapter iSCSI-bootable to the OS
6. Reboot
7. Move the old adapter into another slot
8. Repeat steps 2 - 5 for the old adapter you have just moved

Uninstalling Driver can cause blue screen

If the driver for the device in use for iSCSI Boot is uninstalled via Device Manager, Windows will blue screen on reboot and the OS will have to be re-installed. This is a known Windows issue.

Adapters flashed with iSCSI image are not removed from the Device Manager during uninstall

During uninstallation all other Intel Network Connection Software is removed, but drivers for iSCSI Boot adapters that have boot priority.

I/OAT Offload may stop with Intel® Ethernet iSCSI Boot or with Microsoft Initiator installed

A workaround for this issue is to change the following registry value to "0":

HKEY_LOCAL_MACHINE\SYSTEM\CurrentControlSet\Services\IOATDMA\Start

Only change the registry value if iSCSI Boot is enabled and if you want I/OAT offloading. A blue screen will occur if this setting is changed to "0" when iSCSI Boot is not enabled. It must be set back to "3" if iSCSI Boot is disabled or a blue screen will occur on reboot.

Setting LAA (Locally Administered Address) on an iSCSI Boot-Enabled Port Will Cause System Failure on Next Reboot

Do not set LAA on ports with iSCSI Boot enabled.

Intel® Ethernet iSCSI Boot version does not match between displayed versions on DMIX and the scrolling text during boot

If a device is not set to primary but is enumerated first, the BIOS will still use that device's version of iSCSI Boot. Therefore the user may end up using an earlier version of Intel® Ethernet iSCSI Boot than expected. The solution is that all devices in the system must have the same version of iSCSI Boot. To do this the user should go to the Boot Options Tab and update the devices' flash to the latest version.

IPv6 iSCSI login to Dell EqualLogic arrays using jumbo frames

To establish an iSCSI session using IPv6 and jumbo frames with Dell EqualLogic arrays, TCP/UDP checksum offloads on the Intel iSCSI adapter should be disabled.

Microsoft Windows iSCSI/DCB Known Issues

Automatic creation of iSCSI traffic filters for DCB is only supported on networks which make use of IPv4 addressing

The iSCSI for Data Center Bridging (DCB) feature uses Quality of Service (QOS) traffic filters to tag outgoing packets with a priority. The Intel iSCSI Agent dynamically creates these traffic filters as needed on networks using IPv4 addressing.

IPv6 iSCSI login to Dell EqualLogic arrays using jumbo frames

To establish an iSCSI session using IPv6 and jumbo frames with Dell EqualLogic arrays, TCP/UDP checksum offloads on the Intel iSCSI adapter should be disabled.

Linux Known Issues

Channel Bonding

Linux Channel Bonding has basic compatibility issues with iSCSI Boot and should not be used.

LRO and iSCSI Incompatibility

LRO (Large Receive Offload) is incompatible with iSCSI target or initiator traffic. A panic may occur when iSCSI traffic is received through the ixgbe driver with LRO enabled. The driver should be built and installed with:

```
# make CFLAGS_EXTRA=-DIXGBE_NO_LRO install
```

Firmware

Firmware is a layer of software that is programmed into a device's memory. It provides low level functionality for the device. In most cases you will not notice the firmware on your device at all. Firmware error states usually occur because of an unsuccessful update.

Firmware Security

Intel or your equipment manufacturer will occasionally release a firmware security patch. We recommend that you update your firmware to the latest version available for your device to take advantage of these security patches. Firmware updates for Intel Ethernet devices will have a Security Revision number (SRev).

Minimum Security Revision Enforcement

Firmware security updates can be undone if you install a previous version of the firmware onto your device. Intel firmware releases include a Minimum Security Revision (MinSRev) enforcement feature. This means you can block someone from installing a lower revision of the firmware onto your device. This will limit the rollback capabilities of your device. The firmware update process will block the update if the supplied firmware has a lower security revision (SRev) than the MinSRev value of the firmware currently loaded on the device. Only update the MinSRev value if you are certain you will not need to roll the firmware back to an earlier version.

You can update the MinSRev value during the firmware update process, locking the current security version in as the new MinSRev baseline, by using the `-optinminsrev` command line option.



CAUTION: The MinSRev value on a device can never be decreased. Once the MinSRev is increased, NVM downgrades attempting to install a lower Security revision (SRev) than the current MinSRev will be rejected by the device. Users who want to downgrade firmware without regard to security revisions should not use this feature.

This applies to devices based on the following controllers:

- Intel® Ethernet Controller 800 Series
- Intel® Ethernet Controller 700 Series
- Intel® Ethernet Controller X550

Examples

View your device's current Minimum Security Revision

Use these steps to view the Minimum Security Revision that is set on the card.

1. During system boot, press the F2 key to enter the **System Setup** menu.
2. Under **System Setup Main Menu**, select **Device Settings**.
3. Select your adapter from the list to get to the **Device Configuration Menu**.
4. Under **Main Configuration Page**, select **Firmware Image Properties**.
5. View the **Minimum Security Revision** attribute.

Update your device's MinSRev

To update your device's MinSRev, you must first extract `nvupdate` from the FW DUP. In Windows, this can be done by clicking "Extract" after starting the firmware DUP in GUI mode. In Linux, this can be done by using the `--extract <path>` parameters from the command line. After you extract the package, perform the update from the command line as shown in the following example:

Windows: `nvupdatew64e -u -optinminsrev -l update.log -o update.xml -c nvupdate.cfg`

Linux: `nvupdate64e -u -optinminsrev -l update.log -o update.xml -c nvupdate.cfg`

Where:

-u -- Sets nvupdate to update mode.

-optiminsrev -- Tells the tool to update the MinSRev value.

-l update.log -- Specifies the name of the log file.

-o update.xml -- Specifies the name of the results file. This is an XML file that contains the inventory/update results.

-c nvupdate.cfg -- Specifies the name of the configuration file. This is a text file that contains descriptions of networking devices and firmware versions for those devices.

Firmware Rollback Mode

When a port is in firmware rollback mode it may have reduced functionality. Usually a device enters firmware rollback mode when a firmware update does not complete correctly. Rebooting or power cycling the system may allow the port to use the previous firmware. You may need to reapply the firmware update to regain full functionality on the device. Use the appropriate NVM Update Package to update the device's firmware. Download the latest NVM Update Package from your vendor's support website and follow the instructions in it. After restoring the NVM image, you may need to perform an A/C power cycle of the system.

Firmware Recovery Mode

A device will enter Firmware Recovery mode if it detects a problem that requires the firmware to be reprogrammed. When a device is in Firmware Recovery mode it will not pass traffic or allow any configuration; you can only attempt to recover the device's firmware.

Affected Products

Ethernet Device	New NVM Version	Software Driver and Tools
Intel® Ethernet Controller 800 Series	All firmware versions	All driver versions
Intel® Ethernet Controller 700 Series	Intel® NIC Family Version 18.8.0 Firmware and newer	Intel® NIC Family Version 18.8.0 and newer
Intel® Ethernet Controller X550	Intel® NIC Family Version 18.8.0 Firmware and newer	Intel® NIC Family Version 18.8.0 and newer

Recovery Mode Detection

During initialization, a device can enter recovery mode if the device firmware detects a problem with the LAN device, mandating NVM reprogramming to restore normal operation. After thorough internal testing of the NVM (typically less than 10 minutes, but in some cases longer), the NIC enters Recovery Mode.

Firmware Recovery Mode Errors and Messages

When a device is in Firmware Recovery mode, the device drivers, preboot software, and utilities may log or display messages such as the following:

- **Firmware recovery mode detected. Limiting functionality. Refer to the Intel® Ethernet Adapters and Devices User Guide for details on firmware recovery mode.**
- **Firmware recovery mode detected. The underlying hardware has been deactivated. Refer to the Intel® Ethernet Adapters and Devices User Guide for details on firmware recovery mode.**
- **Firmware recovery mode detected. Initialization failed.**
- **Firmware recovery mode detected. Limiting functionality.**
- **Initialization failure due to repeated FW resets.** This message is usually an indication that the device is about to enter Recovery Mode. The device may be able to return to normal functionality without intervention. This may take several minutes. No action is required unless the device does enter Recovery Mode.

Resolving Firmware Recovery Mode Issues

If your device is in Firmware Recovery mode you can restore it to factory defaults using the process for resolution of Firmware Recovery Mode Issues as outlined in the sub-sections below.

NVM Self Check

The process begins after power-on or reboot. At this time, the firmware will perform tests to assess whether there is damage or corruption of the device NVM image.

Actions:

- If NVM image damage or corruption **is not** detected, the device will initialize and operate normally. No further action is required.
- If NVM image damage or corruption **is** detected, the device will not initialize. Proceed with the additional recovery steps listed under Recovery Mode below.

Recovery Mode

The device NVM image has exhibited damage or corruption.

Actions:

1. Wait 10 minutes for the NVM self-check process to complete. If during this period normal operation is achieved, the device will initialize and operate normally. No further action is required.
2. If after 10 minutes normal operation is *not* achieved:
 - a. Check the System Event log for Windows OSs or driver message and kernel logs for Linux and ESXi based distributions. Recovery Mode is confirmed by presence of message/log entries as listed in the Firmware Recovery Mode Errors and Messages section above.
 - b. Reboot the system and proceed with the additional recovery steps listed under NVM Image Restoration below.



NOTES:

- While in Recovery Mode, for Windows OSs, clicking on the adapter in device manager may present a dialog box indicating that Firmware Recovery Mode is active.
 - Once the dialog is dismissed, while the device appears to be functioning normally, it is in fact limited to only enable NVM image recovery.
- If the system is rebooted (versus power cycled), the driver status may not show a Code 10/yellow bang status as expected. Refer to events logged in System Event log for Windows OSs or driver message and kernel logs for Linux and ESXi based distributions to accurately assess the adapter status.
- When the adapter is in recovery mode, the link LED will not be lit and the adapter will not appear in the following locations:
 - F2 System Setup > Device Settings
 - System BIOS as a NIC for PXE Boot in UEFI boot mode
 - Lifecycle Controller > Network Settings
 - iDRAC Web GUI > Firmware Inventory

NVM Image Restoration

At this point, the device is in Firmware Recovery mode and its functionality is limited to only supporting restoration of the NVM image.

Actions:

1. Before initiating device recovery, the integrity of the host operating system, device drivers and firmware utilities must be verified and reinstalled if necessary. Fully functional operating system, device drivers and tools are required for device recovery. Please consult your operating system specific instructions on how to scan and repair potentially damaged system files.
2. If your device is in Firmware Recovery mode you can restore it to factory defaults using the latest Dell EMC Update Package for Intel Adapter Firmware (FW-DUP) or Intel NIC Family ESXi Firmware Update Package. Download the latest Dell EMC Update Package for Intel Adapter Firmware (FW-DUP) or Intel NIC Family ESXi Firmware Update Package from Dell's support website and follow the instructions in them. The Dell EMC Update Package for Intel Adapter Firmware (FW-DUP) must be executed in an operating system to recover the device, not in the Dell Lifecycle Controller or iDRAC.
3. After restoring the NVM image, perform an A/C power cycle of the system. Details for this are in the **Other General Notes** section below.

**NOTES:**

- If a device is in recovery mode when the Dell DUP package is executed for inventory, the Firmware Family Version (FFV) will display "0.0.0". This is expected behavior.
- Running the FW-DUP in recovery mode does not update the option ROM. A/C power cycling and running the FW-DUP a second time will correct this.
- After running the FW-DUP in recovery mode, the firmware version is incorrect. Updating the firmware via the Dell Lifecycle Controller or iDRAC resolves the issue.
- User configured settings (i.e. iSCSI target information, user defined port/alternate MAC addresses) will not be restored to pre-recovery mode values.

Other General Notes**NOTES:**

- To perform an AC power cycle, do the following:
 - Shut down the system if its is powered up.
 - Unplug all AC power cords from the system.
 - Leave the AC power cords unplugged for 15 seconds to allow the system power supply to discharge completely.
 - Plug in AC power cords to the system.
- If NPAR was enabled when the system was recovered, you may continue to see NIC partitions in the OS even though the HII reports that NPAR is disabled. In this situation, you may also see that some of the partitions do not function properly in the OS. To fix this, re-enable NPAR in HII, as follows:
 1. During system boot, press F2 to enter system setup, then select Device Settings and select the desired device.
 2. Select Device Level Configuration, then turn NPAR on in the Virtualization Mode menu.
 3. Save changes, which will cause the system to restart.

Troubleshooting

Common Problems and Solutions

There are many simple, easy-to-fix problems related to network problems. Review each one of these before going further.

- Check for recent changes to hardware, software, or the network that may have disrupted communications.
- Check the driver software.
 - Make sure you are using the latest appropriate drivers for your adapter from the [Dell support website](#).
 - Disable (or unload), then re-enable (reload) the driver or adapter.
 - Check for conflicting settings. Disable advanced settings to see if it corrects the problem.
 - Re-install the drivers.
- Check the cable. Use the best available cabling for the intended data rate.
 - Check that the cable is securely attached at both points.
 - Make sure the cable length does not exceed specifications.
 - Perform a cable test.
 - Replace the cable.
- Check the link partner (switch, hub, etc.).
 - Make sure the link partner is active and can send and receive traffic.
 - Make sure the adapter and link partner settings match one another, or are set to auto-negotiate.
 - Make sure the port is enabled.
 - Re-connect to another available port or another link partner.
- Look for adapter hardware problems.
 - Re-seat the adapter.
 - Insert the adapter in another slot.
 - Check for conflicting or incompatible hardware devices and settings.
 - Replace the adapter.
- Check the [Dell support website](#) for possible documented issues.
 - Select your adapter from the adapter family list.
 - Check the Frequently Asked questions section.
 - Check the Knowledge Base.
- Check your process monitor and other system monitors.
 - Check to see that there is sufficient processor and memory capacity to perform networking activity.
 - Look for any unusual activity (or lack of activity).
 - Use network testing programs to check for basic connectivity.
- Check your BIOS version and settings.
 - Use the latest appropriate BIOS for your computer.
 - Make sure the settings are appropriate for your computer.

The following troubleshooting table assumes that you have already reviewed the common problems and solutions.

Problem	Solution
Your computer cannot find the adapter	Make sure your adapter slots are compatible for the type of adapter you are using.
Diagnostics pass but the connection fails	Make sure the cable is securely attached, is the proper type and does not exceed the recommended lengths. Make sure the duplex mode and speed setting on the adapter matches the setting on the switch.
Adapter unable to connect to switch at correct speed. Gigabit adapter connects at 100 Mbps and 10 gigabit adapter connects at 1000 Mbps.	<i>This is applicable only to copper-based connections.</i> Make sure the adapter and the link partner are set to auto-negotiate.

Problem	Solution
	Verify that you are running the latest operating system revision for your switch and that the switch is compliant with the proper IEEE standard: • IEEE 802.3ad-compliant (gigabit over copper) • IEEE 802.3an-compliant (10 gigabit over copper)
The device does not connect at the expected speed.	When Gigabit Master/Slave mode is forced to "master" mode on both the Intel adapter and its link partner, the link speed obtained by the Intel adapter may be lower than expected.
The adapter stops working without apparent cause	Run the adapter and network tests described under "Test the Adapter".
The Link indicator light is off	Run the adapter and network tests described under "Test the Adapter". Make sure the proper (and latest) driver is loaded. Make sure that the link partner is configured to auto-negotiate (or forced to match adapter) Verify that the switch is IEEE 802.3ad-compliant.
The link light is on, but communications are not properly established	Make sure the proper (and latest) driver is loaded. Both the adapter and its link partner must be set to either auto-detect or manually set to the same speed and duplex settings.  NOTE: The adapter's link indicator light may be on even if communications between the adapter and its link partner have not been properly established. Technically, the link indicator light represents the presence of a carrier signal but not necessarily the ability to properly communicate with a link partner. This is expected behavior and is consistent with IEEE's specification for physical layer operation.
RX or TX light is off	Network may be idle; try creating traffic while monitoring the lights.
The diagnostic utility reports the adapter is "Not enabled by BIOS"	The PCI BIOS isn't configuring the adapter correctly. See PCI / PCI-X / PCI Express Configuration.
The computer hangs when the drivers are loaded	Try changing the PCI BIOS interrupt settings. See PCI / PCI-X / PCI Express Configuration.
PCI / PCI-X / PCI Express Configuration	If the adapter is not recognized by your OS or if it does not work you may need to change some BIOS settings. Try the following only if you are having problems with the adapter and are familiar with BIOS settings. • Check to see that the "Plug-and-Play" setting is compatible with the operating system you are using. • Make sure the slot is enabled. • Install the adapter in a bus-master slot. • Configure interrupts for level-triggering, as opposed to edge-triggering. • Reserve interrupts and/or memory addresses. This prevents multiple buses or bus slots from using the same interrupts. Check the BIOS for IRQ options for PCI / PCI-X / PCIe.
Driver message: "Rx/Tx is disabled on this device because an unsupported SFP+ module type was detected."	You installed an unsupported module in the device. See Supported SFP+ and QSFP+ Modules for a list of supported modules.

Multiple Adapters

When configuring a multi-adapter environment, you must upgrade all Intel adapters in the computer to the latest software.

If the computer has trouble detecting all adapters, consider the following:

- If you enable Wake on LAN* (WoL) on more than two adapters, the Wake on LAN feature may overdraw your system's auxiliary power supply, resulting in the inability to boot the system and other unpredictable problems. For multiple desktop/management adapters, it is recommended that you install one adapter at a time and use the IBAUtil utility (ibautil.exe in \APPS\BOOTAGNT) to disable the WoL feature on adapters that do not require WoL capabilities. On server adapters, the WoL feature is disabled by default.
- Adapters with Intel Boot Agent enabled will require a portion of the limited start up memory for each adapter enabled. Disable the service on adapters that do not need to boot Pre-Boot Execution Environment (PXE).

Safe Mode

Adapters based on the Intel® Ethernet Controller 800 Series require a [Dynamic Device Personalization \(DDP\)](#) package file to enable advanced and performance features. If the driver detects a missing or incompatible DDP package file, the driver will go into Safe Mode. Safe Mode supports only basic traffic and minimal functionality, such as updating the NVM or downloading a new driver or DDP package.



NOTES:

- Safe Mode only applies to the affected physical function and does not impact any other PFs.
- [Firmware Recovery Mode](#) takes precedence over Safe Mode.

Safe Mode Errors and Messages

When the driver is in Safe Mode, the device drivers and utilities may log or display messages to help with troubleshooting. The following conditions may cause the driver to enter Safe Mode:

- The DDP package file was not found or couldn't be read.
- The DDP package file's version number, signature, or other metadata aren't valid or aren't supported by the driver.
- An unknown error occurred when loading the DDP package.
- The driver couldn't load the DDP package file because a compatible DDP package is already present on the device.
- The device has a DDP package that isn't supported by the driver.

Resolving Safe Mode Issues

The device drivers and utilities may display the action to take to get out of Safe Mode, depending on the underlying cause. Possible actions could include the following:

- Wait for the device to reset.
- Install the latest driver.
- Download a new DDP package.
- Restart the adapter. If the problem persists, install the latest driver.
- Reboot the system. If the problem persists, update the NVM.

You can download the latest drivers and DDP packages from the [Dell support website](#).

PF Message Queue Overflow

The device driver can detect some types of anomalous behavior. When it does, it will log the VF MAC address and associated PF MAC address. Using this information, you can check the virtual machine (VM) that is using the VF MAC address to ensure that the VM is operating correctly.

Possible Misconfiguration of the Ethernet Port

You may see an informational message stating that a potential misconfiguration of the Ethernet port was detected. This is to alert you that your device is being underutilized. If this was intentional, you may ignore this message. For example, setting your Intel® Ethernet 100G 2P E810-C adapter to 2x2x25 is valid, but it does not use the full capabilities of the device. If you see this message, and the configuration was not intentional, you may use the Ethernet Port Configuration Tool (EPCT) to correct the configuration.

Other Performance Issues

Attaining the best speed requires that many components are operating at peak efficiency. Among them are the following:

- **Cable quality and length** - Do not exceed the maximum recommended length for your cable type. Shorter lengths often provide better results. Check for loose or damaged connectors. Check the cable for kinked or damaged sections.
- **Bus speed and traffic** - The PCI bus speed accommodates the slowest PCI card installed. Check to see if you have a card that is slowing down your system.
- **Processor and Memory** - Check your performance monitoring programs to see if traffic is being affected by your processor speed, available memory or other processes.
- **Transmission frame size** - Your network performance may be enhanced by adjusting or maximizing the transmission frame size. Operating systems, switches and adapters will impose varying limits on maximum frame size. See the discussion on Jumbo Frames for your OS.
- **Operating System** - Networking feature implementation will vary by operating system version, such as offloading and multiprocessor threading.

Testing the Adapter

Intel's diagnostic software lets you test the adapter to see if there are problems with the adapter hardware, the cabling, or the network connection.

Testing from Windows

Intel PROSet allows you to run three types of diagnostic tests:

- **Connection Test:** Verifies network connectivity by pinging the DHCP server, WINS server, and gateway.
- **Cable Tests:** Provides information about cable properties.



NOTE: The Cable Test is not supported on all adapters and will not run on Direct Attached Cables (DAC) or Fiber. The Cable Test will only be available on adapters that support it.

- **Hardware Tests:** Determines if the adapter is functioning properly.



NOTE: Hardware tests will fail if the adapter is configured for iSCSI Boot.

To access these tests, select the adapter in Windows Device Manager, click the **Link** tab, and click **Diagnostics**. A Diagnostics window displays tabs for each type of test. Click the appropriate tab and run the test.

In Intel® PROSet ACU, use the Diagnostics panel.

The availability of these tests is dependent on the adapter and operating system. Tests may be disabled if:

- iSCSI Boot is enabled on the port.
- The port is used as a manageability port.
- The tests are being run from a virtual machine.



NOTE: At this time, Windows diagnostics are not supported on ports based on an Intel Ethernet Controller I225 and will fail.

Testing from Windows PowerShell*

Intel provides two [PowerShell cmdlets](#) for testing your adapter.

- Test-IntelNetDiagnostics runs the specified test suite on the specified device. See the Test-IntelNetDiagnostics help inside PowerShell for more information.
- Test-IntelNetIdentifyAdapter blinks the LED on the specified device.

Linux Diagnostics

The driver utilizes the ethtool interface for driver configuration and diagnostics, as well as displaying statistical information. ethtool version 1.6 or later is required for this functionality.

The latest release of ethtool can be found at: <http://sourceforge.net/projects/gkernel>.



NOTE: ethtool 1.6 only supports a limited set of ethtool options. Support for a more complete ethtool feature set can be enabled by upgrading ethtool to the latest version.

Windows* Event Log

Windows Event Log Service Names

Intel® Ethernet Controller	NDIS Driver File Names	Windows Event Log Service Name
I350	E1r*.sys	e1repress
I354	E1s*.sys	e1sexpress
X520	Ixn*.sys	ixgbn
X540	Ixt*.sys	ixgbt
X550	Ixs*.sys	ixgbs
710 Series	I40ea*.sys	i40ea
810 Series	icea.sys	icea

Intel® Network Adapter Messages

Below is a list of custom event messages that appear in the Windows Event Log for Intel® Ethernet adapters:

Event ID	Message	Severity
1	The Hyper-V role was disabled on the system. All Intel® Ethernet devices configured with a Virtualization performance profile were changed to a more appropriate performance profile.	Informational
6	PROBLEM: Unable to allocate the map registers necessary for operation. ACTION: Reduce the number of transmit descriptors and restart.	Error
7	PROBLEM: Could not assign an interrupt for the network adapter. ACTION: Try a different PCIe slot. ACTION: Install the latest driver from http://www.intel.com/support/go/network/adapter/home.htm .	Error
23	PROBLEM: The EEPROM on the network adapter may be corrupt. ACTION: Visit the support web site at http://www.intel.com/support/go/network/adapter/home.htm .	Error
24	PROBLEM: Unable to start the network adapter. ACTION: Install the latest driver from http://www.intel.com/support/go/network/adapter/home.htm .	Error

Event ID	Message	Severity
25	PROBLEM: The MAC address on the network adapter is invalid. ACTION: Visit http://www.intel.com/support/go/network/adapter/home.htm for assistance.	Error
27	Network link has been disconnected.	Warning
30	PROBLEM: The network adapter is configured for auto-negotiation but the link partner is not. This may result in a duplex mismatch. ACTION: Configure the link partner for auto-negotiation.	Warning
31	Network link has been established at 10 Gbps full duplex.	Informational
32	Network link has been established at 1 Gbps full duplex.	Informational
33	Network link has been established at 100 Mbps full duplex.	Informational
34	Network link has been established at 100 Mbps half duplex.	Informational
35	Network link has been established at 10 Mbps full duplex.	Informational
36	Network link has been established at 10 Mbps half duplex.	Informational
37	PROBLEM: PCI Express bandwidth available for this adapter is not sufficient for optimal performance. ACTION: Move the adapter to a x8 PCI Express slot.	Warning
40	Intel Smart Speed has downgraded the link speed from the maximum advertised.	Informational
41	The network adapter driver has been stopped.	Informational
42	The network adapter driver has been started.	Informational
43	PROBLEM: Could not allocate shared memory necessary for operation. ACTION: Reduce the number of transmit and receive descriptors, then restart.	Error
44	PROBLEM: Could not allocate memory necessary for operation. ACTION: Reduce the number of transmit and receive descriptors, then restart.	Error
45	PROBLEM: Could not allocate a resource pool necessary for operation. ACTION: Reduce the number of transmit and receive descriptors, then restart.	Error
46	PROBLEM: Could not initialize scatter-gather DMA resources necessary for operation. ACTION: Reduce the number of transmit descriptors and restart.	Error
47	PROBLEM: Could not map the network adapter flash. ACTION: Install the latest driver from http://www.intel.com/support/go/network/adapter/home.htm . ACTION: Try another slot.	Error
48	PROBLEM: The fan on the network adapter has failed. ACTION: Power off the machine and replace the network adapter.	Error
49	PROBLEM: The driver was unable to load due to an unsupported SFP+ module installed in the adapter. ACTION: Replace the module. ACTION: Install the latest driver from http://www.intel.com/support/go/network/adapter/home.htm .	Error
50	PROBLEM: The network adapter has been stopped because it has overheated. ACTION: Restart the computer. If the problem persists, power off the computer and replace the network adapter.	Error
51	PROBLEM: The network adapter link speed was downshifted because it overheated.	Error

Event ID	Message	Severity
52	PROBLEM: The network adapter has been stopped because it has overheated.	Error
53	Jumbo Frames cannot be configured when MACSec is enabled.	Informational
54	PROBLEM: A malicious VF driver has been detected.	Warning
56	The network driver has been stopped because the network adapter has been removed.	Informational
58	Network link has been established at 25Gbps full duplex.	Informational
60	Network link has been established at 50Gbps full duplex.	Informational
61	Network link has been established at 20Gbps full duplex.	Informational
64	This network adapter's etrack ID is:	Informational
65	PROBLEM: PCI Express bandwidth available for this adapter is not sufficient for optimal performance. ACTION: Move the adapter to a Generation 3 x4 PCI Express slot.	Warning
66	PROBLEM: PCI Express bandwidth available for this adapter is not sufficient for optimal performance. ACTION: Move the adapter to a Generation 3 x8 PCI Express slot.	Warning
67	The partition detected link speed that is less than 10Gbps.	Warning
68	The driver for the device stopped because the NVM image is newer than the driver. You must install the most recent version of the network driver.	Error
69	The driver for the device detected a newer version of the NVM image than expected. Please install the most recent version of the network driver.	Warning
70	The driver for the device detected an older version of the NVM image than expected. Please update the NVM image.	Informational
71	The driver failed to load because an unsupported module type was detected.	Error
72	PROBLEM: The driver failed to load because the adapter was not provided MSI-X interrupt resources. ACTION: Move the adapter to another slot or platform.	Error
73	The 'Speed and Duplex' and 'Flow Control' user settings cannot be changed since this device is operating in virtual connect mode.	Informational

Intel DCB Messages

Below is a list of intermediate driver custom event messages that appear in the Windows Event Log:

Event ID	Message	Severity
256	Service debug string	Informational
257	Enhanced Transmission Selection feature has been enabled on a device.	Informational
258	Enhanced Transmission Selection feature has been disabled on a device.	Informational
259	Priority Flow Control feature has been enabled on a device.	Informational
260	Priority Flow Control feature has been disabled on a device.	Informational

Event ID	Message	Severity
261	Enhanced Transmission Selection feature on a device has changed to operational.	Informational
262	Priority Flow Control feature on a device has changed to operational.	Informational
263	Application feature on a device has changed to operational.	Informational
264	Application feature has been disabled on a device.	Informational
265	Application feature has been enabled on a device.	Informational
269	Logical Link feature on a device has changed to operational.	Informational
270	Logical Link feature has been disabled on a device.	Informational
271	Logical Link feature has been enabled on a device.	Informational
768	Service failed while starting.	Error
770	Service handler failed while installing.	Error
771	Service could not allocate sufficient memory.	Error
772	Service unable to use network adapter.	Error
773	Service rejected configuration - invalid total for transmit bandwidth groups.	Error
774	Service rejected configuration - invalid total for receive bandwidth groups.	Error
775	Service rejected configuration - invalid transmit bandwidth group index.	Error
776	Service rejected configuration - invalid receive bandwidth group index.	Error
777	Service rejected configuration - link strict and non-zero bandwidth on transmit traffic class.	Error
778	Service rejected configuration - link strict and non-zero bandwidth on receive traffic class.	Error
779	Service rejected configuration - zero bandwidth on transmit traffic class.	Error
780	Service rejected configuration - zero bandwidth on receive traffic class.	Error
781	Service rejected configuration - link strict and non-zero bandwidth on transmit bandwidth group.	Error
782	Service rejected configuration - link strict and non-zero bandwidth on receive bandwidth group.	Error
783	Service rejected configuration - invalid total transmit for bandwidth group.	Error
784	Service rejected configuration - invalid total receive for bandwidth group.	Error
785	Service unable to configure needed WMI services.	Error
786	Service experienced a transmit state machine error.	Error
787	Service experienced a receive state machine error.	Error
789	Service connection to LLDP protocol driver failed.	Error
790	Enhanced Transmission Selection feature on a device has changed to non-operational.	Error
791	Priority Flow Control feature on a device has changed to non-operational.	Error
792	Application feature on a device has changed to non-operational.	Error
793	Service rejected configuration - multiple link strict bandwidth groups were detected.	Error

Event ID	Message	Severity
794	Logical Link feature on a device has changed to non-operational.	Error
795	Failed to open device.	Error
796	DCB settings of the network adapter are invalid.	Error
797	DCB settings of the network adapter are invalid - AppSelector.	Error
798	Detected a non-optimal network adapter driver component. Please install network adapter driver version 3.5 or greater.	Error

Intel iSCSI DCB Messages

Below is a list of intermediate driver custom event messages that appear in the Windows Event Log:

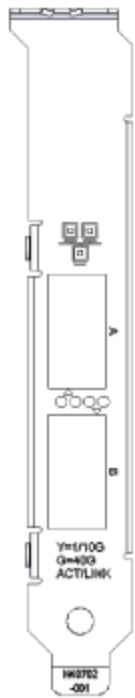
Event ID	Message	Severity
4352	Service debug string:	Informational
4353	iSCSI DCB Agent has added a QOS filter for iSCSI traffic.	Informational
4354	iSCSI DCB Agent has removed a QOS filter for iSCSI traffic.	Informational
4355	iSCSI DCB Agent has modified a QOS filter for iSCSI traffic.	Informational
4356	iSCSI DCB Agent was notified by the QOS service that an iSCSI DCB adapter was closed.	Informational
4357	Priority Flow Control and Application User Priority are configured for iSCSI DCB traffic.	Informational
8704	Some members of the Team configured for iSCSI DCB traffic have an invalid DCB configuration.	Warning
13056	Service failed while starting.	Error
13057	Service handler failed while installing.	Error
13058	Error returned by Traffic Control interface.	Error
13059	Service could not allocate sufficient memory.	Error
13060	iSCSI DCB Agent is unable to add the QOS filter for iSCSI traffic.	Error
13061	iSCSI DCB Agent was notified by the QOS service that all QOS filters for an iSCSI DCB adapter were removed.	Error
13062	Application User Priority or Priority Flow Control is misconfigured for iSCSI DCB traffic.	Error
13063	Priority Flow Control TLV is non-operational for iSCSI DCB traffic.	Error
13064	Application TLV is non-operational for iSCSI DCB traffic.	Error
13065	Detected unsupported Operating System.	Error
13066	No member of the Team configured for iSCSI DCB traffic has a valid DCB configuration.	Error

Indicator Lights

The Intel Server and Desktop network adapters feature indicator lights on the adapter backplate that serve to indicate activity and the status of the adapter board. The following tables define the meaning for the possible states of the indicator lights for each adapter board.

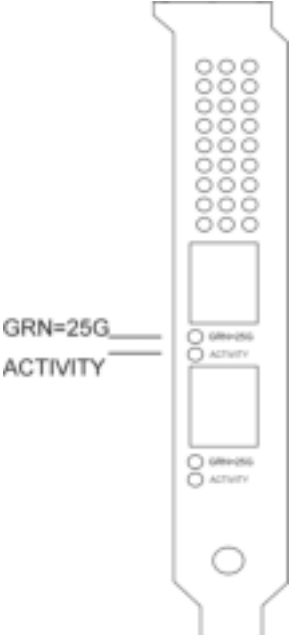
Dual Port QSFP+ Adapters

The **Intel® Ethernet Converged Network Adapter XL710-Q2** has the following indicator lights:

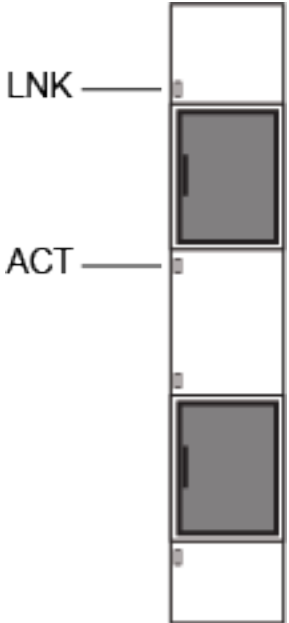
 The diagram shows the backplate of the Intel Ethernet Converged Network Adapter XL710-Q2. It features two QSFP+ ports at the top, a central rectangular area with a small circle labeled 'ACT/LNK', and a bottom section with a circular indicator light. Text on the backplate includes 'Y=1/10G', 'G=40G', 'ACT/LNK', and 'INTEL -001'.	Label	Indication	Meaning
	ACT/LNK	Green	Linked at 40 Gb
		Blinking On/OFF	Actively transmitting or receiving data
		Off	No link.

Dual Port SFP28 Adapters

The **Intel® Ethernet 25G 2P E810-XXV Adapter** and **Intel® Ethernet 25G 2P XXV710 Adapter** have the following indicator lights:

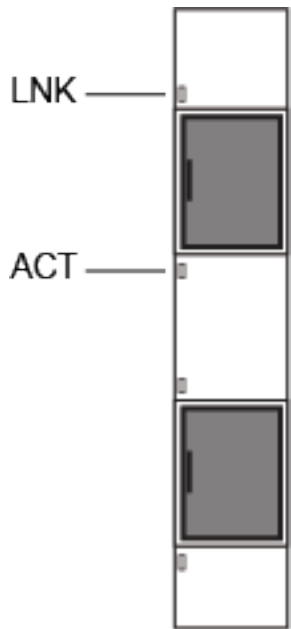
	Label	Indication	Meaning
	GRN 25G	Green	Linked at maximum port speed
		Yellow	Linked at less than maximum port speed
	ACTIVITY	Blinking On/OFF	Actively transmitting or receiving data
		Off	No link

The **Intel® Ethernet 25G 2P E810-XXV OCP** has the following indicator lights:

	Label	Indication	Meaning
	LNK	Green	Operating at maximum port speed
		Yellow	Linked at less than maximum port speed
	ACT	Green flashing	Data activity
		Off	No activity

Dual Port SFP/SFP+ Adapters

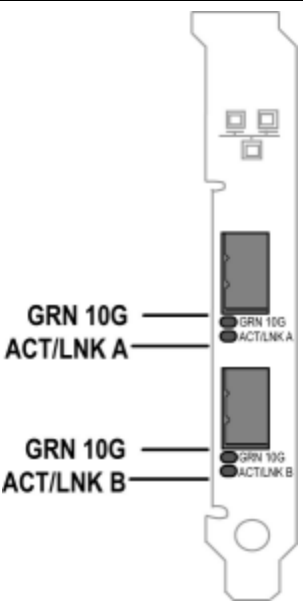
The **Intel® Ethernet 10G 2P X710 OCP** has the following indicator lights:

 LNK ACT	Label	Indication	Meaning
	LNK	Green	Operating at maximum port speed
		Yellow	Linked at less than maximum port speed
	ACT	Green flashing	Data activity
		Off	No activity

The **Intel® Ethernet Converged Network Adapter X710** has the following indicator lights:

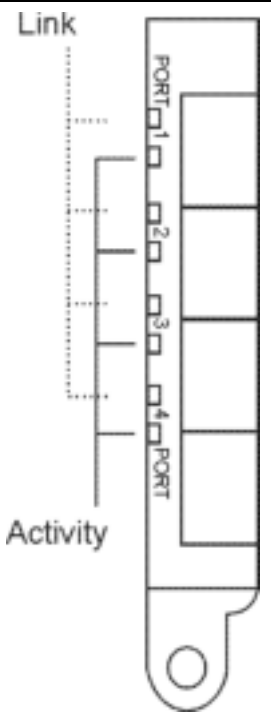
 LNK ACT	Label	Indication	Meaning
	LNK	Green	Linked at 10 Gb
		Yellow	Linked at 1 Gb
	ACT	Blinking On/OFF	Actively transmitting or receiving data
		Off	No link

The **Intel® 10G 2P X520 Adapter** has the following indicator lights:

	Label	Indication	Meaning
	GRN 10G (A or B): Green	On	Linked to the LAN
		Off	Not linked to the LAN
	ACT/LNK (A or B): Green	Blinking On/Off	Actively transmitting or receiving data
		Off	No link

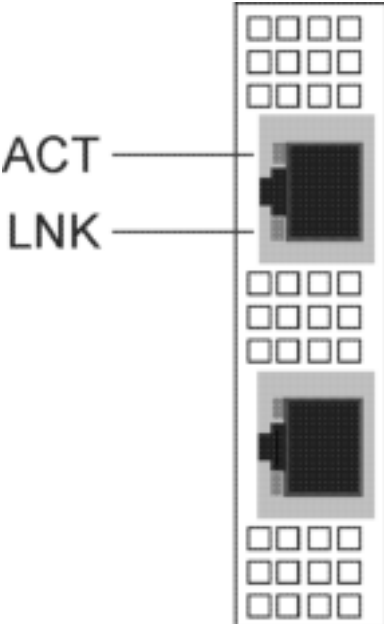
Quad Port SFP/SFP+ Adapters

The **Intel® Ethernet 10G 4P X710 OCP** has the following indicator lights:

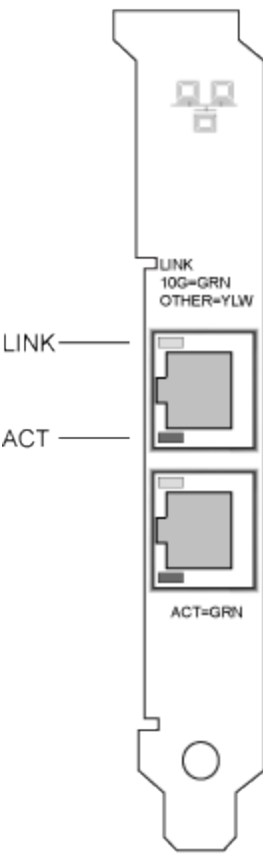
	Label	Indication	Meaning
	Link	Green	Linked at maximum port speed
		Yellow	Linked at less than maximum port speed
	ACTIVITY	Blinking On/OFF	Actively transmitting or receiving data
		Off	No link.

Dual Port Copper Adapters

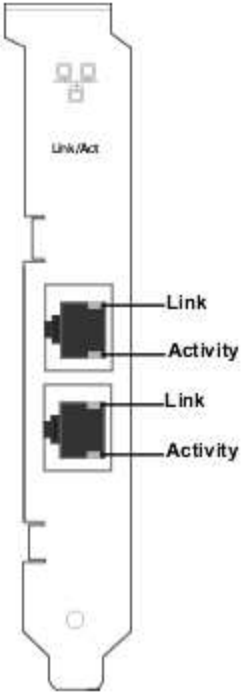
The **Intel® Ethernet 10G 2P X710-T2L-t OCP** has the following indicator lights:

	Label	Indication	Meaning
	Link	Green	Linked at 10 Gbps.
		Yellow	Linked at slower than 10 Gbps.
		Off	No link.
	Activity	Blinking On/Off	Actively transmitting or receiving data.
		Off	No link.

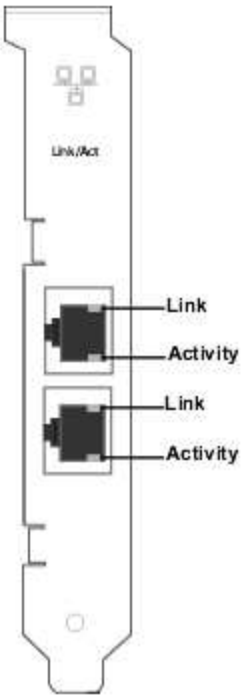
The **Intel® Ethernet 10G 2P X710-T2L-t Adapter** has the following indicator lights:

	Label	Indication	Meaning
	Link	Green	Linked at 10 Gbps.
		Yellow	Linked at slower than 10 Gbps.
		Off	No link.
	Activity	Blinking On/Off	Actively transmitting or receiving data.
		Off	No link.

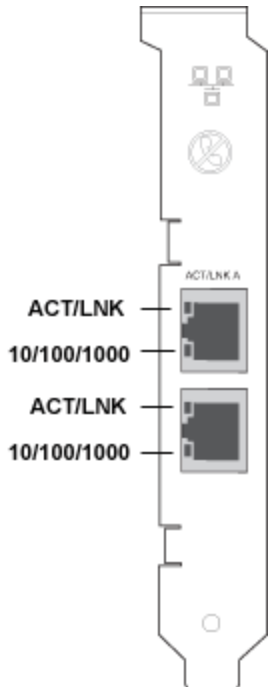
The **Intel® Ethernet 10G 2P X550-t Adapter** has the following indicator lights:

	Label	Indication	Meaning
	Link	Green	Linked at 10 Gbps.
		Yellow	Linked at less than 10 Gbps.
		Off	No link.
	Activity	Blinking On/Off	Actively transmitting or receiving data.
		Off	No link.

The **Intel® Ethernet 10G 2P X540-t Adapter** has the following indicator lights:

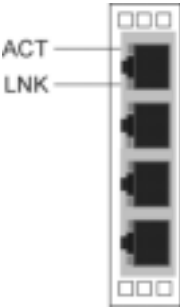
	Label	Indication	Meaning
	Link	Green	Linked at 10 Gb.
		Yellow	Linked at less than 10 Gb.
		Off	No link.
	Activity	Blinking On/Off	Actively transmitting or receiving data.
		Off	No link.

The **Intel® Gigabit 2P I350-t Adapter** has the following indicator lights:

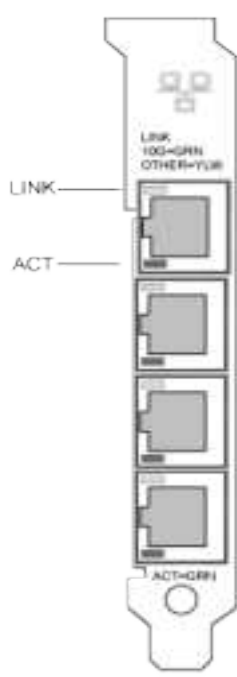
	Label	Indication	Meaning
	ACT/LNK	Green on	The adapter is connected to a valid link partner.
		Green flashing	Data activity
		Off	No link.
	10/100/1000	Off	10 Mbps
		Green	100 Mbps
		Yellow	1000 Mbps
		Orange flashing	Identity. Use the "Identify Adapter" button in Intel PROSet to control blinking. See Intel PROSet Help for more information.

Quad Port Copper Adapters

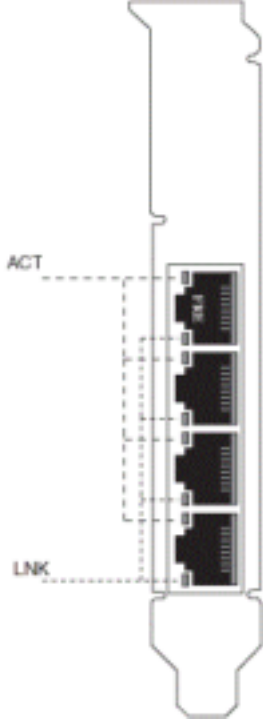
The **Intel® Ethernet 10G 4P X710-T4L-t OCP** has the following indicator lights:

	Label	Indication	Meaning
	Link	Green	Linked at 10 Gbps.
		Yellow	Linked at slower than 1 Gbps.
		Off	No link.
	Activity	Blinking On/Off	Actively transmitting or receiving data.
		Off	No link.

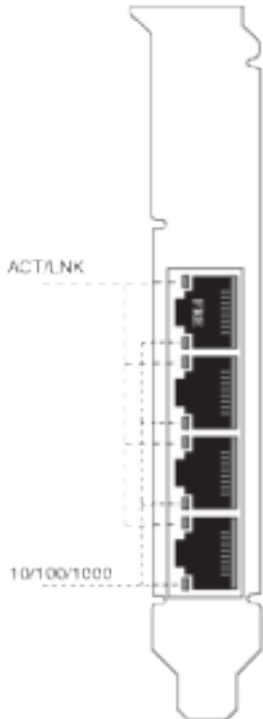
The **Intel® Ethernet 10G 4P X710-T4L-t Adapter** has the following indicator lights:

	Label	Indication	Meaning
	Link	Green	Linked at 10 Gbps.
		Yellow	Linked at slower than 10 Gbps.
		Off	No link.
	Activity	Blinking On/Off	Actively transmitting or receiving data.
		Off	No link.

The **Intel® Ethernet Converged Network Adapter X710** and **Intel® Ethernet Converged Network Adapter X710-T** have the following indicator lights:

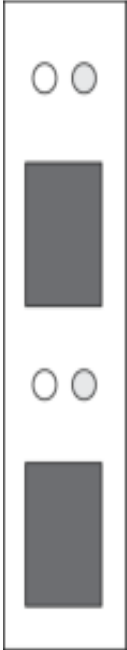
 The diagram shows the front panel of the adapter with two sets of indicator lights. The top set is labeled 'ACT' and the bottom set is labeled 'LNK'.	Label	Indication	Meaning
	ACT	Green on	The adapter is connected to a valid link partner.
		Green flashing	Data activity
		Off	No link.
	LNK	Green	10 Gbps
		Yellow	1 Gbps
		Off	100 Mbps

The **Intel® Gigabit 4P I350-t Adapter** has the following indicator lights:

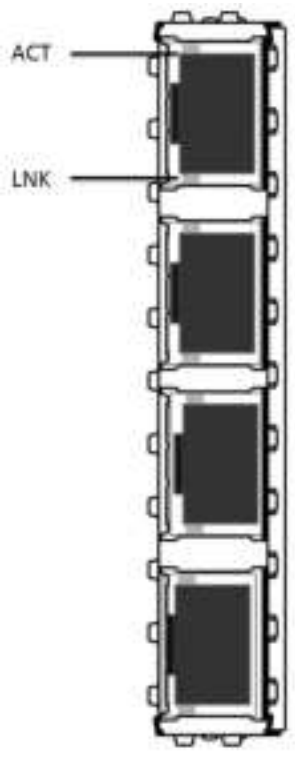
 The diagram shows the front panel of the adapter with two sets of indicator lights. The top set is labeled 'ACT/LNK' and the bottom set is labeled '10/100/1000'.	Label	Indication	Meaning
	ACT/LNK	Green on	The adapter is connected to a valid link partner.
		Green flashing	Data activity
		Off	No link.
	10/100/1000	Green	100 Mbps
		Yellow	1000 Mbps
		Orange flashing	Identify. Use the "Identify Adapter" button in Intel® PROSet to control blinking. See Intel PROSet Help for more information.
		Off	10 Mbps

rNDC (Rack Network Daughter Cards)

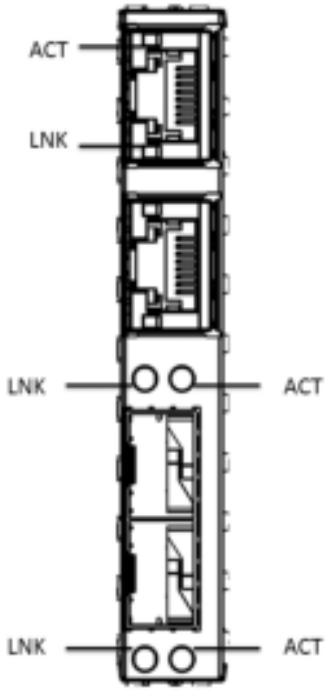
The **Intel® Ethernet 40G 2P XL710 QSFP+ rNDC** has the following indicator lights:

	Label	Indication	Meaning
	LNK (green/yellow)	Green on	Operating at maximum port speed.
		Off	No link.
	ACT (green)	Green flashing	Data activity.
		Off	No activity.

The Intel® Ethernet 1G 4P I350-t OCP, Intel® Ethernet 10G 4P X550/I350 rNDC, Intel® Gigabit 4P X550/I350 rNDC, Intel® Ethernet 10G 4P X550 rNDC, Intel® Ethernet 10G 4P X540/I350 rNDC, Intel® Gigabit 4P X540/I350 rNDC and Intel® Gigabit 4P I350-t rNDC have the following indicator lights:

	Label	Indication	Meaning
	LNK (green/yellow)	Green on	Operating at maximum port speed.
		Yellow on	Operating at lower port speed.
		Off	No link.
	ACT (green)	Green flashing	Data activity.
		Off	No activity.

The **Intel® Ethernet 10G 4P X520/I350 rNDC**, **Intel® Gigabit 4P X520/I350 rNDC**, **Intel® Ethernet Gigabit 4P x710/I350 rNDC**, and **Intel® 10G 4P X710/I350 rNDC** have the following indicator lights:

	Label	Indication	Meaning
	LNK (green/yellow)	Green on	Operating at maximum port speed.
		Yellow on	Operating at lower port speed.
		Off	No link.
	ACT (green)	Green flashing	Data activity.
		Off	No activity.

Transitioning from i40evf to iavf

Overview

Intel created the Intel® Adaptive Virtual Function (iavf) driver to provide a consistent, future-proof virtual function (VF) interface for Intel® Ethernet controllers. Previously, when you upgraded your network hardware, you replaced the drivers in each virtual machine (VM) with new drivers that were capable of accessing the new VF device provided by the new hardware. The iavf driver allows you to upgrade your network hardware without the need to update the VF driver in your existing VMs.

Support for adaptive virtual functions was initially added to the existing i40evf drivers for Microsoft® Windows Server® and Linux® operating systems. Intel renamed the driver from i40evf to iavf to clarify that the iavf driver will be the VF driver for future devices beyond the devices supported by the i40e driver. Release 18.8.0 was the last release that contained i40evf. Release 19.0.0 is the first release that contains iavf.

Supported Devices

The iavf driver supports devices based on the Intel® Ethernet Controller 800 Series and Intel® Ethernet Controller 700 Series.

Supported Operating Systems

- Novell® SUSE® Linux Enterprise Server (SLES) 15 SP1 and higher
- Red Hat® Enterprise Linux® (RHEL) 7.6 and higher
- RHEL 8.0 and higher
- Microsoft® Windows Server® 2016
- Microsoft® Windows Server® 2019 (x64 Edition)

Transition i40evf to iavf on Linux Operating Systems



NOTES:

- Do not use the i40evf device as your primary interface to access the VM. You must have another way to interact with the VM so you don't lose the connection when you disable the i40evf driver.
- If you are using a kernel/distro that does not contain iavf, and you update your kernel/distro, make sure that iavf is still loaded after the update.

Update Kernel

If you only use the drivers in the kernel or distro, you do not need to do anything until you update to a kernel or distro that contains the iavf driver. At that point you will need to update any scripts that call the driver by name.

If you update from kernel.org you will automatically get the iavf driver.

Using Linux RPM

1. Copy the iavf driver tar file to your VM image.
2. Unload the previous driver.

```
rmmod i40evf
```

3. Compile the driver module.

```
rpmbuild -tb /path/to/the/driver/file/iavf-[version].tar.gz
```

4. Install the driver.

a. RHEL:

```
rpm -i /root/rpmbuild/RPMS/x86_64/iavf-[version]-1.x86_64.rpm
```

b. SLES:

```
rpm -i /usr/src/packages/RPMS/x86_64/iavf-[version]-1.x86_64.rpm
```

5. Load the new driver module.

```
modprobe iavf
```

Install Using Linux tarball

1. Copy the iavf driver tar file to your VM image.

2. Untar the file.

```
tar xzf iavf-<x.x.x>.tar.gz
```

where

<x.x.x> is the version number for the driver tar file.

3. Change to the src directory under the unzipped driver file.

4. Compile the driver module.

```
make
```

```
make install
```

5. Make sure that any older i40evf drivers are removed from the kernel before loading the new module.

```
rmmmod i40evf
```

6. Load new driver module

```
modprobe iavf
```



NOTES: The “make install” command:

- Creates `/etc/modprobe.d/iavf-blacklist-i40evf.conf` that contains `blacklist i40evf`.
- Adds the line “alias i40evf iavf” to the modprobe configuration,

Linux Questions

I only use kernel/distro inbox drivers. Do I have to do anything?

There is nothing for you to do until you upgrade the operating system in your VM. If you update your VM kernel/distro to one that contains iavf, you will need to update any scripts that call the driver.

What if I want to keep i40evf and only update the PF driver to the latest version?

This is not a supported configuration. Please transition your driver to iavf when you update your PF driver. Release 18.8.0 was the last release that supported the i40evf driver.

What if I want to keep the old PF and only transition the VM to iavf?

There should be not issues with this scenario. If you run into problems, updating your PF driver may resolve them.

I have scripts that reference the VF drivers by name, do I have to change the scripts?

Yes. You must change the driver name in your scripts, rather than using an alias.

If I decide not to transition to the new iavf drivers, how long will Intel provide fixes for the old i40evf drivers?

Intel changed the name of the i40evf driver to iavf. All future updates and fixes will be published to the iavf driver.

Do I need to uninstall the i40evf driver?

This is not absolutely necessary, but we do recommend you uninstall the i40evf driver.

Is there any possibility of conflicts or situations where both drivers may exist in a system

Both drivers can be installed in the system. Installing the iavf driver tells the system that the iavf driver should be used instead of the i40evf driver. So when the system probes for new devices and finds a device that is supported by i40evf and iavf drivers, the system is told to always use the iavf driver.

Transition i40evf to iavf on Microsoft Windows Operating Systems



NOTES:

- Do not use the i40evf device as your primary interface to access the VM. You must have another way to interact with the VM so you don't lose the connection when you disable the i40evf driver.

1. Copy the iavf installer package to your VM image.
2. Use Add/Remove Programs to remove the i40evf driver.
3. Run the iavf install package to install the iavf driver.

If the i40evf driver does not show up in Add/Remove Programs, use Device Manager to remove it from all virtual NIC devices:

1. Open Device Manager.
2. Under Network Adapters, select the virtual NIC device.
3. Right-click and select Uninstall.
4. In the popup window, select the option to **Delete the driver software for this device**.
5. Click OK.

Windows Questions

What if I want to keep i40evf and only update the PF driver to the latest version?

This is not a supported configuration. Please transition your driver to iavf when you update your PF driver. Release 18.8.0 was the last release that supported the i40evf driver.

What if I want to keep the old PF and only transition the VM to iavf?

There should be no issues with this scenario. If you run into problems, updating your PF driver may resolve them.

Are the existing i40evf registry entries replaced or is a new registry created with the iavf driver names for the same port?

The i40evf registry entries are replaced.

What happens to i40evf driver files in system32? Are they deleted?

Not all of them. Some will remain.

Known Issues



NOTE: [iSCSI Known Issues](#) are located in their own section of this manual.

Fiber optics and auto-negotiation

Modules based on 100GBASE-SR4, 40GBASE-SR4, 25GBASE-SR, active optical cable (AOC), and active copper cable (ACC) do not support auto-negotiation per the IEEE specification. To obtain link with these modules, you must turn off auto-negotiation on the link partner's switch ports.

Link issues at speeds faster than 10 Gbps

If you are having link issues (including no link) at link speeds faster than 10 Gbps, check your switch configuration and/or specifications. Many optical connections and direct attach cables require RS-FEC for connection speeds faster than 10 Gbps. One of the following may resolve the issue:

- Configure your switch to use RS-FEC mode.
- Specify a 10 Gbps, or slower, link speed connection.
- If you are attempting to connect at 25 Gbps, try using an SFP28 CA-S or CS-N Direct Attach cable. These cables do not require RS-FEC.
- If your switch does not support RS-FEC mode, check with your switch vendor for the availability of a SW or FW upgrade.

The get-netadaptervmq PowerShell cmdlet displays less than the expected number of receive queues

After installing the Dell Update Package (DUP), the get-netadaptervmq PowerShell cmdlet reports 31 queues per port. This is expected behavior. The DUP changes the queue pooling default from pairs to groups of four. Pre-DUP, queues are paired into pools of two. After the DUP is installed, queues put into groups of four. This decreases the number of queues displayed by the get-netadaptervmq cmdlet.

NVM update utilities exit with error on Linux kernel 4.16 or higher

On Linux kernel 4.16 or higher, if you update the ixgbe, igb, or i40e driver and then run any of the NVM update utilities (NVMUpdate, NVMCheck, or Bootutil), the utility may exit with the error "The selected adapter cannot be updated due to strict MMIO memory settings in the kernel." To fix this, set the iomem kernel parameter to "relaxed" (i.e., iomem=relaxed) and reboot the system before running the tool again. On kernel 4.16 or higher, the iomem parameter is set to "strict" by default, which prevents the NVM update utilities from accessing the MMIO of the device.

Firmware downgrade to v18.0.x or older fails on X550 based devices

On X550 based devices, downgrading the firmware to version 18.0.x or older will fail and may result in NVM and Option ROM version incompatibility issues. To fix this issue, update to the latest firmware version.

An error occurred when updating a module on Intel® Ethernet 10G 2P X550-t Adapter using FW 17.5.0

If you use the FW DUP (Dell EMC Update Package) v17.5.0 to downgrade the firmware on an Intel® Ethernet 10G 2P X550-t Adapter, the DUP may report "An error occurred when updating a module." Please ignore this error message. The FW was successfully downgraded.

"Rx/Tx is disabled on this device because the module does not meet thermal requirements." error during POST

This error is caused by installing a module in an X710 based device that does not meet thermal requirements for that device. To resolve the issue, please install a module that meets the device's thermal requirements. See the section "[SFP+ and QSFP+ Devices](#)" in this document.

"Rx/Tx is disabled on this device because an unsupported SFP+ module type was detected." error during POST

This error is caused by installing an unsupported module in an X710/XL710 based device. You will not be able to send or receive traffic on this device. To resolve the issue, please install a supported module. See the section "[SFP+ and QSFP+ Devices](#)" in this document.

Missing virtual function ports in VMWare ESX

If you enable NPar and SR-IOV on the same device, the number of virtual functions enabled and displayed in lspci may be 8 or less. ESX limits the number of virtual functions to 8 per device. Also, due to ESXi limitations, the number of virtual functions created may be less than the number requested. See the ESXi documentation for details.

<http://pubs.vmware.com/>

Code 10 yellow bang errors on a Virtual Machine in Windows Device Manager

On a system running Microsoft Windows Server 2016, inside a Virtual Machine running Microsoft Windows Server 2016 or Windows Server 2012 R2, Intel Ethernet connections may have a code 10 yellow bang in Windows Device Manager. Installing a cumulative update that contains Microsoft KB3192366 and KB3176936 will resolve the issue.

Throughput Reduction After Hot-Replace

If an Intel gigabit adapter is under extreme stress and is hot-swapped, throughput may significantly drop. This may be due to the PCI property configuration by the Hot-Plug software. If this occurs, throughput can be restored by restarting the system.

CPU Utilization Higher Than Expected

Setting RSS Queues to a value greater than 4 is only advisable for large servers with several processors. Values greater than 4 may increase CPU utilization to unacceptable levels and have other negative impacts on system performance.

Supported SFP or SFP+ Module Not Recognized by the System

If you try to install an unsupported module, the port may no longer install any subsequent modules, regardless of whether the module is supported or not. The port will show a yellow bang under Windows Device Manager and an event id 49 (unsupported module) will be added to the system log when this issue occurs. To resolve this issue, the system must be completely powered off.

Windows Known Issues

Unable to shutdown virtual machine

Multiple VF failover events may leave a VM in an unstable state. You may not be able to shutdown the VM. Rebooting the host will resolve the issue.

Traffic does not transmit through VXLAN tunnel

On a system running Microsoft* Windows Server* 2016, traffic may fail to transmit through a VXLAN tunnel. Enabling Transmit Checksum Offloads for the appropriate traffic type will resolve the issue. For example, set "TCP Checksum Offload (IPv4)" to "TX Enabled" or "RX & TX Enabled"

hv_vmbus probe error on a Linux guest in a Windows Server system

On a system running Microsoft Windows Server 2019 or Windows Server 2016 on the host and Linux in the VF, you may see an "hv_vmbus: probe failed for device X" error in dmesg after you change a vSwitch from VMQ to SRIOV. This is due to a known timing issue in the operating system. There is no functionality loss, and the VF will successfully start after a few failed probes.

Incomplete branding string displayed in the event log

Some branding strings are too long to be displayed fully in the event log. In these cases, the branding string will be truncated and the port's PCI Bus/Device/Function are appended to the string. For example: Intel(R) Ethernet Converged Network Ad... [129,0,1].

PcieLinkSpeed is Unknown

When you install an Intel® Ethernet Controller 800 series device in a PCI Gen 4 slot, the operating system may report PcieLinkSpeed as Unknown. This does not affect the operation of the device.

Out of box driver will not uninstall via Web Console Apps & Features

You cannot use the Microsoft Windows Server 2016 Web Console Apps & Features menu to uninstall out-of-box drivers. Instead, use the Programs and Features selection in the Windows Control Panel.

Port is missing from Lifecycle Controller : Network Settings

If a port is configured for iSCSI boot, and it successfully connected to its boot target, then you cannot modify the port settings in the Lifecycle Controller.

Procedure for Installing and Upgrading Drivers and Utilities

Intel does not recommend installing or upgrading drivers and Intel® PROSet software over a network connection. Instead, install or upgrade drivers and utilities from each system.

Advanced Properties Settings Change While Traffic is Running

In the Advanced Properties tab of Intel® PROSet, parameters should not be modified under heavy network loads. Otherwise, a reboot may be required to make the changes effective.

In a Microsoft Hyper-V environment, Virtual Machines bound to NPAR partitions will not communicate with each other

In a Microsoft Hyper-V environment, if you have NPAR enabled on a port, and Virtual Machines (VMs) bound to partitions on that port, the VMs may not be able to communicate with each other. This happens because the virtual switch inside Hyper-V sends the packets to the physical port, which sends the packets to the switch that is connected to the port. The physical switch may not be configured for reflective relay (also called hairpin mode), so it may not send the packets back on the same connection from which it received them. Connecting the port to a Virtual Ethernet Port Aggregator (VEPA) capable switch will resolve the issue.

Intel drivers must be installed by Dell EMC Update Package before configuring Microsoft Hyper-V features

Prior to configuring the Microsoft* Hyper-V features, the Intel® NIC drivers must be installed by the Dell EMC Update Package. If the Microsoft* Hyper-V feature is configured on an unsupported NIC partition on an Intel® X710 device prior to using the Dell EMC Update Package to install Intel® NIC Drivers, the driver installation may not complete. To recover, you must uninstall Microsoft* Hyper-V, uninstall 'Intel® Network Connections' from 'Programs and Features', and use the Dell EMC Update Package to install Intel® NIC Drivers.

Application Error Event IDs 789 and 790 in the Event Log

If Data Center Bridging (DCB) is enabled, and the enabled port loses link, the following three events may be logged in the event log:

- Event ID 789: Enhanced Transmission Selection feature on a device has changed to non-operational
- Event ID 790: Priority Flow Control feature on a device has changed to non-operational

This is the expected behavior when a DCB enabled port loses link. DCB will begin working again as soon as link is reestablished. A port will lose link if the cable is disconnected, the driver or software package is updated, if the link partner goes down, or for other reasons.

"Malicious script detected" Warning from Norton AntiVirus During PROSet Uninstall

The Intel PROSet uninstall process uses a Visual Basic script as part of the process. Norton AntiVirus and other virus scanning software may mistakenly flag this as a malicious or dangerous script. Letting the script run allows the uninstall process to complete normally.

Unexpected Connectivity Loss

If you uncheck the "Allow the computer to turn off this device to save power" box on the Power Management tab and then put the system to sleep, you may lose connectivity when you exit sleep. You must disable and enable the NIC to resolve the issue. Installing Intel® PROSet for Windows Device Manager will also resolve the issue.

RSS Load Balancing Profile Advanced Setting

Setting the "RSS load balancing profile" Advanced Setting to "ClosestProcessor" may significantly reduce CPU utilization. However, in some system configurations (such as a system with more Ethernet ports than processor cores), the "ClosestProcessor" setting may cause transmit and receive failures. Changing the setting to "NUMAScalingStatic" will resolve the issue.

Opening Windows Device Manager property sheet takes longer than expected

The Windows Device Manager property sheet may take 60 seconds or longer to open. The driver must discover all Intel Ethernet devices and initialize them before it can open the property sheet. This data is cached, so subsequent openings of the property sheet are generally quicker.

Linux Known Issues

HeaderDataSplit is not supported in 82599-based adapters.

Configuring the Driver on Different Distributions

Configuring a network driver to load properly when the system is started (0=legacy, 1=MSI, 2=MSI-X) is distribution dependent. Typically, the configuration process involves adding an alias line to /etc/modules.conf or /etc/modprobe.conf as well as editing other system startup scripts and/or configuration files. Many popular Linux distributions ship with tools to make these changes for you. To learn the proper way to configure a network device for your system, refer to your distribution documentation.

Enabling WOL in Linux Using Ethtool and BootUtil

By default, WOL is disabled. In a Linux environment, WOL is enabled using ethtool and, in some instances, using BootUtil is also required. Only port A (port 0) can be enabled through ethtool without using BootUtil. To enable WOL using ethtool on other ports, WOL must be enabled with BootUtil first.

Link LEDs May Be Off During Offline Installation of Linux OS

When performing an offline installation (installation without a NIC connected to a valid network) of SUSE Linux Enterprise Server 15, the PHY link is disabled in the OS even if the Ethernet cable is plugged in. The PHY link is disabled due to the Intel driver disabling the link for power savings when the network is not in use. Configuring the network setting during or after the installation restores link LEDs.

Power Management Known Issues

Intel® Ethernet Controller X710 series devices do not support Wake-On-LAN in multicast mode

Devices based on the Intel® Ethernet Controller X710 do not support Wake-On-Lan (WOL) in multicast mode.

System does not wake on link

On a driver-only installation, if you change 'Wake on Link Settings' to Forced and change 'Wake on Magic Packet' and 'Wake on Pattern Match' to Disabled, the system may not wake up when expected. In order to "Wake on Link" successfully, check the Power Management tab and make sure that "Allow this device to wake the computer" is checked. You may also need to change 'Wake on Magic Packet' or 'Wake on Pattern Match' to Enabled.

Directed Packets may not wake the system

On some systems, quad port server adapters may not wake when configured to wake on directed packet. If you experience problems waking on directed packets, you must configure the adapter to use Magic Packets*.

Power Management options are unavailable or missing

If you install only the base drivers, later install Intel® PROSet for Windows Device Manager, then remove Intel PROSet, the settings on the Power Management tab on the Adapter Property Sheet may be unavailable or missing altogether. You must reinstall Intel PROSet to resolve the issue.

System Wakes-Up from a Removed VLAN

If a system goes into standby mode, and a directed packet is sent to the IP address of the removed VLAN, the system will wake-up. This occurs because a directed packet bypasses VLAN filtering.

Intel Adapters ignore consecutive Wake Up signals while transitioning into standby mode

While sending a system into standby, occasionally a wake up packet arrives before the system completes the transition into standby mode. When this happens, the system ignores consecutive wake up signals and remains in standby mode until manually powered up using the mouse, keyboard, or power button.

Other Intel 10GbE Network Adapter Known Issues

The System H/W Inventory (iDRAC) indicates that Auto-negotiation on the Embedded NIC is Disabled, but elsewhere link speed and duplex auto-negotiation is Enabled

If an optical module is plugged into the Intel® Ethernet 10G X520 LOM on a PowerEdge-C6320, the System H/W Inventory (iDRAC) will indicate that Auto-negotiation is Disabled. However, Windows Device Manager and HII indicate that link speed and duplex Auto-negotiation is Enabled. This is because the driver contains an algorithm that allows the LOM to link with SFP partners at 10 Gbps or 1 Gbps. This is reported to Windows Device Manager and HII, but it is not true auto-negotiation. iDRAC reads the device's firmware, which has no knowledge of the algorithm, and therefore reports that auto-negotiation is disabled.

ETS Bandwidth Allocations Don't Match Settings

When Jumbo Frames is set to 9K with a 10GbE adapter, a 90%/10% ETS traffic split will not actually be attained on any particular port, despite settings being made on the DCB switch. When ETS is set to a 90%/10% split, an actual observed split of 70%/30% is more likely.

Link Loss on 10GbE Devices with Jumbo Frames Enabled

You must not lower Receive_Buffers or Transmit_Buffers below 256 if jumbo frames are enabled on an Intel® 10GbE Device. Doing so will cause loss of link.

Continuous PFC pause frames sent from Intel® Ethernet X520 based devices

If you have an Intel® Ethernet X520 based device connected to a switch port and modify the DCB bandwidth settings on the switch port, the Intel® Ethernet X520 device may perpetually send pause frames, causing a storm, and fail to transfer data to and from the storage targets it was using. To recover from this issue, disable the X520 ports, re-enable them, and then reconnect to the iSCSI target volumes. To avoid the issue, if the DCB bandwidth settings need to be changed, do one of the following:

- Power down the server that contains the Intel® Ethernet X520 device prior to modifying the DCB bandwidth settings.
- Disable the switch ports connected to Intel X520 based device.
- Have no traffic running on the Intel X520 based device.

Intel® Ethernet 10G 2P/4P X710-k bNDC does not have link and is not displayed in Windows Device Manager

If you install an Intel® Ethernet 10G 2P X710-k bNDC or an Intel® Ethernet 10G 4P X710-k bNDC onto a Dell EMC PowerEdge M630/M830 blade server, and install that blade into an M1000e chassis, the bNDC may not have link and may display a yellow bang, or may not be displayed at all, in Windows Device Manager. This is limited to the 1.0 version of the M1000e Midplane.

Intel® Ethernet 10G X520 LOM links at 10 Gbps when 1.0 Gbps Full Duplex is selected

When connected with a direct attach cable, the Intel® Ethernet 10G X520 LOM will always connect at 10 Gbps.

Intel X540-t and Dell Force10 will not establish link at 100 Mbps full duplex if set manually on both ends

For an X540-t based adapter coupled with a Force10 component, in order to run at 100Mbps, the properties for BOTH components must be set to Auto-Negotiation ON.

When trying to identify the adapter, the Activity LED blinks and the Link LED is solid

If you use the Identify Adapter feature with the following adapters, the Activity LED blinks instead of the Link LED. The Link LED may display a solid green light for 10G ports even if a network link is not present.

- All Intel® Ethernet X520 10GbE devices
- All Intel® Ethernet X540 10GbE devices
- All Intel® Ethernet X550 10GbE devices
- Some Intel® Gigabit I350 LOM devices

Intel® Ethernet Controller 700 Series Devices Known Issues

Some Intel® X710 based devices report a subdevice ID of 0x0000 and may display a generic branding string. Port 0 reports the correct subvendor ID and displays the correct branding string.

Intel X710 based devices may maintain link on any and all ports as long as power is provided to the device, regardless of the device's or system's power state.

Unexpected IntelDCB errors in the Windows Application Event Log

After upgrading your X710 drivers, you may see several IntelDCB errors in the Windows Application Event Log. These errors are erroneous and can be ignored.

Lower than expected throughput on X710/XL710 based devices

If you have an X710 or XL710 based device installed in a four CPU socket system. Receive and transmit traffic may be significantly lower than expected. Setting your interrupt rate to High may mitigate the issue.

Cable tests unavailable with Broadcom BCM84886 transceiver

The cable diagnostic tests may be unavailable if you have a Broadcom BCM84886 transceiver installed on a port. The transceiver does not support TDR diagnostics. See the BCM84886 datasheet (84886-DS103; October 27, 2017 revision). This issue affects the following devices:

- Intel® Ethernet Converged Network Adapter X710-T
- Intel® Ethernet 25G 2P XXV710 Mezz
- Intel® Ethernet 10G 2P X710-T2L-t Adapter
- Intel® Ethernet 10G 4P X710-T4L-t Adapter
- Intel® Ethernet 10G 2P X710-T2L-t OCP
- Intel® Ethernet 10G 4P X710-T4L-t OCP

Wake on LAN erroneously available in iDRAC/racadm

The Intel® Ethernet Converged Network Adapter X710-2 only supports WoL on port 1. If NPAR is enabled, WoL is only supported on the first partition of port 1. If you view WoL status with iDRAC/racadm, WoL may erroneously appear as available on other ports and partitions.

Intel® Gigabit 4P I350-t Adapter Known Issues

Downshifting

When connecting to any Gigabit switch via a faulty CAT 5 cable where one pair is broken, the adapter does not downshift from 1 Gig to 100Mbps. For the adapter to downshift, it must identify two broken pairs in the cable.

System does not boot

Your system may run out of I/O resources and fail to boot if you install more than four quad port server adapters. Moving the adapters to different slots or rebalancing resources in the system BIOS may resolve the issue. This issue affects the following Adapters:

- Intel® Gigabit 4P I350-t Adapter

Regulatory Compliance Statements

FCC Class A Products

40 Gigabit Ethernet Products

- Intel® Ethernet 40G 2P XL710 QSFP+ rNDC
- Intel® Ethernet Converged Network Adapter XL710-Q2

25 Gigabit Ethernet Products

- Intel® Ethernet 25G 2P E810-XXV OCP
- Intel® Ethernet 25G 2P XXV710 Mezz
- Intel® Ethernet 25G 2P XXV710 Adapter

10 Gigabit Ethernet Products

- Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz
- Intel® Ethernet 10G 2P X540-t Adapter
- Intel® Ethernet 10G 2P X550-t Adapter
- Intel® Ethernet 10G 4P X550 rNDC
- Intel® Ethernet 10G 4P X550/I350 rNDC
- Intel® Ethernet 10G 4P X540/I350 rNDC
- Intel® Ethernet 10G 4P X520/I350 rNDC
- Intel® Ethernet 10G 2P X520-k bNDC
- Intel® Ethernet 10G 4P X710-k bNDC
- Intel® Ethernet 10G 2P X710-k bNDC
- Intel® Ethernet 10G X710-k bNDC
- Intel® Ethernet Converged Network Adapter X710
- Intel® Ethernet Converged Network Adapter X710-T
- Intel® Ethernet 10G 4P X710/I350 rNDC
- Intel® Ethernet 10G 4P X710 SFP+ rNDC
- Intel® Ethernet 10G X710 rNDC
- Intel® Ethernet Server Adapter X710-DA2 for OCP
- Intel® Ethernet 10G 2P X710 OCP
- Intel® Ethernet 10G 4P X710 OCP
- Intel® Ethernet 10G 2P X710-T2L-t OCP
- Intel® Ethernet 10G 4P X710-T4L-t OCP

Gigabit Ethernet Products

- Intel® Ethernet 1G 4P I350-t OCP
- Intel® Gigabit 4P X550/I350 rNDC
- Intel® Gigabit 4P I350-t rNDC
- Intel® Gigabit 4P X540/I350 rNDC
- Intel® Gigabit 4P X520/I350 rNDC
- Intel® Gigabit 4P I350-t Mezz
- Intel® Gigabit 4P X710/I350 rNDC
- Intel® Gigabit 4P I350 bNDC

FCC Class B Products

25 Gigabit Ethernet Products

- Intel® Ethernet 25G 2P E810-XXV Adapter

10 Gigabit Ethernet Products

- Intel® Ethernet 10G 2P X710-T2L-t Adapter
- Intel® Ethernet 10G 4P X710-T4L-t Adapter
- Intel® Ethernet 10G 2P X520 Adapter
- Intel® Ethernet 10G X520 LOM

Gigabit Ethernet Products

- Intel® Gigabit 2P I350-t Adapter
- Intel® Gigabit 4P I350-t Adapter

Safety Compliance

The following safety standards apply to all products listed above:

- UL 60950-1, 2nd Edition, 2011-12-19 (Information Technology Equipment - Safety - Part 1: General Requirements)
- UL 62368-1 2nd Edition (Information Technology Equipment - Safety requirements)
- CSA C22.2 No. 60950-1-07, 2nd Edition, 2011-12 (Information Technology Equipment - Safety - Part 1: General Requirements)
- CAN/CSA C22.2 European Group Differences and National Differences according to 62368-1-14 - Audio/video, information and communication technology equipment - Part 1: Safety requirements
- EN 60950-1:2006/A11:2009/A1:2010/A12:2011 (European Union)
- IEC 60950-1:2005 (2nd Edition); Am 1:2009 (International)
- EU LVD Directive 2006/95/EC

EMC Compliance

The following standards may apply.

Class A Products

- FCC Part 15 – Radiated & Conducted Emissions (USA)
- CAN ICES-3(A)/NMB-3(A) – Radiated & Conducted Emissions (Canada)
- CISPR 22 – Radiated & Conducted Emissions (International)
- EN55022: 2010 – Radiated & Conducted Emissions (European Union)
- EN55024: 2010 +A1:2001+A2:2003 – Immunity (European Union)
- EN55032: 2015 Class A Radiated and Conducted Emissions requirements (European Union)
- EMC Directive 2004/108/EC (European Union)
- VCCI (Class A)– Radiated & Conducted Emissions (Japan)
- CNS13438 – Radiated & Conducted Emissions (Taiwan)
- AS/NZS CISPR 22:2009 + A1:2010 Class A and CISPR 32:2015 for Radiated and Conducted Emissions requirements (Australia/New Zealand)
- NRRA No. 2012-13 (2012.06.28), NRRA Notice No. 2012-14 (2012.06.28) (Korea)

Class B Products

- FCC Part 15 (Class B) – Radiated & Conducted Emissions (USA)
- CAN ICES-3(B)/NMB-3(B) – Radiated & Conducted Emissions (Canada)
- CISPR 22 – Radiated & Conducted Emissions (International)

- EN55022: 2010 – Radiated & Conducted Emissions (European Union)
- EN55024: 2010 – Immunity (European Union)
- EN55032: 2015 Class B Radiated and Conducted Emissions requirements (European Union)
- EMC Directive 2004/108/EC (European Union)
- VCCI (Class B)– Radiated & Conducted Emissions (Japan) (excluding optics)
- CNS13438 (Class B)-2006 – Radiated & Conducted Emissions (Taiwan) (excluding optics)
- AS/NZS CISPR 22:2009 + A1:2010 Class B and CISPR 32:2015 for Radiated and Conducted Emissions requirements (Australia/New Zealand)
- KN22; KN24 – Korean emissions and immunity
- NRR No. 2012-13 (2012.06.28), NRR Notice No. 2012-14 (2012.06.28) (Korea)

Hazardous Substances Compliance

The following standards may apply:

- EU REACH directive
- EU WEEE directive
- EU RoHS directive
- China RoHS directive
- BSMI CNS15663: Taiwan RoHS

Regulatory Compliance Markings

When required, these products are provided with the following Product Certification Markings:

- UL Recognition Mark for USA and Canada
- CE Mark
- EU WEEE Logo
- FCC markings
- VCCI marking
- Australian C-Tick Mark
- Korea MSIP mark
- Taiwan BSMI mark
- People's Republic of China "EFUP" mark

FCC Class A User Information

The Class A products listed above comply with Part 15 of the FCC Rules. Operation is subject to the following two conditions:

1. This device may not cause harmful interference.
2. This device must accept any interference received, including interference that may cause undesired operation.



NOTE: This equipment has been tested and found to comply with the limits for a Class A digital device, pursuant to Part 15 of the FCC rules. These limits are designed to provide reasonable protection against harmful interference when the equipment is operated in a commercial environment. This equipment generates, uses and can radiate radio frequency energy and, if not installed and used in accordance with the instructions, may cause harmful interference to radio communications. Operation of this equipment in a residential area is likely to cause harmful interference in which case the user will be required to correct the interference at his own expense.



CAUTION: If the device is changed or modified without permission from Intel, the user may void his or her authority to operate the equipment.

Canadian Compliance (Industry Canada)

CAN ICES-3(A)/NMB-3(A)

VCCI Class A Statement

この装置は、クラスA情報技術装置です。この装置を家庭環境で使用すると電波妨害を引き起こすことがあります。この場合には使用者が適切な対策を講ずるよう要求されることがあります。	VCCI-A
---	--------

BSMI Class A Statement

警告使用者:
此為甲類資訊技術設備，於居住環境中使用時，可能會造成射頻擾動，在此種情況下，使用者會被要求採取某些適當的對策。

KCC Notice Class A (Republic of Korea Only)

A급 기기 (업무용 방송통신기기)	이 기기는 업무용(A급)으로 전파관리처장령을 한 기기이오니 판매자 또는 사용자는 이 것을 주의하시기 바라며, 가정외의 지역에서 사용하는 것을 목적으로 합니다.
CLASS A device (commercial broadcasting and communication equipment)	This device has been approved by EMC registration. Distributors or users pay attention to this point. This device is usually aimed to be used in other area except at home.

BSMI Class A Notice (Taiwan)

警告使用者:
此為甲類資訊技術設備，於居住環境中使用時，可能會造成射頻擾動，在此種情況下，使用者會被要求採取某些適當的對策。

FCC Class B User Information

This equipment has been tested and found to comply with the limits for a Class B digital device pursuant to Part 15 of the FCC rules. These limits are designed to provide reasonable protection against harmful interference in a residential installation. This equipment generates, uses and can radiate radio frequency energy and, if not installed and used in accordance with the instructions, may cause harmful interference to radio communications. However, there is no guarantee that interference will not occur in a particular installation.

If this equipment does cause harmful interference to radio or television reception, which can be determined by turning the equipment off and on, the user is encouraged to try to correct the interference by one or more of the following measures:

- Reorient or relocate the receiving antenna.
- Increase the separation between the equipment and receiver.
- Connect the equipment into an outlet on a circuit different from that to which the receiver is connected.
- Consult the dealer or an experienced radio/TV technician for help.



CAUTION: If the device is changed or modified without permission from Intel, the user may void his or her authority to operate the equipment.



NOTE: This device complies with Part 15 of the FCC Rules. Operation is subject to the following two conditions: (1) this device may not cause harmful interference, and (2) this device must accept any interference received, including interference that may cause undesired operation.

Electromagnetic Compatibility Notices

FCC Declaration of Conformity Statement

The following products have been tested to Comply with FCC Standards for Home or Office Use:

PRO/1000 PT, PRO/1000 GT, Gigabit PT, Gigabit ET, I210-T1, I340-T2/T4, and I350-T2/T4.

Canadian Compliance (Industry Canada)

CAN ICES-3 (B)/NMB-3 (B)

VCCI Class B Statement (Japan)

この装置は、クラスB情報技術装置です。この装置は、家庭環境で使用することを目的としていますが、この装置がラジオやテレビジョン受信機に近接して使用されると、受信障害を引き起こすことがあります。

取扱説明書に従って正しい取り扱いをして下さい。

VCCI-B

KCC Notice Class B (Republic of Korea Only)

B급 기기 (가정용 방송통신기기) CLASS B device residential broadcasting and communication equipment	이 기기는 가정용(B급)으로 전자파적합등록을 한 기기로서 주로 가정에서 사용하는 것을 목적으로 하며, 모든 지역에서 사용할 수 있습니다. This device has been approved by EMC Registration and is usually aimed to be used in a residential area so that it can be used in all other location as well as at home.
---	---

EU WEEE Logo



Manufacturer Declaration European Community



Manufacturer Declaration

Intel Corporation declares that the equipment described in this document is in conformance with the requirements of the European Council Directive listed below:

- Low Voltage Directive 2006/95/EC
- EMC Directive 2004/108/EC
- RoHS Directive 2011/65/EU

These products follow the provisions of the European Directive 1999/5/EC.

Dette produkt er i overensstemmelse med det europæiske direktiv 1999/5/EC.

Dit product is in navolging van de bepalingen van Europees Directief 1999/5/EC.

Tämä tuote noudattaa EU-direktiivin 1999/5/EC määräyksiä.

Ce produit est conforme aux exigences de la Directive Européenne 1999/5/EC.

Dieses Produkt entspricht den Bestimmungen der Europäischen Richtlinie 1999/5/EC.

Þessi vara stenst reglugerð Evrópska Efnahags Bandalagsins númer 1999/5/EC.

Questo prodotto è conforme alla Direttiva Europea 1999/5/EC.

Dette produktet er i henhold til bestemmelsene i det europeiske direktivet 1999/5/EC.

Este produto cumpre com as normas da Diretiva Europeia 1999/5/EC.

Este producto cumple con las normas del Directivo Europeo 1999/5/EC.

Denna produkt har tillverkats i enlighet med EG-direktiv 1999/5/EC.

This declaration is based upon compliance of the Class A products listed above to the following standards:

EN 55022:2010 (CISPR 22 Class A) RF Emissions Control.

EN 55024:2010 (CISPR 24) Immunity to Electromagnetic Disturbance.

EN 60950-1:2006/A11:2009/A1:2010/A12:2011 Information Technology Equipment- Safety-Part 1: General Requirements.

EN 50581:2012 - Technical documentation for the assessment of electrical and electronic products with respect to the restriction of hazardous substances.

This declaration is based upon compliance of the Class B products listed above to the following standards:

EN 55022:2010 (CISPR 22 Class B) RF Emissions Control.

EN 55024:2010 (CISPR 24) Immunity to Electromagnetic Disturbance.

EN 60950-1:2006/A11:2009/A1:2010/A12:2011 Information Technology Equipment- Safety-Part 1: General Requirements.

EN 50581:2012 - Technical documentation for the assessment of electrical and electronic products with respect to the restriction of hazardous substances.



WARNING: In a domestic environment, Class A products may cause radio interference, in which case the user may be required to take adequate measures.

Responsible Party

Intel Corporation, Mailstop JF3-446
5200 N.E. Elam Young Parkway
Hillsboro, OR 97124-6497
Phone 1-800-628-8686

China RoHS Declaration

关于符合中国《电子信息产品污染控制管理办法》的声明

**Management Methods on Control of Pollution From
Electronic Information Products
(China RoHS declaration)**

产品中有毒有害物质的名称及含量

部件名称	有毒有害物质或元素					
	铅 (Pb)	汞 (Hg)	镉 (Cd)	六价铬 (Cr(VI))	多溴联苯 (PBB)	多溴二苯醚 (PBDE)
印刷板组件	X	○	○	○	○	○
○：表示该有毒有害物质在该部件所有均质材料中的含量均在 SJ/T 11363-2006 标准规定的限量要求以下。 X：表示该有毒有害物质至少在该部件的某一均质材料中的含量超出 SJ/T 11363-2006 标准规定的限量要求。						

Class 1 Laser Products

Server adapters listed above may contain laser devices for communication use. These devices are compliant with the requirements for Class 1 Laser Products and are safe in the intended use. In normal operation the output of these laser devices does not exceed the exposure limit of the eye and cannot cause harm.

For continued safe operation in case of an abnormal circumstance, always have the provided laser connector cover in place or a compatible fiber optics cable properly connected when power is available to the product.

The Laser device must be factory serviced ONLY by the responsible manufacturer! NO adjustments, service or maintenance is to be performed otherwise.



CAUTION: Use of controls or adjustments or performance of procedures other than those specified herein may result in hazardous radiation exposure.

These Class 1 Laser devices:

Comply with FDA/CDRH per CFR21, subchapter J.
Comply with IEC 60825-1:2007

End-of-Life / Product Recycling

Product recycling and end-of-life take-back systems and requirements vary by country.

Contact the retailer or distributor of this product for information about product recycling and/or take-back.

Customer Support

Web and Internet Sites

<http://support.dell.com/>

Customer Support Technicians

If the troubleshooting procedures in this document do not resolve the problem, please contact Dell, Inc. for technical assistance (refer to the "Getting Help" section in your system documentation).

Before you call...

You need to be at your computer with your software running and the product documentation at hand.

The technician may ask for the following:

- Your address and telephone number
- The name and model number of the product you are calling about
- The serial number and service tag of the product
- The names and version numbers of the software you are using to operate the product
- The name and version number of the operating system you are using
- The computer type (manufacturer and model number)
- Expansion boards or add-in cards in your computer
- The amount of memory in your computer

Adapter Specifications

Intel® 40 Gigabit Network Adapter Specifications

Feature	Intel® Ethernet Converged Network Adapter XL710-Q2
Bus Connector	PCI Express 3.0
Bus Speed	x8
Transmission Mode/Connector	QSFP+
Cabling	40GBase-SR4, Twinax DAC (7m max)
Power Requirements	6.5 W Maximum @ +12 V
Dimensions (excluding bracket)	5.21 x 2.71 in 13.3 x 6.9 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	159 years
Available Speeds	40 Gbps
Duplex Modes	Full only
Indicator Lights	<i>Two per port:</i> Link and Activity
Standards Conformance	IEEE 802.3ba SFF-8436 PCI Express 3.0
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No. 60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)

Intel® 40GbE Network Daughter Cards (NDC) Specifications

Feature	Intel® Ethernet 40G 2P XL710 QSFP+ rNDC
Bus Connector	PCI Express 3.0

Bus Speed	x8
Transmission Mode/Connector	QSFP+
Cabling	40GBase-SR4, Twinax DAC (7m max)
Power Requirements	6.2 W Maximum @ +12 V
Dimensions (excluding bracket)	3.66 x 6.081 in 9.3 x 15.5 cm
Operating Temperature	32 - 140 deg. F (0 - 60 deg. C)
MTBF at 55°C	112 years
Available Speeds	40 Gbps
Duplex Modes	Full only
Indicator Lights	<i>Two per port:</i> Link and Activity
Standards Conformance	IEEE 802.3ba SFF-8436 PCI Express 3.0
Regulatory and Safety	Safety Compliance • UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) • EN 60 950 (European Union) • IEC 60 950 (International) EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55022-1998 - Radiated & Conducted Emissions (European Union) • EN55024 - 1998 - (Immunity) (European Union) • CE - EMC Directive (89/336/EEC) (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) • MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)

Intel® 25 Gigabit Network Adapter Specifications

Feature	Intel® Ethernet 25G 2P E810-XXV Adapter
Bus Connector	PCI Express 4.0 PCI Express 3.0
Bus Speed	x8
Transmission Mode/Connector	SFP28
Cabling	25GBase-CR, Twinax DAC (3m max)
Power Requirements	25 W maximum @ +12 V

	3.63 W maximum @ +3.3 V
Dimensions (excluding bracket)	2.54 in. x 6.6 in. 6.44 cm x 16.76 cm
Operating Temperature	32 - 140°F (0 - 60°C)
MTBF at 55°C	271 years
Available Speeds	25 Gbps/10 Gbps/1 Gbps
Duplex Modes	Full only
Indicator Lights	<i>Two per port:</i> Link and Activity
<u>Standards Conformance</u>	PCI Express 4.0 SFF-8419 IEEE 802.3
Regulatory and Safety	EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55032-2015- Radiated & Conducted Emissions (European Union) • EN55024 - 2010- (Immunity) (European Union) • REACH, WEEE, RoHS Directives (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS CISPR - Radiated & Conducted Emissions (Australia/New Zealand) • KN22 -Radiated & Conducted Emissions (Korea) • RoHS (China)

Feature	Intel® Ethernet 25G 2P XXV710 Adapter
Bus Connector	PCI Express 3.0
Bus Speed	x8
Transmission Mode/Connector	SFP28
Cabling	25GBase-CR, Twinax DAC (3m max)
Power Requirements	6.5 W Maximum @ +12 V
Dimensions (excluding bracket)	2.70 x 2.02 in 6.86 x 5.12 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	239 years
Available Speeds	25 Gbps/10 Gbps/1 Gbps
Duplex Modes	Full only
Indicator Lights	<i>Two per port:</i> Link and Activity
<u>Standards Conformance</u>	IEEE 802.3-2015

	SFF-8431 PCI Express 3.0
Regulatory and Safety	Safety Compliance • UL/CSA 60950-1-07 2nd Edition • EN 60 950 (European Union) • IEC 60 950 (International) EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55032-2015-Radiated & Conducted Emissions (European Union) • EN55024 - 2010- (Immunity) (European Union) • REACH, WEEE, RoHS Directives (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS CISPR - Radiated & Conducted Emissions (Australia/New Zealand) • KN32 -Radiated & Conducted Emissions (Korea) • KN35 - (Immunity) (Korea) • RoHS

Intel® 25 Gigabit Network Mezzanine Card Specifications

Feature	Intel® Ethernet 25G 2P E810-XXV OCP
Bus Connector	PCI Express 4.0 PCI Express 3.0
Bus Speed	x8
Transmission Mode/Connector	SFP28
Cabling	25GBase-CR, Twinax DAC (3m max)
Power Requirements	25 W maximum @ +12 V 3.63 W maximum @ +3.3 V
Dimensions (excluding bracket)	4.53 in. x 2.99 in. 11.5 cm x 7.6 cm
Operating Temperature	41 - 149°F (5 - 65°C)
MTBF at 55°C	266 years
Available Speeds	25 Gbps/10 Gbps/1 Gbps
Duplex Modes	Full only
Indicator Lights	<i>Two per port:</i> Link and Activity
<u>Standards Conformance</u>	PCI Express 4.0 SFF-8431 IEEE 802.3 OCP NIC 3.0
Regulatory and Safety	EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada)

	• CISPR 22 - Radiated & Conducted Emissions (International) • EN55032-2015- Radiated & Conducted Emissions (European Union) • EN55024 - 2010- (Immunity) (European Union) • REACH, WEEE, RoHS Directives (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS CISPR - Radiated & Conducted Emissions (Australia/New Zealand) • KN22 -Radiated & Conducted Emissions (Korea) • RoHS (China)
--	--

Feature	Intel® Ethernet 25G 2P XXV710 Mezz
Bus Connector	PCI Express 3.0
Bus Speed	x8
Transmission Mode/Connector	SFP28
Cabling	25GBase-CR, Twinax DAC (3m max)
Power Requirements	9.78W @ +12V
Dimensions (excluding bracket)	3.78 x 3.15 9.60 x 8.001
Operating Temperature	105° F max
MTBF at 55°c	353years
Available Speeds	25 Gbps
Duplex Modes	Full only
Indicator Lights	NA Link and Activity
Standards Conformance	IEEE 802.3 backplane standards
Regulatory and Safety	EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55032-2015- Radiated & Conducted Emissions (European Union) • EN55024 - 2010- (Immunity) (European Union) • REACH, WEEE, RoHS Directives (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS CISPR - Radiated & Conducted Emissions (Australia/New Zealand) • KN22 -Radiated & Conducted Emissions (Korea) • RoHS (China)

Intel® 10 Gigabit Network Adapter Specifications

Feature	Intel® Ethernet 10G 2P X710 OCP	Intel® Ethernet 10G 2P X710-T2L-t OCP	Intel® Ethernet 10G 4P X710-T4L-t OCP
Bus Connector	PCI Express 3.0	OCP NIC 3.0	OCP NIC 3.0
Bus Speed	x8	x8 PCI Express v3.0	x8 PCI Express v3.0
Transmission Mode/Connector	10G/SFP+	RJ45 BASE-T Connector	RJ45 BASE-T Connector
Cabling	10GBASE-SR 10GBASE-LR SFP+ Direct Attach Cables	10GBASE-T: CAT6A (100m max), CAT6 (55m max) 1000BASE-T: CAT6A, CAT6, CAT5e (100m max)	10GBASE-T: CAT6A (100m max), CAT6 (55m max) 1000BASE-T: CAT6A, CAT6, CAT5e (100m max)
Power Requirements	6 W max, optics not included	9.0 W Maximum @ +12 V	16.3 W Maximum @ +12 V
Dimensions (excluding bracket)	Standard OCP3.0 Small Form Factor 2.99 x 4.53 in (7.6 x 11.5 cm)	Standard OCP3.0 Small Form Factor 2.99 x 4.53 in (7.6 x 11.5 cm)	Standard OCP3.0 Small Form Factor 2.99 x 4.53 in (7.6 x 11.5 cm)
Operating Temperature	23 - 149 deg. F (-5 - 65 deg. C)	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	376	376	223
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps	10 Gbps/1 Gbps
Duplex Modes	Full only	Full only	Full only
Indicator Lights	Link Activity	Link Activity	Link Activity
Standards Conformance	PCI Express 3.0 SFF-8431 IEEE 802.3ae OCP NIC 3.0	IEEE 802.3 PCI Express 3.0	IEEE 802.3 PCI Express 3.0
Regulatory and Safety	Safety Compliance - UL/ CSA 62368-1: 2014 2nd Edition - EN 62368 (European Union) - IEC 62368 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 32 - Radiated & Conducted Emissions (International) - EN55032-2015 - Radiated & Conducted Emissions (European Union) - EN55035: 2017 - (Immunity) (European Union) - REACH, WEEE, RoHS Directives (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZ CISPR - Radiated & Conducted Emissions (Australia/New Zealand) - KN32 -Radiated & Conducted Emissions (Korea)KN35 - (Immunity) (Korea)		

Feature	Intel® Ethernet 10G 2P X710-T2L-t Adapter	Intel® Ethernet 10G 4P X710-T4L-t Adapter
Bus Connector	PCI Express 3.0	PCI Express 3.0
Bus Speed	x8	x8
Transmission Mode/Connector	RJ45 BASE-T Connector	RJ45 BASE-T Connector
Cabling	10GBASE-T: CAT6A (100m max), CAT6 (55m max) 1000BASE-T: CAT6A, CAT6, CAT5e (100m max)	10GBASE-T: CAT6A (100m max), CAT6 (55m max) 1000BASE-T: CAT6A, CAT6, CAT5e (100m max)
Power Requirements	9.6 W Maximum @ +12 V	14.2 W Maximum @ +12 V
Dimensions (excluding bracket)	2.70 x 6.74 in (6.86 x 17.12 cm)	2.70 x 6.63 in (6.86 x 16.84 cm)
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	356	237
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps
Duplex Modes	Full only	Full only
Indicator Lights	Link Activity	Link Activity
Standards Conformance	IEEE 802.3 PCI Express 3.0	IEEE 802.3 PCI Express 3.0
Regulatory and Safety	Safety Compliance - UL/ CSA 62368-1: 2014 2nd Edition - EN 62368 (European Union) - IEC 62368 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 32 - Radiated & Conducted Emissions (International) - EN55032-2015 - Radiated & Conducted Emissions (European Union) - EN55035: 2017 - (Immunity) (European Union) - REACH, WEEE, RoHS Directives (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZ CISPR - Radiated & Conducted Emissions (Australia/New Zealand) - KN32 -Radiated & Conducted Emissions (Korea)KN35 - (Immunity) (Korea)	

Feature	Intel® Ethernet Converged Network Adapter X710-T	Intel® Ethernet Converged Network Adapter X710	Ethernet Server Adapter X710-DA2 for OCP
Bus Connector	PCI Express 3.0	PCI Express 3.0	PCI Express 3.0
Bus Speed	x8	x8	x8
Transmission Mode/Connector	10GBase-T/RJ-45	SFP+	SFP+
Cabling	10GBase-T (Category 6A)	Twinax	Direct Attach

Feature	Intel® Ethernet Converged Network Adapter X710-T	Intel® Ethernet Converged Network Adapter X710	Ethernet Server Adapter X710-DA2 for OCP
		10GBase-SR/LR	10GBASE-SR
Power Requirements	8.53 W (idle) @ 12V Main	6.7 Watts (maximum) @ 12 V	3.08 Watts (max) @ 5V Main
Dimensions (excluding bracket)	6.578 x 4.372 in 16.708 x 11.107 cm	6.578 x 4.372 in 16.708 x 11.107 cm	2.67 x 4.59 in 6.78 x 11.658 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)	41 - 131 deg. F (5 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)
MTBF	493 years	491 years	1276 years
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps	10 Gbps
Duplex Modes	Full only	Full Only	Full only
Indicator Lights	Link/Activity 1Gig/10Gig	Link/Activity 1Gig/10Gig	Link/Activity 1Gig/10Gig
Standards Conformance	PCI Express 3.0 SFF-8431 IEEE 802.3z IEEE 802.3ae	PCI Express 3.0 SFF-8431 IEEE 802.3z IEEE 802.3ae	PCI Express 3.0 SFF-8431 IEEE 802.3z IEEE 802.3ae
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)		

Feature	Intel® Ethernet 10G 2P X540-t Adapter	Intel® Ethernet 10G 2P X520 Adapter	Intel® Ethernet 10G 2P X550-t Adapter
Bus Connector	PCI Express 2.0	PCI Express 2.0	PCI Express 3.0
Bus Speed	x8	x8	x8
Transmission Mode/Connector	10GBase-T/RJ-45	Twinaxial copper/SFP+ 10GBase-SR/LR	10GBase-T/RJ-45
Cabling	10GBase-T (Category 6A)	10 Gigabit Ethernet over SFP+ Direct Attach Copper (10GSFP+Cu) 10GBASE-SR/LR	10GBase-T (Category 6A)
Power	15 W Maximum @ +12 V	6.2 W Maximum @ +3.3 V	13W Maximum @ +12 V

Feature	Intel® Ethernet 10G 2P X540-t Adapter	Intel® Ethernet 10G 2P X520 Adapter	Intel® Ethernet 10G 2P X550-t Adapter
Requirements			
Dimensions (excluding bracket)	5.7 x 2.7 in 14.5 x 6.9 cm	5.7 x 2.7 in 14.5 x 6.9 cm	5.13 x 2.7 in 13.0 x 6.9 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	108 years	83.9 years	127 years
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps	10 Gbps/1 Gbps
Duplex Modes	Full only	Full only	Full only
Indicator Lights	<i>Two per port:</i> Link and Activity	<i>Two per port:</i> Link and Activity	Link Activity
Standards Conformance	IEEE 802.1p IEEE 802.1Q IEEE 802.3an IEEE 802.3ac IEEE 802.3ad IEEE 802.3an IEEE 802.3x ACPI v1.0 PCI Express 2.0	IEEE 802.1p IEEE 802.1Q IEEE 802.3an IEEE 802.3ac IEEE 802.3ad IEEE 802.3x ACPI v1.0 PCI Express 2.0	IEEE 802.1p IEEE 802.1Q IEEE 802.3an IEEE 802.3ac IEEE 802.3ad IEEE 802.3x ACPI v1.0 PCI Express 3.0
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)		

Intel® 10 Gigabit Network Mezzanine Card Specifications

Feature	Intel® Ethernet X520 10GbE Dual Port KX4-KR Mezz
Bus Connector	PCI Express 2.0
Bus Speed	x8
Power Requirements	7.4 Watts (maximum) @ 3.3 V
Dimensions	3.65 x 3.3 in.

Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	147 years
Available Speeds	10 Gbps/1 Gbps
Duplex Modes	Full only
Standards Conformance	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0 PCI Express 2.0
Regulatory and Safety	Safety Compliance • UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) • EN 60 950 (European Union) • IEC 60 950 (International) EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55022-1998 - Radiated & Conducted Emissions (European Union) • EN55024 - 1998 - (Immunity) (European Union) • CE - EMC Directive (89/336/EEC) (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) • MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)

Intel® 10GbE Network Daughter Cards (NDC) Specifications

Feature	Intel® Ethernet 10G 4P X550/I350 rNDC	Intel® Ethernet 10G 4P X550 rNDC
Bus Connector	PCI Express 3.0	PCI Express 3.0
Bus Speed	x8	x8
Transmission Mode/Connector	Twisted copper/RJ-45	Twisted copper/RJ-45
Cabling	Cat 6A (10 Gbps)/Cat 5e (1 Gbps)	Cat 6A
Power Requirements	15.39 Watts (max) @12 V	33.6 Watts (maximum) @ 12 V
Dimensions	4.34 x 4.012 in 11.04 x 10.19 cm	4.37 x 5.86 in 11.10 x 14.883 cm
Operating Temperature	60° F	60° F
MTBF at 55°C	445	436
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps

Duplex Modes	Full only	Full only
Standards Conformance	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0 PCI Express 3.0	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0 PCI Express 3.0
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition-CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)	

Feature	Intel® Ethernet 10G 4P X540/I350 rNDC	Intel® Ethernet 10G 4P X520/I350 rNDC	Intel® Ethernet 10G 2P X520-kbNDC
Bus Connector	PCI Express 2.0	PCI Express 2.0	PCI Express 2.0
Bus Speed	x8	x8	x8
Transmission Mode/Connector	Twisted copper/RJ-45	SFP+	Copper/Backplane
Cabling	Cat 6A (10 Gbps)/Cat 5e (1 Gbps)	SFP+ SR/DA	10GBase-KR and 1000Base-KX
Power Requirements	5.5 Watts (maximum) @ 3.3 V	10.1 Watts (maximum) @ 12 V	0.6 Watts @ 3.3 V (AUX), 6.3 Watts @ 1.2 V (VCORE)
Dimensions	3.93 x 3.67 in.	4.3 x 3.7 in.	3.0 x 2.5 in.
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	68 years	65 years	147 years
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps	10 Gbps/1 Gbps
Duplex Modes	Full only	Full only	Full only
Standards Conformance	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ap IEEE 802.3x ACPI v1.0

	PCI Express 2.0a	PCI Express 2.0a	PCI Express 2.0a
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)		

Feature	Intel® Ethernet 10G 4P X710-k bNDC	Intel® Ethernet 10G 4P X710/I350 rNDC	Intel® Ethernet 10G 4P X710 SFP+ rNDC
Bus Connector	Dell EMC bNDC	PCI Express 3.0	PCI Express 3.0
Bus Speed	x8	x8	x8
Transmission Mode/Connector	KX/KR	SFP+	SFP+
Cabling	Backplane	Twinax 10GBase-SR/LR	Twinax 10GBase-SR/LR
Power Requirements	3.3 Watts @ 3.3 V (AUX), 12.6 Watts @ 12 V (AUX)	10.7 Watts Maximum @ +12 V	9.5 Watts Maximum @ +12 V
Dimensions	3.000x2.449 in 7.62x6.220cm	4.331x3.661 in 11.0x9.298 cm	4.331x3.661 in 11.0x9.298 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	828 years	108 years	505 years
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps	10 Gbps/1 Gbps
Duplex Modes	Full only	Full only	Full only
Indicator Lights	None	Link/Activity Speed	Link/Activity Speed
Standards Conformance	PCI Express 3.0 IEEE 802.3ap	PCI Express 3.0 SFF-8431 IEEE 802.3z IEEE 802.3ae	PCI Express 3.0 SFF-8431 IEEE 802.3z IEEE 802.3ae
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance		

	• FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55022-1998 - Radiated & Conducted Emissions (European Union) • EN55024 - 1998 - (Immunity) (European Union) • CE - EMC Directive (89/336/EEC) (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) • MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)
--	---

Intel® Gigabit Network Adapter Specifications

Feature	Intel® Ethernet 1G 4P I350-t OCP
Bus Connector	PCI Express 2.1
Bus Speed	x4
Transmission Mode/Connector	1GBase-T/RJ-45
Cabling	Cat 5e
Power Requirements	25.2 W Maximum @ +12 V
Dimensions (excluding bracket)	Standard OCP3.0 Small Form Factor 2.99 x 4.53 in (7.6 x 11.5 cm)
Operating Temperature	23 - 149 deg. F (-5 to 65 deg. C))
MTBF at 55°C	335 years
Available Speeds	10/100/1000 Mbps auto-negotiate
Duplex Modes	Full or half at 10/100 Mbps; full only at 1000 Mbps
Standards Conformance	IEEE 802.3 IEEE 802.3ab IEEE 802.3u PCI Express 2.1 OCP NIC 3.0
Indicator Lights	<i>Two per port:</i> Link and Activity
Regulatory and Safety	Safety Compliance • UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) • EN 60 950 (European Union) • IEC 60 950 (International) EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55022-1998 - Radiated & Conducted Emissions (European Union) • EN55024 - 1998 - (Immunity) (European Union)

	• CE - EMC Directive (89/336/EEC) (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan) • AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) • MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)
--	--

Feature	Intel® Gigabit 2P I350-t Adapter and Intel® Gigabit 4P I350-t Adapter
Bus Connector	PCI Express 2.0
Bus Speed	x4
Transmission Mode/Connector	Twisted copper/RJ-45
Cabling	1000Base-T (Category 3 or Category 5)
Power Requirements	<i>Intel® Gigabit 2P I350-t Adapter: 4.8 Watts @ 12 V</i> <i>Intel® Gigabit 4P I350-t Adapter: 6.0 Watts @ 12 V</i>
Dimensions (excluding bracket)	5.3 x 2.7 in. 13.5 x 6.9 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	68 years
Available Speeds	10/100/1000 auto-negotiate
Duplex Modes	Full or half at 10/100 Mbps; full only at 1000 Mbps
Standards Conformance	IEEE 802.1p IEEE 802.1Q IEEE 802.3ab IEEE 802.3ac IEEE 802.3ad IEEE 802.3az IEEE 802.3u IEEE 802.3x IEEE 802.3z ACPI v1.0 PCI Express 2.0
Indicator Lights	<i>Two per port:</i> Activity and Speed
Regulatory and Safety	Safety Compliance • UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) • EN 60 950 (European Union) • IEC 60 950 (International) EMC Compliance • FCC Part 15 - Radiated & Conducted Emissions (USA) • ICES-003 - Radiated & Conducted Emissions (Canada) • CISPR 22 - Radiated & Conducted Emissions (International) • EN55022-1998 - Radiated & Conducted Emissions (European Union) • EN55024 - 1998 - (Immunity) (European Union) • CE - EMC Directive (89/336/EEC) (European Union) • VCCI - Radiated & Conducted Emissions (Japan) • CNS13438 - Radiated & Conducted Emissions (Taiwan)

	- AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)
--	---

Intel® Gigabit Network Mezzanine Card Specifications

Feature	Intel® Gigabit 4P I350-t Mezz
Bus Connector	PCI Express 2.0
Bus Speed	x4
Power Requirements	3.425 Watts (maximum) @ 3.3 V
Dimensions	3.65 x 3.3 in.
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)
MTBF at 55°C	108 years
Available Speeds	Full only at 1000 Mbps
Duplex Modes	Full at 1000 Mbps
Standards Conformance	IEEE 802.1p IEEE 802.1Q IEEE 802.3ab IEEE 802.3ac IEEE 802.3ad IEEE 802.3x ACPI v1.0 PCI Express 2.0
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)

Intel® Gigabit Network Daughter Cards Specifications

Feature	Intel® Gigabit 4P X710/I350 rNDC	Intel® Ethernet Gigabit 4P X550/I350 rNDC	Intel® Gigabit 4P I350-t rNDC
Bus Connector	PCI Express 2.0	PCI Express 2.0	PCI Express 2.0

Bus Speed	x8	x8	x8
Transmission Mode/Connector	Twisted copper/RJ-45	Twisted copper/RJ-45	Twisted copper/RJ-45
Cabling	Cat-5e	Cat-5e	Cat-5e
Power Requirements	10.7W Maximum @ +12 V	15.39 W (max) @ +12 V	5.5W (max)@ +3.3 V
Dimensions (excluding bracket)	4.331 x 3.661 in 11.007 x 9.298 cm	5.86 x 4.35 in 14.882 x 11.04 cm	5.33 x 2.71 in 13.54 x 6.59 cm
Operating Temperature	32 - 131 deg. F (0 - 55 deg. C)	32 - 60 deg. F (0 - 16 deg C.)	32 - 60 deg. F (0 - 16 deg. C)
MTBF at 55°C	108 years	251 years	117 years
Available Speeds	10 Gbps/1 Gbps	10 Gbps/1 Gbps	1 Gbps
Duplex Modes	Full only	Full only	Full only
<u>Standards Conformance</u>	PCI Express 2.1 IEEE 802.3i IEEE 802.3ab IEEE 802.3u IEEE 802.3ad IEEE 802.3az	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0 PCI Express 2.1	IEEE 802.1p IEEE 802.1Q IEEE 802.3ac IEEE 802.3ad IEEE 802.3ae IEEE 802.3x ACPI v1.0 PCI Express 2.1
Regulatory and Safety	Safety Compliance - UL 60950 Third Edition- CAN/CSA-C22.2 No.60950-00 (USA/Canada) - EN 60 950 (European Union) - IEC 60 950 (International) EMC Compliance - FCC Part 15 - Radiated & Conducted Emissions (USA) - ICES-003 - Radiated & Conducted Emissions (Canada) - CISPR 22 - Radiated & Conducted Emissions (International) - EN55022-1998 - Radiated & Conducted Emissions (European Union) - EN55024 - 1998 - (Immunity) (European Union) - CE - EMC Directive (89/336/EEC) (European Union) - VCCI - Radiated & Conducted Emissions (Japan) - CNS13438 - Radiated & Conducted Emissions (Taiwan) - AS/NZS3548 - Radiated & Conducted Emissions (Australia/New Zealand) - MIC notice 1997-41, EMI and MIC notice 1997-42 - EMS (Korea)		

Standards

- IEEE 802.1p: Priority Queuing (traffic prioritizing) and Quality of Service levels
- IEEE 802.1Q: Virtual LAN identification
- IEEE 802.3ab: Gigabit Ethernet over copper
- IEEE 802.3ac: Tagging
- IEEE 802.3ad: SLA (FEC/GEC/Link Aggregation - static mode)
- IEEE 802.3ad: Dynamic mode
- IEEE 802.3ae: 10 Gbps Ethernet
- IEEE 802.3an: 10GBase-T 10 Gbps Ethernet over unshielded twisted pair
- IEEE 802.3ap: Backplane Ethernet
- IEEE 802.3u: Fast Ethernet
- IEEE 802.3x: Flow Control
- IEEE 802.3z: Gigabit Ethernet over optical fiber
- ACPI: Advanced Configuration and Power Management
- PCI Express: system bus specification: 32/64-bit, x1, x2, x4, x8, x16

More information on IEEE 802 standards available at <http://www.ieee802.org>.

IEEE 802.3ac VLANs:

VLANs require VLAN-capable switches either implicit (switch only) or explicit (IEEE 802.3ac). IEEE 802.3ac VLANs allow multiple VLANs per adapter since both the switch and adapter use a tag in the packet header to sort VLANs.

Intel gigabit and 10 gigabit network adapters fully support implicit and explicit VLANs.

X-UEFI Attributes

This section contains information about X-UEFI attributes and their expected values.

List of Multi-controller Devices

The adapters listed below contain more than one controller. On these adapters, configuring controller based settings will not affect all ports. Only ports bound to the same controller will be affected.

The following settings apply to all ports on a given controller:

- Virtualization Mode
- NParEP Mode
- PCI Virtual Functions Advertised

Multi-controller Devices	Number of controllers on the device	Controller 1	Controller 2
Intel® Ethernet 10G 4P X520/I350 rNDC Intel® Gigabit 4P X520/I350 rNDC	2	10G Ports 1 and 2	1G Ports 3 and 4
Intel® Ethernet 10G 4P X540/I350 rNDC Intel® Gigabit 4P X540/I350 rNDC	2	10G Ports 1 and 2	1G Ports 3 and 4
Intel® Ethernet 10G 4P X550 rNDC	2	10G Ports 1 and 2	10G Ports 3 and 4
Intel® Ethernet 10G 4P X550/I350 rNDC Intel® Gigabit 4P X550/I350 rNDC	2	10G Ports 1 and 2	1G Ports 3 and 4
Intel® Ethernet 10G 4P X710/I350 rNDC Intel® Ethernet 10G X710 rNDC Intel® Gigabit 4P X710/I350 rNDC	2	10G Ports 1 and 2	1G Ports 3 and 4

Table of X-UEFI Attributes

Display Name	X-UEFI Name	Supported Adapters								User Configurable	User Configurable Values	Values that can be displayed	Dependencies for Values	I/O Identity Optimization (iDRAC 8/9)	Information
		I350	X520	X540	X550	X710	XL710	XXV710	E810						
Virtualization Mode	VirtualizationMode	X	X	X	X	X	X	X	X	Yes	None/NPAR/SR-IOV/NPAR + SR-IOV	None/NPAR/SR-IOV/NPAR + SR-IOV		No	Specifies the virtualization mode setting of the controller. "NPAR" and "NPAR + SR-IOV" are only supported on X710 and XXV710 devices. They are not supported on client operating systems. The attribute setting applies to all ports on a given controller.
Number of Virtual Functions	NumberVFSupported	X	X	X	X	X	X	X	X	No		0-256		No	The number of virtual functions supported on this port.

Display Name	X-UEFI Name	Supported Adapters								User Configurable	User Configurable Values	Values that can be displayed	Dependencies for Values	I/O Identity Optimization (iDRAC 8/9)	Information
		I350	X520	X540	X550	X710	XL710	XXV710	E810						
Supported															
Partition State Interpretation	PartitionStateInterpretation					X		X		No		Variable/Fixed		No	Describes how partitioning is implemented and how the PartitionState attribute is used in the controller. Fixed is the only value used.
RDMA Support	RDMASupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether any RDMA protocol is supported by the controller. Unavailable is the only value used.
LLDP Agent	INTEL_LLDPAgent					X	X	X	X	Yes	Disabled/Enabled	Disabled/Enabled		No	Persistently enables or disables firmware's LLDP Agent. Note that disabling firmware's LLDP Agent also disables DCB functionality. Disabling LLDP Agent allows LLDP packets from the switch to pass unobstructed to the OS. Some OS layer LLDP agents and software defined network layer LLDP agents need these packets to function properly. Note that this option is not available in NPAR mode.
SR-IOV Support	SRIOVSupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether SR-IOV capability is supported.
VF Allocation Basis	VFAllocBasis	X	X	X	X	X	X	X	X	No		Device/Port		No	Defines the domain in which Virtual Functions are allocated. Port is the only value used.
VF Allocation Multiple	VFAllocMult	X	X	X	X	X	X	X	X	No		1-255		No	Virtual Functions must be allocated to a port in multiples of this number.
NParEP Mode	NParEP					X		X		Yes	Disabled/Enabled	Disabled/Enabled	VirtualizationMode - NPAR or NPAR + SR-IOV	No	NParEP mode enables more than 8 partitions on the device. It must not be enabled if the system and OS do not support devices with more than 8 PCI Physical Functions. The attribute setting applies to all ports on a given controller.
Partition n Maximum TX Bandwidth	MaxBandwidth[Partition:n]					X		X		Yes	1-100	1-100		Yes	Represents the maximum transmit bandwidth of the partition as a percentage of the full physical port link speed. The Maximum Bandwidth range is 1-100 percent for each enabled partition. If the remotely configured Maximum Bandwidth value on partition n is lower than Minimum Bandwidth on partition n, Minimum Bandwidth will be used.
Partition n Minimum TX Bandwidth	MinBandwidth[Partition:n]					X		X		Yes	1-100	1-100		Yes	Represents the minimum transmit bandwidth of the partition as a percentage of the full physical port link speed. The Minimum Bandwidth range is 1-100 percent for each enabled partition. The minimum bandwidth across all enabled partitions on a port must add up to 100%. If the remotely configured Minimum Bandwidth percentages do not add up to 100, the firmware will automatically normalize them to 100.
Boot LUN	FirstTgtBootLun	X	X	X	X	X	X	X	X	Yes	0-255	0-255		Yes	Specifies the boot Logical Unit Number (LUN) on the first iSCSI storage target.

Display Name	X-UEFI Name	Supported Adapters								User Configurable	User Configurable Values	Values that can be displayed	Dependencies for Values	I/O Identity Optimization (iDRAC 8/9)	Information
		I350	X520	X540	X550	X710	XL710	XXV710	E810						
CHAP Secret	FirstTgtChapPwd	X	X	X	X	X	X	X	X	Yes	string	string		Yes	Specifies the Challenge-Handshake Authentication Protocol secret (CHAP password) of the first iSCSI storage target. The string value is limited to alphanumeric characters, '.' (dot), ':' (colon), and '-' (dash).
IP Address	FirstTgtIpAddress	X	X	X	X	X	X	X	X	Yes	X.X.X.X	X.X.X.X		Yes	Specifies the IP address of the first iSCSI target.
iSCSI Name	FirstTgtIscsiName	X	X	X	X	X	X	X	X	Yes	string	string		Yes	Specifies the iSCSI Qualified Name (IQN) of the first iSCSI storage target. The string value is limited to alphanumeric characters, '.' (dot), ':' (colon), and '-' (dash).
TCP Port	FirstTgtTcpPort	X	X	X	X	X	X	X	X	Yes	1024-65535	1024-65535		Yes	Specifies the TCP Port number of the first iSCSI target.
CHAP ID	IscsiInitiatorChapId	X	X	X	X	X	X	X	X	Yes	string	string	ChapAuthEnable - Enabled	Yes	Specifies the first iSCSI storage target Challenge-Handshake Authentication Protocol (CHAP) ID. The string value is limited to alphanumeric characters, '.' (dot), ':' (colon), and '-' (dash).
CHAP Authentication	ChapAuthEnable	X	X	X	X	X	X	X	X	Yes	Enabled/Disabled	Enabled/Disabled		No	Enables the initiator to use CHAP authentication when connecting to the iSCSI target.
TCP/IP Parameters via DHCP	TcpIpViaDHCP	X	X	X	X	X	X	X	X	Yes	Disabled/Enabled	Disabled/Enabled		No	Controls the source of the initiator IP address, DHCP or static assignment. This option is specific to IPv4.
IP Version	IpVer	X	X	X	X	X	X	X	X	No		IPv4		No	Controls whether IPv4 or IPv6 network addressing will be used for iSCSI initiator and targets. Currently only IPv4 is supported.
CHAP Mutual Authentication	ChapMutualAuth	X	X	X	X	X	X	X	X	Yes	Disabled/Enabled	Disabled/Enabled	ChapAuthEnable - Enabled	No	Enables or disables CHAP Mutual Authentication. To use mutual CHAP authentication, you must specify an initiator secret on the Initiator Parameters page and configure that secret on the target.
iSCSI Parameters via DHCP	IscsiViaDHCP	X	X	X	X	X	X	X	X	Yes	Disabled/Enabled	Disabled/Enabled	TcpIpViaDHCP - Enabled	No	Enables the acquisition of iSCSI target parameters from DHCP.
iSCSI Name	IscsiInitiatorName	X	X	X	X	X	X	X	X	Yes	string	string		Yes	Specifies the initiator's iSCSI Qualified Name (IQN). The attribute setting applies to all ports on a given controller. It is recommended to use the same IscsiInitiatorName on all ports for a given device.
CHAP Secret	IscsiInitiatorChapPwd	X	X	X	X	X	X	X	X	Yes	string	string	ChapAuthEnable - Enabled	Yes	Sets the iSCSI initiator Challenge-Handshake Authentication Protocol (CHAP) secret (password). The string value is limited to alphanumeric characters, '.' (dot), ':' (colon), and '-' (dash).
Default Gateway	IscsiInitiatorGateway	X	X	X	X	X	X	X	X	Yes	X.X.X.X	X.X.X.X	TcpIpViaDHCP - Disabled	Yes	Specifies the IP address of the default Gateway used by the iSCSI initiator.
IP Address	IscsiInitiatorIpAddress	X	X	X	X	X	X	X	X	Yes	X.X.X.X	X.X.X.X	TcpIpViaDHCP - Disabled	Yes	Specifies the IP address of the iSCSI initiator.

Display Name	X-UEFI Name	Supported Adapters								User Configurable	User Configurable Values	Values that can be displayed	Dependencies for Values	I/O Identity Optimization (iDRAC 8/9)	Information
		I350	X520	X540	X550	X710	XL710	XXV710	E810						
Subnet Mask	IscsiInitiatorSubnet	X	X	X	X	X	X	X	X	Yes	X.X.X.X	X.X.X.X	TcpIpViaDHCP - Disabled	Yes	Specifies the IPv4 Subnet Mask of the iSCSI initiator.
Blink LEDs	BlinkLeds	X	X	X	X	X	X	X	X	Yes	0-15	0-15		No	Specifies the number of seconds the LEDs on the physical network port should blink to assist with port identification.
Virtual MAC Address	VirtMacAddr	X	X	X	X	X	X	X	X	Yes	XX:XX:XX:XX:XX:XX	XX:XX:XX:XX:XX:XX		Yes	Sets the programmatically assignable MAC address for port.
FlexAddressing	FlexAddressing	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether the Dell FlexAddressing feature is supported.
iSCSI Boot Support	iSCSIBootSupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether iSCSI Boot is supported.
iSCSI Dual IP Version Support	iSCSIDualIPVersionSupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates support for simultaneous IPv4 and IPv6 configurations of iSCSI initiator and iSCSI primary and secondary targets. Unavailable is the only value used.
iSCSI Offload Support	iSCSIOffloadSupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether the iSCSI Offload capability is supported. Unavailable is the only value used.
Link Status	LinkStatus	X	X	X	X	X	X	X	X	No		Disconnected/Connected		No	Reports the physical link status of the network ports reported by the controller.
MAC Address	MacAddr	X	X	X	X	X	X	X	X	No		XX:XX:XX:XX:XX:XX		No	Reports the permanent MAC address assigned during manufacturing.
NIC Partitioning Support	NicPartitioningSupport					X		X		No		Available/Unavailable		No	Indicates whether NIC Partitioning capability is supported.
OS BMC Management Pass Through	OSBMCManagementPassThrough					X	X	X	X	No		Available/Unavailable		No	Indicates whether OS-BMC Management Pass Through capability is supported.
PCI Device ID	PCIDeviceID	X	X	X	X	X	X	X		No		XXXX		No	Reports the PCI Device ID of the controller.
PXE Boot Support	PXEBootSupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether PXE Boot capability is supported.
RX Flow Control	RXFlowControl					X	X	X	X	No		Available/Unavailable		No	Indicates whether Receive (RX) Flow control capability is supported. Unavailable is the only value used.
TOE Support	TOESupport	X	X	X	X	X	X	X	X	No		Available/Unavailable		No	Indicates whether TCP/IP Offload Engine capability is supported. Unavailable is the only value used.
TX Bandwidth Control Maximum	TXBandwidthControlMaximum							X	X	No		Available/Unavailable		No	Indicates whether Transmit (TX) Bandwidth Control Maximum capability is supported.
TX Flow Control	TXFlowControl					X	X	X	X	No		Available/Unavailable		No	Indicates whether Transmit (TX) Flow Control capability is supported. Unavailable is the only value used.
Legacy Boot Protocol	LegacyBootProto	X	X	X	X	X	X	X	X	Yes	None/PXE/iSCSI Primary/iSCSI Secondary	None/PXE/iSCSI Primary/iSCSI Secondary		No	Selects a boot protocol to be used in legacy BIOS (non-UEFI) boot mode.

Display Name	X-UEFI Name	Supported Adapters								User Configurable	User Configurable Values	Values that can be displayed	Dependencies for Values	I/O Identity Optimization (iDRAC 8/9)	Information
		I350	X520	X540	X550	X710	XL710	XXV710	E810						
Legacy Virtual LAN ID	VLANid	X	X	X	X	X	X	X	X	Yes	0-4094	0-4094		No	Specifies the ID (tag) to be used for PXE VLAN Mode. The VLAN ID must be in the range from 0 to 4094. PXE VLAN is disabled if value is set to 0.
Wake On LAN	WakeOnLan	X	X	X	X	X	X	X	X	Yes	Disabled/Enabled/'N/A'	Disabled/Enabled/'N/A'		No	Enables power on of the system via LAN. Note that configuring Wake on LAN in the operating system does not change the value of this setting, but does override the behavior of Wake on LAN in OS controlled power states.
Link Speed	LnkSpeed	X	X	X	X	X	X	X	X	*Yes	Auto Negotiated/1000 Mbps Full/10 Mbps Half/10 Mbps Full/100 Mbps Half/100 Mbps Full	Auto Negotiated/1000 Mbps Full/10 Mbps Half/10 Mbps Full/100 Mbps Half/100 Mbps Full		No	Specifies the port link speed to be used when booting the selected protocol. *The attribute is only writable on 1G (I350) devices.
PCI Virtual Functions Advertised	NumberVFAdvertised	X	X	X	X	X	X	X	X	Yes	I350: 1-8, X520/X540/X550: 1-64, X710/XL710/XXV710: 0-127	I350: 1-8, X520/X540/X550: 1-64, X710/XL710/XXV710: 0-127	VirtualizationMode - SR-IOV	No	Specifies the number of PCI Virtual Functions (VFs) to be advertised in non-NPAR mode. Available values vary between product families. On I350, X520, X540, and X550 based devices, the value represents the total number of PCI VFs that will be shared across all ports on a given controller. On all other devices, the value represents the number of PCI VFs that will be dedicated to each port.
PCI Virtual Functions Advertised	NumberVFAdvertised					X		X		Yes	0-128	0-128	VirtualizationMode -NPAR+ SR-IOV	No	Specifies the number of PCI Virtual Functions (VFs) to be advertised on this port in NPAR mode. This attribute is present only in HII Browser. Virtual Functions in NPAR mode can be assigned only to the first partition on a port. Use the VFDistribution attribute for remote configuration.
Number of PCI Physical Functions Currently Enabled per Port	NumberPCIFunctionsEnabled					X	X	X	X	No		1-8		No	Reports the number of PCI Physical Functions currently enabled on this port.
Number of PCI Physical Functions Supported	NumberPCIFunctionsSupported					X	X	X	X	No		1-8		No	Reports the number of PCI Physical Functions supported on this port. This value may change depending on the support and configuration of NParEP.
Partition n	PartitionState[Partition:n]					X		X		No		Enabled/Disabled		No	Reports the current enablement state of the partition.
Virtual MAC Address	VirtMacAddr[Partition:n]					X		X		Yes	XX:XX:XX:XX:XX:XX	XX:XX:XX:XX:XX:XX			Reports the programmatically assignable MAC address for partition.
MAC Address	MacAddr[Partition:n]					X		X		No		XX:XX:XX:XX:XX:XX		No	Reports the permanent MAC address assigned during manufacturing.
NIC Mode	NicMode[Partition:n]					X		X		No		Disabled/Enabled		No	Specifies use of the partition for L2-Ethernet traffic. Enabled is the only value used.

Display Name	X-UEFI Name	Supported Adapters								User Configurable	User Configurable Values	Values that can be displayed	Dependencies for Values	I/O Identity Optimization (iDRAC 8/9)	Information
		I350	X520	X540	X550	X710	XL710	XXV710	E810						
PCI Device ID	PCIDeviceID[Partition:n]					X		X		No		XXXX		No	Reports the PCI Device ID of the partition.
Port Number	PortNumber[Partition:n]					X		X		No		1-4		No	Reports the port to which the partition belongs, where n is the number of the partitions.
VF Distribution	VFDistribution					X		X		Yes	X:0:0:0.....0:0 (The number of zeros depends on the number of partitions currently enabled on the port)	X:0:0:0.....0:0 (The number of zeros depends on the number of partitions currently enabled on the port)	VirtualizationMode - NPAR + SR-IOV	No	Defines the distribution of VFs to PFs within the domain specified by VFAllocBasis. A value appears in the colon-separated list for each Physical Function that can be potentially present within the allocation domain. Values in the list from left to right apply to function numbers in the domain from least to greatest.
Permit Total Port Shutdown	PermitTotalPortShutdown					X	X	X	X	Yes	Enabled/Disabled	Enabled/Disabled		No	Specifies whether or not to allow the port to be completely disabled when a Port Down command is received from the Host OS. Note: Use with caution, port shutdown will halt all operations configured on the port including WakeOnLAN and shared LOM.
Forward Error Correction	ForwardErrorCorrection							X	X	Yes	Auto, Fire Code, RS-FEC, Disabled	Auto, Fire Code, RS-FEC, Disabled			Specifies the forward error correction mode.

Legal Disclaimers

Software License Agreement

INTEL SOFTWARE LICENSE AGREEMENT (Final, License)

IMPORTANT - READ BEFORE COPYING, INSTALLING OR USING

Do not use or load this software and any associated materials (collectively, the "Software") until you have carefully read the following terms and conditions. By loading or using the Software, you agree to the terms of this Agreement. If you do not wish to so agree, do not install or use the Software.

LICENSES

Please Note:

- If you are a network administrator, the "Site License" below shall apply to you.
- If you are an end user, the "Single User License" shall apply to you.

SITE LICENSE. You may copy the Software onto your organization's computers for your organization's use, and you may make a reasonable number of back-up copies of the Software, subject to these conditions:

1. **This Software is licensed for use only in conjunction with Intel component products. Use of the Software in conjunction with non-Intel component products is not licensed hereunder.**
2. You may not copy, modify, rent, sell, distribute or transfer any part of the Software except as provided in this Agreement, and you agree to prevent unauthorized copying of the Software.
3. You may not reverse engineer, decompile, or disassemble the Software.
4. You may not sublicense or permit simultaneous use of the Software by more than one user.
5. The Software may include portions offered on terms in addition to those set out here, as set out in a license accompanying those portions.

SINGLE USER LICENSE. You may copy the Software onto a single computer for your personal, noncommercial use, and you may make one back-up copy of the Software, subject to these conditions:

1. **This Software is licensed for use only in conjunction with Intel component products. Use of the Software in conjunction with non-Intel component products is not licensed hereunder.**
2. You may not copy, modify, rent, sell, distribute or transfer any part of the Software except as provided in this Agreement, and you agree to prevent unauthorized copying of the Software.
3. You may not reverse engineer, decompile, or disassemble the Software.
4. You may not sublicense or permit simultaneous use of the Software by more than one user.
5. The Software may include portions offered on terms in addition to those set out here, as set out in a license accompanying those portions.

OWNERSHIP OF SOFTWARE AND COPYRIGHTS. Title to all copies of the Software remains with Intel or its suppliers. The Software is copyrighted and protected by the laws of the United States and other countries, and international treaty provisions. You may not remove any copyright notices from the Software. Intel may make changes to the Software, or to items referenced therein, at any time without notice, but is not obligated to support or update the Software. Except as otherwise expressly provided, Intel grants no express or implied right under Intel patents, copyrights, trademarks, or other intellectual property rights. You may transfer the Software only if the recipient agrees to be fully bound by these terms and if you retain no copies of the Software.

LIMITED MEDIA WARRANTY. If the Software has been delivered by Intel on physical media, Intel warrants the media to be free from material physical defects for a period of ninety days after delivery by Intel. If such a defect is found, return the media to Intel for replacement or alternate delivery of the Software as Intel may select.

EXCLUSION OF OTHER WARRANTIES. EXCEPT AS PROVIDED ABOVE, THE SOFTWARE IS PROVIDED "AS IS" WITHOUT ANY EXPRESS OR IMPLIED WARRANTY OF ANY KIND INCLUDING WARRANTIES OF MERCHANTABILITY, NONINFRINGEMENT, OR FITNESS FOR A PARTICULAR PURPOSE. Intel does not warrant or assume responsibility for the accuracy or completeness of any information, text, graphics, links or other items contained within the Software.

LIMITATION OF LIABILITY. IN NO EVENT SHALL INTEL OR ITS SUPPLIERS BE LIABLE FOR ANY DAMAGES WHATSOEVER (INCLUDING, WITHOUT LIMITATION, LOST PROFITS, BUSINESS INTERRUPTION, OR LOST INFORMATION) ARISING OUT OF THE USE OF OR INABILITY TO USE THE SOFTWARE, EVEN IF INTEL HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGES. SOME JURISDICTIONS PROHIBIT EXCLUSION OR LIMITATION OF LIABILITY FOR IMPLIED WARRANTIES OR CONSEQUENTIAL OR INCIDENTAL DAMAGES, SO THE ABOVE LIMITATION MAY NOT APPLY TO YOU. YOU MAY ALSO HAVE OTHER LEGAL RIGHTS THAT VARY FROM JURISDICTION TO JURISDICTION.

TERMINATION OF THIS AGREEMENT. Intel may terminate this Agreement at any time if you violate its terms. Upon termination, you will immediately destroy the Software or return all copies of the Software to Intel.

APPLICABLE LAWS. Claims arising under this Agreement shall be governed by the laws of California, excluding its principles of conflict of laws and the United Nations Convention on Contracts for the Sale of Goods. You may not export the Software in violation of applicable export laws and regulations. Intel is not obligated under any other agreements unless they are in writing and signed by an authorized representative of Intel.

GOVERNMENT RESTRICTED RIGHTS. The Software is provided with "RESTRICTED RIGHTS." Use, duplication, or disclosure by the Government is subject to restrictions as set forth in FAR52.227-14 and DFAR252.227-7013 *et seq.* or its successor. Use of the Software by the Government constitutes acknowledgment of Intel's proprietary rights therein. Contractor or Manufacturer is Intel.

Third-party Licenses

Portions of this release may include software distributed under the following licenses.

Open Toolkit Library (OpenTK)

The Open Toolkit library license

Copyright (c) 2006 - 2009 The Open Toolkit library.

Permission is hereby granted, free of charge, to any person obtaining a copy of this software and associated documentation files (the "Software"), to deal in the Software without restriction, including without limitation the rights to use, copy, modify, merge, publish, distribute, sublicense, and/or sell copies of the Software, and to permit persons to whom the Software is furnished to do so, subject to the following conditions

The above copyright notice and this permission notice shall be included in all copies or substantial portions of the Software.

THE SOFTWARE IS PROVIDED "AS IS", WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO THE WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND NONINFRINGEMENT. IN NO EVENT SHALL THE AUTHORS OR COPYRIGHT HOLDERS BE LIABLE FOR ANY CLAIM, DAMAGES OR OTHER LIABILITY, WHETHER IN AN ACTION OF CONTRACT, TORT OR OTHERWISE, ARISING FROM, OUT OF OR IN CONNECTION WITH THE SOFTWARE OR THE USE OR OTHER DEALINGS IN THE SOFTWARE.

Third parties

* The Open Toolkit library includes portions of the Mono class library, which are covered by the following license:

Copyright (c) 2004 Novell, Inc.

Permission is hereby granted, free of charge, to any person obtaining a copy of this software and associated documentation files (the "Software"), to deal in the Software without restriction, including without limitation the rights to use, copy, modify, merge, publish, distribute, sublicense, and/or sell copies of the Software, and to permit persons to whom the Software is furnished to do so, subject to the following conditions:

The above copyright notice and this permission notice shall be included in all copies or substantial portions of the Software.

THE SOFTWARE IS PROVIDED "AS IS", WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO THE WARRANTIES OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND NONINFRINGEMENT. IN NO EVENT SHALL THE AUTHORS OR COPYRIGHT HOLDERS BE LIABLE FOR ANY CLAIM, DAMAGES OR OTHER LIABILITY, WHETHER IN AN ACTION OF CONTRACT, TORT OR OTHERWISE, ARISING FROM, OUT OF OR IN CONNECTION WITH THE SOFTWARE OR THE USE OR OTHER DEALINGS IN THE SOFTWARE.

* Half-to-Single and Single-to-Half conversions are covered by the following license:

Copyright (c) 2002, Industrial Light & Magic, a division of Lucas Digital Ltd. LLC. All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

- Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
- Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.
- Neither the name of Industrial Light & Magic nor the names of its contributors may be used to endorse or promote products derived from this software without specific prior written permission.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS "AS IS" AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE COPYRIGHT OWNER OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

RSA Data Security-MD5 Message

RSA Data Security

Copyright (C) 1991-2, RSA Data Security, Inc. Created 1991. All rights reserved.

License to copy and use this software is granted provided that it is identified as the "RSA Data Security, Inc. MD5 Message-Digest Algorithm" in all material mentioning or referencing this software or this function.

License is also granted to make and use derivative works provided that such works are identified as "derived from the RSA Data Security, Inc. MD5 Message-Digest Algorithm" in all material mentioning or referencing the derived work.

RSA Data Security, Inc. makes no representations concerning either the merchantability of this software or the suitability of this software for any particular purpose. It is provided "as is" without express or implied warranty of any kind.

These notices must be retained in any copies of any part of this documentation and/or software.

Restrictions and Disclaimers

Information in this document is subject to change without notice.

Copyright © 2008-2021, Intel Corporation. All rights reserved.

Trademarks used in this text: *Dell EMC* and the *Dell EMC* logo are trademarks of Dell, Inc.; Intel is a trademark of Intel Corporation or its subsidiaries in the U.S. and/or other countries.

* Other trademarks and trade names may be used in this document to refer to either the entities claiming the marks and names or their products. Intel Corporation disclaims any proprietary interest in trademarks and trade names other than its own.

Restrictions and Disclaimers

The information contained in this document, including all instructions, cautions, and regulatory approvals and certifications, is provided by the supplier and has not been independently verified or tested by Dell. Dell EMC cannot be responsible for damage caused as a result of either following or failing to follow these instructions.

All statements or claims regarding the properties, capabilities, speeds or qualifications of the part referenced in this document are made by the supplier and not by Dell EMC. Dell EMC specifically disclaims knowledge of the accuracy, completeness or substantiation for any such statements. All questions or comments relating to such statements or claims should be directed to the supplier.

Export Regulations

Customer acknowledges that these Products, which may include technology and software, are subject to the customs and export control laws and regulations of the United States (U.S.) and may also be subject to the customs and export laws and regulations of the country in which the Products are manufactured and/or received. Customer agrees to abide by those laws and regulations. Further, under U.S. law, the Products may not be sold, leased or otherwise transferred to restricted end users or to restricted countries. In addition, the Products may not be sold, leased or otherwise transferred to, or utilized by an end-user engaged in activities related to weapons of mass destruction, including without limitation, activities related to the design, development, production or use of nuclear weapons, materials, or facilities, missiles or the support of missile projects, and chemical or biological weapons.

January 11, 2021