

Veritas InfoScale™ 7.4

Release Notes - Solaris

Last updated: 2019-09-25

Legal Notice

Copyright © 2019 Veritas Technologies LLC. All rights reserved.

Veritas and the Veritas Logo are trademarks or registered trademarks of Veritas Technologies LLC or its affiliates in the U.S. and other countries. Other names may be trademarks of their respective owners.

This product may contain third-party software for which Veritas is required to provide attribution to the third-party ("Third-Party Programs"). Some of the Third-Party Programs are available under open source or free software licenses. The License Agreement accompanying the Software does not alter any rights or obligations you may have under those open source or free software licenses. Refer to the third-party legal notices document accompanying this Veritas product or available at:

<https://www.veritas.com/about/legal/license-agreements>

The product described in this document is distributed under licenses restricting its use, copying, distribution, and decompilation/reverse engineering. No part of this document may be reproduced in any form by any means without prior written authorization of Veritas Technologies LLC and its licensors, if any.

THE DOCUMENTATION IS PROVIDED "AS IS" AND ALL EXPRESS OR IMPLIED CONDITIONS, REPRESENTATIONS AND WARRANTIES, INCLUDING ANY IMPLIED WARRANTY OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE OR NON-INFRINGEMENT, ARE DISCLAIMED, EXCEPT TO THE EXTENT THAT SUCH DISCLAIMERS ARE HELD TO BE LEGALLY INVALID. VERITAS TECHNOLOGIES LLC SHALL NOT BE LIABLE FOR INCIDENTAL OR CONSEQUENTIAL DAMAGES IN CONNECTION WITH THE FURNISHING, PERFORMANCE, OR USE OF THIS DOCUMENTATION. THE INFORMATION CONTAINED IN THIS DOCUMENTATION IS SUBJECT TO CHANGE WITHOUT NOTICE.

The Licensed Software and Documentation are deemed to be commercial computer software as defined in FAR 12.212 and subject to restricted rights as defined in FAR Section 52.227-19 "Commercial Computer Software - Restricted Rights" and DFARS 227.7202, et seq. "Commercial Computer Software and Commercial Computer Software Documentation," as applicable, and any successor regulations, whether delivered by Veritas as on premises or hosted services. Any use, modification, reproduction release, performance, display or disclosure of the Licensed Software and Documentation by the U.S. Government shall be solely in accordance with the terms of this Agreement.

Veritas Technologies LLC
500 E Middlefield Road
Mountain View, CA 94043

<http://www.veritas.com>

Technical Support

Technical Support maintains support centers globally. All support services will be delivered in accordance with your support agreement and the then-current enterprise technical support policies. For information about our support offerings and how to contact Technical Support, visit our website:

<https://www.veritas.com/support>

You can manage your Veritas account information at the following URL:

<https://my.veritas.com>

If you have questions regarding an existing support agreement, please email the support agreement administration team for your region as follows:

Worldwide (except Japan)

CustomerCare@veritas.com

Japan

CustomerCare_Japan@veritas.com

Documentation

Make sure that you have the current version of the documentation. Each document displays the date of the last update on page 2. The latest documentation is available on the Veritas website:

<https://sort.veritas.com/documents>

Documentation feedback

Your feedback is important to us. Suggest improvements or report errors or omissions to the documentation. Include the document title, document version, chapter title, and section title of the text on which you are reporting. Send feedback to:

infoscaledocs@veritas.com

You can also see documentation information or ask a question on the Veritas community site:

<http://www.veritas.com/community/>

Veritas Services and Operations Readiness Tools (SORT)

Veritas Services and Operations Readiness Tools (SORT) is a website that provides information and tools to automate and simplify certain time-consuming administrative tasks. Depending on the product, SORT helps you prepare for installations and upgrades, identify risks in your datacenters, and improve operational efficiency. To see what services and tools SORT provides for your product, see the data sheet:

https://sort.veritas.com/data/support/SORT_Data_Sheet.pdf

Contents

Chapter 1	Introduction	12
	About this document	12
Chapter 2	Requirements	14
	VCS system requirements	14
	Supported Solaris operating systems	14
	Supported Oracle VM Server for SPARC	15
	Storage Foundation for Databases features supported in database environments	16
	Storage Foundation memory requirements	17
	Supported database software	17
	Supported hardware and software	17
	Number of nodes supported	18
Chapter 3	Changes introduced in 7.4	19
	Changes related to installation and upgrades	19
	Change related to licensing	19
	Using Chef to deploy and configure InfoScale	20
	Change in upgrade path	20
	Changes related to InfoScale in cloud environments	21
	Changes related to Veritas File System	21
	Changes in VxFS Disk Layout Versions (DLV)	21
	Changes related to replication	21
	Veritas Volume Replicator performance improvements for UDP protocol	21
	Change in documentation	21
Chapter 4	Fixed issues	23
	Cluster Server fixed issues	23
Chapter 5	Limitations	24
	Storage Foundation software limitations	24
	Dynamic Multi-Pathing software limitations	24

Veritas Volume Manager software limitations	26
Veritas File System software limitations	27
SmartIO software limitations	28
Replication software limitations	29
VVR Replication in a shared environment	29
VVR IPv6 software limitations	29
VVR support for replicating across Storage Foundation versions	30
Cluster Server software limitations	30
Limitations related to bundled agents	30
Limitations related to VCS engine	33
Veritas cluster configuration wizard limitations	34
Limitations related to the VCS database agents	34
Systems in a cluster must have same system locale setting	35
Limitations with DiskGroupSnap agent [1919329]	35
Cluster Manager (Java console) limitations	36
Limitations related to LLT	36
Limitations related to I/O fencing	36
Limitations related to global clusters	38
CP Server 6.0.5 client fails to communicate with CP Server 7.0 with certificates having 2048-bit keys and SHA256 hashing [IIP-5803]	38
Clusters must run on VCS 6.0.5 and later to be able to communicate after upgrading to 2048 bit key and SHA256 signature certificates [3812313]	39
Storage Foundation Cluster File System High Availability software limitations	39
cfsmntadm command does not verify the mount options (2078634)	39
Stale SCSI-3 PR keys remain on disk after stopping the cluster and deporting the disk group	39
Unsupported FSS scenarios	40
Storage Foundation for Oracle RAC software limitations	40
Supportability constraints for normal or high redundancy ASM disk groups with CVM I/O shipping and FSS (3600155)	40
Limitations of CSSD agent	40
Oracle Clusterware/Grid Infrastructure installation fails if the cluster name exceeds 14 characters	41
Policy-managed databases not supported by CRSResource agent	41
Health checks may fail on clusters that have more than 10 nodes	41

Cached ODM not supported in Veritas InfoScale environments	41
Storage Foundation for Databases (SFDB) tools software limitations	41
Parallel execution of <code>vxsfdm</code> is not supported (2515442)	41
Creating point-in-time copies during database structural changes is not supported (2496178)	42
Oracle Data Guard in an Oracle RAC environment	42

Chapter 6	Known issues	43
	Issues related to installation and upgrade	43
	Switch fencing in enable or disable mode may not take effect if VCS is not reconfigured [3798127]	44
	After the upgrade to version 7.4, the installer may fail to stop the Asynchronous Monitoring Framework (AMF) process [3781993]	44
	LLT may fail to start after upgrade on Solaris 11 [3770835]	44
	On SunOS, drivers may not be loaded after a reboot [3798849]	44
	On Oracle Solaris, drivers may not be loaded after stop and then reboot [3763550]	45
	During an upgrade process, the <code>AMF_START</code> or <code>AMF_STOP</code> variable values may be inconsistent [3763790]	45
	Uninstallation fails on global zone on Solaris 11 if product packages are installed on both global zone and local zone [3762814]	45
	On Solaris 11, when you install the operating system together with SFHA products using Automated Installer, the local installer scripts do not get generated (3640805)	45
	Stopping the installer during an upgrade and then resuming the upgrade might freeze the service groups (2574731)	46
	Installing VRTSvlic package during live upgrade on Solaris system non-global zones displays error messages [3623525]	46
	VCS installation with CPI fails when a non-global zone is in installed state and zone root is not mounted on the node (2731178)	46
	Log messages are displayed when VRTSvcs is uninstalled on Solaris 11 [2919986]	47
	Cluster goes into <code>STALE_ADMIN_WAIT</code> state during upgrade from VCS 5.1 to 6.1 or later [2850921]	47
	Flash Archive installation not supported if the target system's root disk is encapsulated	48

The Configure Sybase ASE CE Instance in VCS option creates duplicate service groups for Sybase binary mount points (2560188)	48
The Installer fails to unload GAB module while installation of SF packages [3560458]	48
On Solaris 11 non-default ODM mount options will not be preserved across package upgrade (2745100)	48
Upgrade fails because there is zone installed on the VxFS file system which is offline. The packages in the zone are not updated. (3319753)	49
If you choose to upgrade nodes without zones first, the rolling upgrade or phased upgrade is not blocked in the beginning, but fails later (3319961)	49
Upgrades from previous SF Oracle RAC versions may fail on Solaris systems (3256400)	49
After a locale change restart the vxconfig daemon (2417547, 2116264)	50
Verification of Oracle binaries incorrectly reports as failed during Oracle Grid Infrastructure installation	50
Live upgrade of the InfoScale product may detect the wrong product or ask for the license repeatedly (3870685)	50
In RHEV environment, if you stop the SF service and then start it by installer, the permission on dmpnode will get lost (3870111)	51
The installer fails to upgrade the product packages on Solaris 11 during an upgrade to InfoScale 7.4 (3896530)	51
Resource faults during Rolling upgrade due to perl changes (3930605)	52
Storage Foundation known issues	52
Dynamic Multi-Pathing known issues	52
Veritas Volume Manager known issues	53
Veritas File System known issues	69
Replication known issues	74
The secondary vradmind may appear hung and the vradmind commands may fail (3940842,3944301)	74
Data corruption may occur if you perform a rolling upgrade of InfoScale Storage or InfoScale Enterprise from 7.3.1 or earlier to 7.4 or later during replication (3951527)	75
After the product upgrade on secondary site, replication may fail to resume with "Secondary SRL missing" error [3931763]	75
vradmind repstatus command reports secondary host as "unreachable"(3896588)	76

RVGPrimary agent operation to start replication between the original Primary and the bunker fails during failback (2036605)	76
A snapshot volume created on the Secondary, containing a VxFS file system may not mount in read-write mode and performing a read-write mount of the VxFS file systems on the new Primary after a global clustering site failover may fail [3761497]	77
In an IPv6-only environment RVG, data volumes or SRL names cannot contain a colon (1672410, 1672417)	78
vxassist relayout removes the DCM (145413)	78
vradmin functionality may not work after a master switch operation [2158679]	78
Cannot relayout data volumes in an RVG from concat to striped-mirror (2129601)	79
vradmin verifydata operation fails when replicating between versions 5.1 and 6.0 or later (2360713)	80
vradmin verifydata may report differences in a cross-endian environment (2834424)	80
vradmin verifydata operation fails if the RVG contains a volume set (2808902)	80
Bunker replay does not occur with volume sets (3329970)	81
SmartIO does not support write-back caching mode for volumes configured for replication by Volume Replicator (3313920)	81
During moderate to heavy I/O, the vradmin verifydata command may falsely report differences in data (3270067)	81
The vradmin repstatus command does not show that the SmartSync feature is running [3343141]	81
While vradmin commands are running, vradmin may temporarily lose heartbeats (3347656, 3724338)	82
Write I/Os on the primary logowner may take a long time to complete (2622536)	82
DCM logs on a disassociated layered data volume results in configuration changes or CVM node reconfiguration issues (3582509)	83
After performing a CVM master switch on the secondary node, both rlinks detach (3642855)	83
The RVGPrimary agent may fail to bring the application service group online on the new Primary site because of a previous primary-elect operation not being run or not completing successfully (3761555, 2043831)	83

A snapshot volume created on the Secondary, containing a VxFS file system may not mount in read-write mode and performing a read-write mount of the VxFS file systems on the new Primary after a global clustering site failover may fail (1558257)	84
DCM plex becomes inaccessible and goes into DISABLED(SPARSE) state in case of node failure. (3931775)	85
Cluster Server known issues	85
Operational issues for VCS	85
Issues related to the VCS engine	88
Issues related to the bundled agents	94
Issues related to the VCS database agents	106
Issues related to the agent framework	110
Issues related to Intelligent Monitoring Framework (IMF)	113
Issues related to global clusters	117
Issues related to the Cluster Manager (Java Console)	117
VCS Cluster Configuration wizard issues	118
LLT known issues	118
I/O fencing known issues	119
GAB known issues	126
Storage Foundation and High Availability known issues	127
Cache area is lost after a disk failure (3158482)	127
NFS issues with VxFS Storage Checkpoints (2027492)	127
Some SmartTier for Oracle commands do not work correctly in non-POSIX locales (2138030)	128
In an IPv6 environment, db2icrt and db2idrop commands return a segmentation fault error during instance creation and instance removal (1602444)	128
Oracle 11gR1 may not work on pure IPv6 environment (1819585)	129
Not all the objects are visible in the VOM GUI (1821803)	129
An error message is received when you perform off-host clone for RAC and the off-host node is not part of the CVM cluster (1834860)	130
A volume's placement class tags are not visible in the Veritas Enterprise Administrator GUI when creating a dynamic storage tiering placement policy (1880081)	130
Storage Foundation Cluster File System High Availability known issues	131
In an FSS environment, creation of mirrored volumes may fail for SSD media [3932494]	131

Mount command may fail to mount the file system (3913246)	131
After the local node restarts or panics, the FSS service group cannot be online successfully on the local node and the remote node when the local node is up again (3865289)	132
In the FSS environment, if DG goes to the dgdisable state and deep volume monitoring is disabled, successive node joins fail with error 'Slave failed to create remote disk: retry to add a node failed' (3874730)	133
DG creation fails with error "V-5-1-585 Disk group punedatadg: cannot create: SCSI-3 PR operation failed" on the VSCSI disks (3875044)	133
Write back cache is not supported on the cluster in FSS scenario [3723701]	134
CVMVOLDg agent is not going into the FAULTED state. [3771283]	134
CFS commands might hang when run by non-root (3038283)	134
The fsappadm subfilemove command moves all extents of a file (3258678)	134
Certain I/O errors during clone deletion may lead to system panic. (3331273)	135
Panic due to null pointer de-reference in vx_bmap_lookup() (3038285)	135
In a CFS cluster, that has multi-volume file system of a small size, the fsadm operation may hang (3348520)	135
Storage Foundation for Oracle RAC known issues	136
Oracle RAC known issues	136
Storage Foundation Oracle RAC issues	136
Storage Foundation for Databases (SFDB) tools known issues	143
Clone operations fail for instant mode snapshot (3916053)	143
Sometimes SFDB may report the following error message: SFDB remote or privileged command error (2869262)	144
SFDB commands do not work in IPV6 environment (2619958)	144
When you attempt to move all the extents of a table, the dbdst_obj_move(1M) command fails with an error (3260289)	144
Attempt to use SmartTier commands fails (2332973)	145
Attempt to use certain names for tiers results in error (2581390)	145

Clone operation failure might leave clone database in unexpected state (2512664)	145
Clone command fails if PFILE entries have their values spread across multiple lines (2844247)	146
Data population fails after datafile corruption, rollback, and restore of offline checkpoint (2869259)	146
Flashsnap clone fails under some unusual archivelog configuration on RAC (2846399)	146
vxdbd process is online after Flash archive installation (2869269)	147
On Solaris 11.1 SPARC, setting up the user-authentication process using the sfae_auth_op command fails with an error message (3556996)	147
In the cloned database, the seed PDB remains in the mounted state (3599920)	147
Cloning of a container database may fail after a reverse resync commit operation is performed (3509778)	148
If one of the PDBs is in the read-write restricted state, then cloning of a CDB fails (3516634)	148
Cloning of a CDB fails for point-in-time copies when one of the PDBs is in the read-only mode (3513432)	148
If a CDB has a tablespace in the read-only mode, then the cloning fails (3512370)	149
If any SFDB installation with authentication setup is upgraded to 7.4, the commands fail with an error (3644030)	149
Error message displayed when you use the vxsfadm -a oracle -s filesnap -o destroyclone command (3901533)	150

Introduction

This chapter includes the following topics:

- [About this document](#)

About this document

This document provides information that is specific to version 7.4 of the Veritas InfoScale products.

Review this entire document before using the following products:

- Veritas InfoScale Foundation
- Veritas InfoScale Storage
- Veritas InfoScale Availability
- Veritas InfoScale Enterprise

The information in this document supersedes the information provided in the product-specific documents.

You can download the latest version of this document from the Veritas Service and Operations Readiness Tools (SORT) website at:

<https://sort.veritas.com/documents>

The following documents provide further information that is common to all the InfoScale for Solaris products:

- *Veritas InfoScale Getting Started Guide*
- *Veritas InfoScale Installation Guide*

For information about the InfoScale product components and their capabilities, refer to the corresponding configuration and upgrade guides and administrator's guides.

For information about installing and configuring and your databases with the InfoScale products, refer to the database-specific installation and configuration guides.

For the latest information on updates, patches, and known issues regarding this release, see the Late Breaking News (LBN) at:

https://www.veritas.com/support/en_US/article.100042983.html

Requirements

This chapter includes the following topics:

- [VCS system requirements](#)
- [Supported Solaris operating systems](#)
- [Supported Oracle VM Server for SPARC](#)
- [Storage Foundation for Databases features supported in database environments](#)
- [Storage Foundation memory requirements](#)
- [Supported database software](#)
- [Supported hardware and software](#)
- [Number of nodes supported](#)

VCS system requirements

This section describes system requirements for VCS.

The following information does not apply to SF Oracle RAC installations.

VCS requires that all nodes in the cluster use the same processor architecture and all nodes in the cluster must run the same VCS version. However, the nodes can have different versions of the supported operating system.

Supported Solaris operating systems

For current updates, visit the Veritas Services and Operations Readiness Tools Installation and Upgrade page: https://sort.veritas.com/land/install_and_upgrade.

[Table 2-1](#) shows the supported operating systems for this release.

Table 2-1 Supported operating systems

Operating systems	Levels	Chipsets	Supported Products
Solaris 11	Solaris 11.1 and up to Support Repository Updates (SRU) 11.1.21.4.1 Solaris 11.2 and up to Support Repository Updates (SRU) 11.2.15.5.1 Solaris 11.3 and up to Support Repository Updates (SRU) 11.3.20.5.0	SPARC	Veritas InfoScale Foundation Veritas InfoScale Storage Veritas InfoScale Availability Veritas InfoScale Enterprise
Solaris 11	Solaris 11.1 and up to Support Repository Updates (SRU) 11.1.21.4.1 Solaris 11.2 and up to Support Repository Updates (SRU) 11.2.15.5.1 Solaris 11.3 and up to Support Repository Updates (SRU) 11.3.20.5.0	x64	Veritas InfoScale Availability

This release (version 7.4) supports native zones on the Solaris 11 operating system. This release does not support the Kernel Zones feature on Solaris 11 x64 operating system. However, it supports the Kernel Zones feature on Solaris 11 SPARC operating system.

For the SF Oracle RAC component, all nodes in the cluster need to have the same operating system version and update level.

Supported Oracle VM Server for SPARC

Supported Oracle VM Server for SPARC (OVM) versions are OVM 2.2, OVM 3.0, OVM 3.1, OVM 3.1.1, OVM 3.2, OVM 3.3.4, and OVM 3.4.0.3.

For supported OS version for Oracle VM Server for SPARC, refer to *Oracle VM server for SPARC Release Notes*.

If necessary, upgrade Solaris before you install the Veritas InfoScale products.

See <http://www.veritas.com/docs/000087666> before you upgrade to Oracle Solaris 11.1.

Storage Foundation for Databases features supported in database environments

Storage Foundation for Databases (SFDB) product features are supported for the following database environments:

Table 2-2 SFDB features supported in database environments

Storage Foundation feature	DB2	Oracle	Oracle RAC	Sybase	Sybase ASE CE
Oracle Disk Manager	No	Yes	Yes	No	No
Cached Oracle Disk Manager	No	Yes	No	No	No
Quick I/O	Yes	Yes	Yes	Yes	Yes
Cached Quick I/O	Yes	Yes	Yes	Yes	Yes
Concurrent I/O	Yes	Yes	Yes	Yes	Yes
Storage Checkpoints	Yes	Yes	Yes	Yes	Yes
Flashsnap	Yes	Yes	Yes	Yes	Yes
SmartTier	Yes	Yes	Yes	Yes	Yes
Database Storage Checkpoints Note: Requires Enterprise license	No	Yes	Yes	No	No
Database Flashsnap Note: Requires Enterprise license	No	Yes	Yes	No	No
SmartTier for Oracle Note: Requires Enterprise license	No	Yes	Yes	No	No

Notes:

- SmartTier is an expanded and renamed version of Dynamic Storage Tiering (DST).

- Storage Foundation for Databases (SFDB) tools Database Storage Checkpoint, Database Flashsnap, and SmartTier for Oracle are supported with an Enterprise product license.

For the most current information on Storage Foundation products and single instance Oracle versions supported, see:

For 6.2 and earlier versions: <http://www.veritas.com/docs/000002658>

For 7.0 and later versions:: <http://www.veritas.com/docs/000115952>

Review the current Oracle documentation to confirm the compatibility of your hardware and software.

Storage Foundation memory requirements

Veritas recommends 2 GB of memory over the minimum requirement for the operating system.

Supported database software

For the latest information on supported database, see the following TechNote:<http://www.veritas.com/docs/000002658>

Additionally, see the following Oracle support site for information on patches that may be required by Oracle for each release.<https://support.oracle.com>

Supported hardware and software

For the latest information on the supported hardware and software, see the appropriate compatibility list at:

[https://www.veritas.com/content/support/en_US/DocumentBrowsing.html?product=InfoScale & Storage Foundation](https://www.veritas.com/content/support/en_US/DocumentBrowsing.html?product=InfoScale%20Storage%20Foundation)

Select 7.4 in the **Version** filter and select **Compatibility Lists** under **Document Type**.

Before installing or upgrading the InfoScale products, review the current compatibility list to confirm the compatibility of your hardware and software. For information on specific setup requirements, see the corresponding Configuration and Upgrade Guide.

In addition to the compatibility list, a Late Breaking News (LBN) is available for the latest updates, patches, and software issues regarding this release:

https://www.veritas.com/support/en_US/article.100042983.html

Number of nodes supported

Veritas InfoScale supports cluster configurations up to 64 nodes.

SFHA, SFCFSHA, SF Oracle RAC: Flexible Storage Sharing (FSS) only supports cluster configurations with up to 8 nodes.

SFHA, SFCFSHA: SmartIO writeback caching only supports cluster configurations with up to 2 nodes.

Changes introduced in 7.4

This chapter includes the following topics:

- [Changes related to installation and upgrades](#)
- [Changes related to InfoScale in cloud environments](#)
- [Changes related to Veritas File System](#)
- [Changes related to replication](#)
- [Change in documentation](#)

Changes related to installation and upgrades

The following changes are introduced to the installation and upgrade of Veritas InfoScale 7.4.

Change related to licensing

Veritas InfoScale introduces a file-based license key format that will be used instead of the text-based license key format used in versions before 7.4. Licenses are now distributed using license key files with an **slf** extension.

Slf license key files are required while installing or upgrading to version 7.4. Ensure that the license key file is downloaded on the local host, where you want to install or upgrade the product.

Note: The text-based license keys that are used in previous product versions are not supported when upgrading to version 7.4. If you plan to upgrade any of the InfoScale products to version 7.4, first contact Customer Care for your region to procure an applicable slf license key file. Refer to the following link for contact information of the Customer Care center for your region:

https://www.veritas.com/content/support/en_US/contact-us.html

For more information about managing licenses, see the *Veritas InfoScale Installation Guide*.

Using Chef to deploy and configure InfoScale

You can now deploy and configure the InfoScale product suite in your environment using Chef. Chef's powerful automation platform transforms infrastructure into code and allows you to automate various deployment and configuration operations performed in InfoScale.

You can download the `veritas_infoscale` cookbook from the Chef Supermarket website at:

<https://supermarket.chef.io/>

Use the Chef resources to perform the following operations in a InfoScale environment:

- Installation
- Configuration
- Full upgrade
- Rolling upgrade
- Patch upgrade
- License
- Security
- Fencing

Note: Note that only disk-based, majority-based, and disabled-based fencing can be configured using Chef.

Prerequisites and steps for using the `veritas_infoscale` cookbook are provided in the `README.md` file, which is located in the `veritas_infoscale.tar.gz` cookbook package.

Change in upgrade path

You can upgrade to Veritas InfoScale 7.4 only if the base version of your currently installed product is 6.1.1 or later.

Changes related to InfoScale in cloud environments

The following sections describe the changes that are introduced to support InfoScale deployments in cloud environments.

Changes related to Veritas File System

The following changes are introduced to Veritas File System (VxFS) of Veritas InfoScale 7.4.

Changes in VxFS Disk Layout Versions (DLV)

The following DLV changes are now applicable:

- Added support for DLV 14
- Default DLV is DLV 14
- Support deprecated for DLV 9

With this change, you can create and mount VxFS only on DLV 10 and later. DLV 6 to 9 can be used for local mount only.

Changes related to replication

The following changes are introduced to replication of Veritas InfoScale 7.4.

Veritas Volume Replicator performance improvements for UDP protocol

Veritas Volume Replicator uses UDP or TCP transport protocols to communicate between the Primary and Secondary. These protocols are used to exchange data messages. In this current 7.4 release UDP connection is also tuned to achieve better performance. UDP also uses 16 multiple data structures, where all the data is transmitted in parallel to multiple network sockets and distributes the network load.

Change in documentation

From version 7.4 onwards, the InfoScale documentation set will not include a separate *What's New* document. All the changes that are introduced in the latest

release will be listed in the *Veritas InfoScale Release Notes* document for each supported platform.

Fixed issues

This chapter includes the following topics:

- [Cluster Server fixed issues](#)

Cluster Server fixed issues

This section lists the Cluster Server (VCS) component issues that are fixed in this release.

Table 4-1 Veritas Cluster Server fixed issues

Incident	Description
3931526	After a node reboots, it fails to join its cluster if the fencing module fails to start.
3932288	On cluster nodes that run RHEL 7.4, the NFS agent fails to come online or go offline when NFSSecurity is enabled.

Limitations

This chapter includes the following topics:

- [Storage Foundation software limitations](#)
- [Replication software limitations](#)
- [Cluster Server software limitations](#)
- [Storage Foundation Cluster File System High Availability software limitations](#)
- [Storage Foundation for Oracle RAC software limitations](#)
- [Storage Foundation for Databases \(SFDB\) tools software limitations](#)

Storage Foundation software limitations

These software limitations apply to the following products:

- Veritas InfoScale Foundation
- Veritas InfoScale Storage
- Veritas InfoScale Enterprise

Dynamic Multi-Pathing software limitations

These software limitations apply to the following products:

- Veritas InfoScale Foundation
- Veritas InfoScale Storage
- Veritas InfoScale Enterprise

DMP support for the Solaris format command (2043956)

When DMP is enabled to support Solaris ZFS pools, the Solaris `format` command displays either a path or the corresponding `dmnode`. The result depends on the order in which the `format` command parses the entries in the `/dev/rdisk` directory.

DMP settings for NetApp storage attached environment

To minimize the path restoration window and maximize high availability in the NetApp storage attached environment,change the default values for the DMP tunable parameters.

Table 5-1 describes the DMP tunable parameters and the new values.

Table 5-1 DMP settings for NetApp storage attached environment

Parameter name	Definition	New value	Default value
dmp_restore_interval	DMP restore daemon cycle	60 seconds.	300 seconds.
dmp_path_age	DMP path aging tunable	120 seconds.	300 seconds.

The change is persistent across reboots.

To change the tunable parameters

- 1 Issue the following commands:
- ```
vxddmpadm settune dmp_restore_interval=60

vxddmpadm settune dmp_path_age=120
```
- 2 To verify the new settings, use the following commands:
- ```
# vxddmpadm gettune dmp_restore_interval

# vxddmpadm gettune dmp_path_age
```

ZFS pool in unusable state if last path is excluded from DMP (1976620)

When a DMP device is used by a ZFS pool, do not exclude the last path to the device. This can put the ZFS pool in an unusable state.

When an I/O domain fails, the vxdisk scandisks or vxctl enable command take a long time to complete (2791127)

When an I/O domain fails, the vxdisk scandisks or vxctl enable from the Oracle VM Server for SPARC guest take a long time to complete. `vdc_ioctl`s like `DKIOCGGEOM` and `DKIOCINFO` also take more time to return. These issues seem to be due to retry operations performed at the Solaris operating system layer.

Reducing the `vdc_timeout` value to lower value might help to bring down time. Dynamic multi-pathing (DMP) code is optimized to avoid making such `vdc_ioctl` calls in an Oracle VM Server for SPARC guest environment as much possible. This change considerably reduces delays.

A complete resolution to this issue may require changes at the Solaris operating system level.

Veritas Volume Manager software limitations

The following are software limitations in this release of Veritas Volume Manager.

Snapshot configuration with volumes in shared disk groups and private disk groups is not supported (2801037)

A snapshot configuration with volumes in the shared disk groups and private disk groups is not a recommended configuration. In this release, this configuration is not supported.

SmartSync is not supported for Oracle databases running on raw VxVM volumes

SmartSync is not supported for Oracle databases that are configured on raw volumes, because Oracle does not support the raw volume interface.

Veritas InfoScale does not support thin reclamation of space on a linked mirror volume (2729563)

The thin reclamation feature does not support thin reclamation for a linked mirror volume.

A 1 TB disk that is not labeled using operating system commands goes into an error state after the vxconfigd daemon is restarted

A 1 TB disk that is not labeled using operating system commands goes into an error state after the vxconfigd daemon is restarted. `vxconfigd` daemon is restarted.

Currently, a solution from the vendor is not available.

Converting a multi-pathed disk

When converting a multi-pathed disk that is smaller than 1 TB from a VTOC label to an EFI label, you must issue the `format -e` command for each path. For example, if a node has two paths, `c1t2d0s2` and `c2t2sd0s2`, you must run the `format -e` command on each of the two paths.

Thin reclamation requests are not redirected even when the ioship policy is enabled (2755982)

Reclamation requests fail from nodes that do not have local connectivity to the disks, even when the ioship policy is enabled. Reclamation I/Os are not redirected to another node.

Veritas Operations Manager does not support disk, disk group, and volume state information related to CVM I/O shipping feature (2781126)

The Veritas Operations Manager (VOM) does not support disk, disk group, and volume state information related to the I/O shipping feature introduced in this release of Cluster Volume Manager. New states such as `lfailed`, `lmissing` or `LDISABLED` are introduced when I/O shipping is active because of storage disconnectivity.

Veritas File System software limitations

The following are software limitations in this release of Veritas File System.

Recommended limit of number of files in a directory

To maximize VxFS performance, do not exceed 100,000 files in the same directory. Use multiple directories instead.

The vxlist command cannot correctly display numbers greater than or equal to 1 EB

The `vxlist` command and all of the other commands that use the same library as the `vxlist` command cannot correctly display numbers greater than or equal to 1 EB.

Limitations with delayed allocation for extending writes feature

The following limitations apply to the delayed allocation for extending writes feature:

- In the cases where the file data must be written to disk immediately, delayed allocation is disabled on that file. Examples of such cases include Direct I/O, concurrent I/O, FDD/ODM access, and synchronous I/O.
- Delayed allocation is not supported on memory mapped files.
- Delayed allocation is not supported with BSD quotas. When BSD quotas are enabled on a file system, delayed allocation is turned off automatically for that file system.
- Delayed allocation is not supported for shared mounts in a cluster file system.

FlashBackup feature of NetBackup 7.5 (or earlier) does not support disk layout Version 8, 9, or 10

The FlashBackup feature of NetBackup 7.5 (or earlier) does not support disk layout Version 8, 9, or 10.

Compressed files that are backed up using NetBackup 7.1 or prior become uncompressed when you restore the files

The NetBackup 7.1 release and prior does not support the file compression feature. If you back up compressed files using NetBackup 7.1 or a prior release, the files become uncompressed when you restore the files.

SmartIO software limitations

The following are the SmartIO software limitations in this release.

Cache is not online after a reboot

Generally, the SmartIO cache is automatically brought online after a reboot of the system.

If the SSD driver module is not loaded automatically after the reboot, you need to load the driver and bring the cache disk group online manually.

To bring a cache online after a reboot

- 1 Perform a scan of the OS devices:

```
# vxdisk scandisks
```

- 2 Bring the cache online manually:

```
# vxdg import cachedg
```

The `sfcache` operations may display error messages in the caching log when the operation completed successfully (3611158)

The `sfcache` command calls other commands to perform the caching operations. If a command fails, additional commands may be called to complete the operation. For debugging purposes, the caching log includes all of the success messages and failure messages for the commands that are called.

If the `sfcache` command has completed successfully, you can safely ignore the error messages in the log file.

Replication software limitations

These software limitations apply to the following products:

- Veritas InfoScale Storage
- Veritas InfoScale Enterprise

VVR Replication in a shared environment

Currently, replication support is limited to 8-node cluster applications.

VVR IPv6 software limitations

VVR does not support the following Internet Protocol configurations:

- A replication configuration from an IPv4-only node to an IPv6-only node and from an IPv6-only node to an IPv4-only node is not supported, because the IPv6-only node has no IPv4 address configured on it and therefore VVR cannot establish communication between the two nodes.
- A replication configuration in which an IPv4 address is specified for the `local_host` attribute of a primary RLINK and an IPv6 address is specified for the `remote_host` attribute of the same RLINK.

- A replication configuration in which an IPv6 address is specified for the `local_host` attribute of a primary RLINK and an IPv4 address is specified for the `remote_host` attribute of the same RLINK.
- IPv6 is not supported in a CVM and VVR cluster where some nodes in the cluster are IPv4-only and other nodes in the same cluster are IPv6-only, or all nodes of a cluster are IPv4-only and all nodes of a remote cluster are IPv6-only.
- VVR does not support Edge and NAT-PT routers that facilitate IPv4 and IPv6 address translation.

VVR support for replicating across Storage Foundation versions

VVR supports replication between Storage Foundation 6.1 and the prior major releases of Storage Foundation (6.0 and 6.0.1). Replication between versions is supported for disk group versions 170, 180, and 190 only. Both the Primary and Secondary hosts must be using a supported disk group version.

Cluster Server software limitations

These software limitations apply to the following products:

- Veritas InfoScale Availability
- Veritas InfoScale Enterprise

Limitations related to bundled agents

Programs using networked services may stop responding if the host is disconnected

Programs using networked services (for example, NIS, NFS, RPC, or a TCP socket connection to a remote host) can stop responding if the host is disconnected from the network. If such a program is used as an agent entry point, a network disconnect can cause the entry point to stop responding and possibly time out.

For example, if the host is configured to use NIS maps as a client, basic commands such as `ps -ef` can hang if there is network disconnect.

Veritas recommends creating users locally. To reflect local users, configure:

```
/etc/nsswitch.conf
```

Volume agent clean may forcibly stop volume resources

When the attribute `FaultOnMonitorTimeouts` calls the Volume agent clean entry point after a monitor time-out, the `vxvol -f stop` command is also issued. This command forcibly stops all volumes, even if they are still mounted.

False concurrency violation when using PidFiles to monitor application resources

The PID files created by an application contain the PIDs for the processes that are monitored by Application agent. These files may continue to exist even after a node running the application crashes. On restarting the node, the operating system may assign the PIDs listed in the PID files to other processes running on the node.

Thus, if the Application agent monitors the resource using the `PidFiles` attribute only, the agent may discover the processes running and report a false concurrency violation. This could result in some processes being stopped that are not under VCS control.

Volumes in a disk group start automatically irrespective of the value of the StartVolumes attribute in VCS [2162929]

Volumes in a disk group are started automatically when the disk group is imported, irrespective of the value of the `StartVolumes` attribute in VCS. This behavior is observed if the value of the system-level attribute `autostartvolumes` in Veritas Volume Manager is set to `On`.

Workaround: If you do not want the volumes in a disk group to start automatically after the import of a disk group, set the `autostartvolumes` attribute to `Off` at the system level.

Online for LDom resource fails [2517350]

Online of LDom resource fails when the boot disk configured in the guest domain that is a part of the virtual disk multi-pathing group (`mpgroup`) and also the primary path to the virtual disk is not available.

This is due to the limitations in Oracle VM Server that do not allow retrying of other device paths that exist for the virtual disks, which are part of a virtual disk multi-pathing group, when booting a guest domain.

Workaround: None.

Zone agent registered to IMF for Directory Online event

The Directory Online event monitors the Zone root directory. If the parent directory of the Zone root directory is deleted or moved to another location, AMF does not

provide notification to the Zone agent. In the next cycle of the zone monitor, it detects the change and reports the state of the resource as offline.

LDom resource calls clean entry point when primary domain is gracefully shut down

LDom agent sets failure policy of the guest domain to stop when primary domain stops. Thus when primary domain is shut down, guest domain is stopped. Moreover, when primary domain is shutdown, ldmd daemon is stopped abruptly and LDom configuration cannot be read. These operations are not under VCS control and VCS may call clean entry point.

Workaround: No workaround.

Application agent limitations

- ProPCV fails to prevent execution of script-based processes configured under MonitorProcesses.

Interface object name must match net<x>/v4static for VCS network reconfiguration script in Solaris 11 guest domain [2840193]

If the Solaris 11 guest domain is configured for DR and its interface object name does not match the `net<x>/v4static` pattern then the VCS guest network reconfiguration script (VRTSvcsnr) running inside the guest domain adds a new interface object and the existing entry will remain as is.

Share agent limitation (2717636)

If the Share resource is configured with VCS to share a system directory (for example, /usr) or Oracle Solaris 11 which gets mounted at boot time, the VCS share resource detects it online once VCS starts on the node after a panic or halt. This can lead to a concurrency violation if the share resource is a part of a failover service group, and the group has failed over to another node in the cluster. VCS brings down the Share resource subsequently. This is due to the share command behavior or Oracle Solaris 11, where a directory shared with share command remains persistently on the system across reboots.

Campus cluster fire drill does not work when DSM sites are used to mark site boundaries [3073907]

The campus cluster FireDrill agent currently uses the SystemZones attribute to identify site boundaries. Hence, campus cluster FireDrill is not supported in DSM enabled environment.

Workaround: Disable DSM and configure the SystemZones attribute on the application service group to perform the fire drill.

Mount agent reports resource state as OFFLINE if the configured mount point does not exist [3435266]

If a configured mount point does not exist on a node, then the Mount agent reports the resource state as OFFLINE instead of UNKNOWN on that particular node. If an attempt is made for onlining the resource, it fails on that node as the mount point does not exist.

Workaround: Make sure that configured mount point exists on all nodes of the cluster or alternatively set the CreateMntPt attribute value of Mount agent to 1. This will ensure that if a mount point does not exist then it will create while onlining the resource.

Limitations related to VCS engine

Loads fail to consolidate and optimize when multiple groups fault [3074299]

When multiple groups fault and fail over at the same time, the loads are not consolidated and optimized to choose the target systems.

Workaround: No workaround.

Preferred fencing ignores the forecasted available capacity [3077242]

Preferred fencing in VCS does not consider the forecasted available capacity for fencing decision. The fencing decision is based on the system weight configured.

Workaround: No workaround.

Failover occurs within the SystemZone or site when BiggestAvailable policy is set [3083757]

Failover always occurs within the SytemZone or site when the BiggestAvailable failover policy is configured. The target system for failover is always selected based on the biggest available system within the SystemZone.

Workaround: No workaround.

Load for Priority groups is ignored in groups with BiggestAvailable and Priority in the same group[3074314]

When there are groups with both BiggestAvailable and Priority as the failover policy in the same cluster, the load for Priority groups are not considered.

Workaround: No workaround.

Veritas cluster configuration wizard limitations

Environment variable used to change log directory cannot redefine the log path of the wizard [3609791]

By default, the Veritas cluster configuration wizard writes the logs in `/var/VRTSvcS/log` directory. VCS provides a way to change the log directory through environment variable `VCS_LOG`, but this does not apply to the logs of VCS wizards.

Workaround: No workaround.

Cluster configuration wizard takes long time to configure a cluster on Solaris systems [3582495]

Some times the VCS cluster configuration wizard takes a long time (10 to 15 minutes) to configure a VCS cluster on Solaris systems. The wizard may appear stuck but it completes the configuration in some time.

Workaround: No workaround.

Limitations related to the VCS database agents

DB2 RestartLimit value [1234959]

When multiple DB2 resources all start at the same time with no dependencies, they tend to interfere or race with each other. This is a known DB2 issue.

The default value for the DB2 agent RestartLimit is 3. This higher value spreads out the re-start of the DB2 resources (after a resource online failure), which lowers the chances of DB2 resources all starting simultaneously.

Sybase agent does not perform qrmutil based checks if Quorum_dev is not set (2724848)

If you do not set the Quorum_dev attribute for Sybase Cluster Edition, the Sybase agent does not perform the qrmutil-based checks. This error in configuration may lead to undesirable results. For example, if qrmutil returns failure pending, the agent does not panic the system. Thus, the Sybase agent does not perform qrmutil-based checks because the Quorum_dev attribute is not set.

Therefore, setting Quorum_Dev attribute is mandatory for Sybase cluster edition.

Pluggable database (PDB) online may timeout when started after container database (CDB) [3549506]

PDB may take long time to start when it is started for the first time after starting CDB. As a result, the PDB online initiated using VCS may cause ONLINE timeout and the PDB online process may get cancelled.

Workaround: Increase the OnlineTimeout attribute value of the Oracle type resource.

Systems in a cluster must have same system locale setting

VCS does not support clustering of systems with different system locales. All systems in a cluster must be set to the same locale.

Limitations with DiskGroupSnap agent [1919329]

The DiskGroupSnap agent has the following limitations:

- The DiskGroupSnap agent does not support layered volumes.
- If you use the Bronze configuration for the DiskGroupSnap resource, you could end up with inconsistent data at the secondary site in the following cases:
 - After the fire drill service group is brought online, a disaster occurs at the primary site during the fire drill.
 - After the fire drill service group is taken offline, a disaster occurs at the primary while the disks at the secondary are resynchronizing.

Veritas recommends that you use the Gold configuration for the DiskGroupSnap resource.

Cluster Manager (Java console) limitations

This section covers the software limitations for Cluster Manager (Java Console).

Cluster Manager does not work if the hosts file contains IPv6 entries

VCS Cluster Manager fails to connect to the VCS engine if the `/etc/hosts` file contains IPv6 entries.

Workaround: Remove IPv6 entries from the `/etc/hosts` file.

VCS Simulator does not support I/O fencing

When running the Simulator, be sure the `UseFence` attribute is set to the default, "None".

Limitations related to LLT

This section covers LLT-related software limitations.

Limitation of LLT support over UDP using alias IP [3622175]

When configuring the VCS cluster, if alias IP addresses are configured on the LLT links as the IP addresses for LLT over UDP, LLT may not work properly.

Workaround: Do not use alias IP addresses for LLT over UDP.

Limitations related to I/O fencing

This section covers I/O fencing-related software limitations.

Preferred fencing limitation when VxFEN activates RACER node re-election

The preferred fencing feature gives preference to more weighted or larger subclusters by delaying the smaller subcluster. This smaller subcluster delay is effective only if the initial RACER node in the larger subcluster is able to complete the race. If due to some reason the initial RACER node is not able to complete the race and the VxFEN driver activates the racer re-election algorithm, then the smaller subcluster delay is offset by the time taken for the racer re-election and the less weighted or smaller subcluster could win the race. This limitation though not desirable can be tolerated.

Stopping systems in clusters with I/O fencing configured

The I/O fencing feature protects against data corruption resulting from a failed cluster interconnect, or “split brain.” See the *Cluster Server Administrator's Guide* for a description of the problems a failed interconnect can create and the protection I/O fencing provides.

In a cluster using SCSI-3 based fencing, I/O fencing implements data protection by placing the SCSI-3 PR keys on both the data disks and coordinator disks. In a cluster using CP server-based fencing, I/O fencing implements data protection by placing the SCSI-3 PR keys on data disks and similar registrations on CP server. The VCS administrator must be aware of several operational changes needed when working with clusters protected by I/O fencing. Specific shutdown procedures ensure keys are removed from coordination points and data disks to prevent possible difficulties with subsequent cluster startup.

Using the reboot command rather than the shutdown command bypasses shutdown scripts and can leave keys on the coordination points and data disks. Depending on the order of reboot and subsequent startup events, the cluster may warn of a possible split brain condition and fail to start up.

Workaround: Use the shutdown -r command on one node at a time and wait for each node to complete shutdown.

Uninstalling VRTSvxvm causes issues when VxFEN is configured in SCSI3 mode with dmp disk policy (2522069)

When VxFEN is configured in SCSI3 mode with dmp disk policy, the DMP nodes for the coordinator disks can be accessed during system shutdown or fencing arbitration. After uninstalling VRTSvxvm package, the DMP module will no longer be loaded in memory. On a system where VRTSvxvm package is uninstalled, if VxFEN attempts to access DMP devices during shutdown or fencing arbitration, the system panics.

Node may panic if HAD process is stopped by force and then node is shut down or restarted [3640007]

A node may panic if the HAD process running on it is stopped by force and then it is shut down or restarted. This limitation is observed when you perform the following steps on a cluster node:

- 1 Stop the HAD process with the `force` flag.

```
# hstop -local -force
```

or

```
# hstop -all -force
```

- 2 Restart or shut down the node.

The node panics because forcefully stopping VCS on the node leaves all the applications, file systems, CVM, and other process online on that node. If the same node is restarted in this state, VCS triggers a fencing race to avoid data corruption. However, the restarted node loses the fencing race and panics.

Workaround: No workaround.

Limitations related to global clusters

- Cluster address for global cluster requires resolved virtual IP.
The virtual IP address must have a DNS entry if virtual IP is used for heartbeat agents.
- Total number of clusters in a global cluster configuration can not exceed four.
- Cluster may not be declared as faulted when Symm heartbeat agent is configured even when all hosts are down.
The Symm agent is used to monitor the link between two Symmetrix arrays. When all the hosts are down in a cluster but the Symm agent is able to see the replication link between the local and remote storage, it would report the heartbeat as ALIVE. Due to this, DR site does not declare the primary site as faulted.
- Configuring Veritas Volume Replicator for Zone Disaster Recovery is not supported for zone root replication. Oracle Solaris 11 supports zone root only on ZFS file system.

CP Server 6.0.5 client fails to communicate with CP Server 7.0 with certificates having 2048-bit keys and SHA256 hashing [IIP-5803]

On Solaris 11x64, if you upgrade a CP Server to 7.4 and accept to upgrade the certificates or configure a new 7.4 CP server, CPS clients on version 6.0.5 or below cannot communicate with the CP Server even after re-establishing trust relationships.

Workaround:

Upgrade the version of the CPS client to 6.1 or later.

Clusters must run on VCS 6.0.5 and later to be able to communicate after upgrading to 2048 bit key and SHA256 signature certificates [3812313]

In global clusters, when you install or upgrade VCS to 7.4 and you upgrade to 2048 bit key and SHA256 signature certificates on one site and the other site is on VCS version lower than 6.0.5, the clusters fail to communicate. The cluster communication will not be restored even if you restore the trust between the clusters. This includes GCO, Steward and CP server communication.

Workaround: You must upgrade VCS to version 6.0.5 or later to enable the global clusters to communicate.

Storage Foundation Cluster File System High Availability software limitations

These software limitations apply to the following products:

- Veritas InfoScale Storage
- Veritas InfoScale Enterprise

cfsmntadm command does not verify the mount options (2078634)

You must confirm the mount options are correct which are then passed to the `cfsmntadm` command. If the mount options are not correct, the mount fails and the CFSSMount resource will not come online. You can check the VCS engine log file for any mount failure messages.

Stale SCSI-3 PR keys remain on disk after stopping the cluster and deporting the disk group

When all nodes present in the Veritas InfoScale cluster are removed from the cluster, the SCSI-3 Persistent Reservation (PR) keys on the data disks may not get preempted. As a result, the keys may be seen on the disks after stopping the cluster or after the nodes have booted up. The residual keys do not impact data disk fencing as they will be reused or replaced when the nodes rejoin the cluster. Alternatively, the keys can be cleared manually by running the `vxfcntlclearpre` utility.

For more information on the `vxfcntlclearpre` utility, see the *Veritas InfoScale Administrator's Guide*.

Unsupported FSS scenarios

The following scenario is not supported with Flexible Storage Sharing (FSS):

Veritas NetBackup backup with FSS disk groups

Storage Foundation for Oracle RAC software limitations

These software limitations apply to Veritas InfoScale Enterprise.

Supportability constraints for normal or high redundancy ASM disk groups with CVM I/O shipping and FSS (3600155)

Normal or high redundancy ASM disk groups are not supported in FSS environments or if CVM I/O shipping is enabled.

Configure ASM disk groups with external redundancy in these scenarios.

Limitations of CSSD agent

The limitations of the CSSD agent are as follows:

- For Oracle RAC 11g Release 2 and later versions: The CSSD agent restarts Oracle Grid Infrastructure processes that you may manually or selectively take offline outside of VCS.
Workaround: First stop the CSSD agent if operations require you to manually take the processes offline outside of VCS.
 For more information, see the topic "Disabling monitoring of Oracle Grid Infrastructure processes temporarily" in the *Storage Foundation for Oracle RAC Configuration and Upgrade Guide*.
- The CSSD agent detects intentional offline only when you stop Oracle Clusterware/Grid Infrastructure outside of VCS using the following command: `crsctl stop crs [-f]`. The agent fails to detect intentional offline if you stop Oracle Clusterware/Grid Infrastructure using any other command.
Workaround: Use the `crsctl stop crs [-f]` command to stop Oracle Clusterware/Grid Infrastructure outside of VCS.

Oracle Clusterware/Grid Infrastructure installation fails if the cluster name exceeds 14 characters

Setting the cluster name to a value that exceeds 14 characters during the installation of Oracle Clusterware/Grid Infrastructure causes unexpected cluster membership issues. As a result, the installation may fail.

Workaround: Restart the Oracle Clusterware/Grid Infrastructure installation and set the cluster name to a value of maximum 14 characters.

Policy-managed databases not supported by CRSResource agent

The CRSResource agent supports only admin-managed database environments in this release. Policy-managed databases are not supported.

Health checks may fail on clusters that have more than 10 nodes

If there are more than 10 nodes in a cluster, the health check may fail with the following error:

```
vxgettext ERROR V-33-1000-10038
Arguments exceed the maximum limit of 10
```

The health check script uses the `vxgettext` command, which does not support more than 10 arguments.[2142234]

Cached ODM not supported in Veritas InfoScale environments

Cached ODM is not supported for files on Veritas local file systems and on Cluster File System.

Storage Foundation for Databases (SFDB) tools software limitations

The following are the SFDB tools software limitations in this release.

Parallel execution of `vxsfadm` is not supported (2515442)

Only one instance of the `vxsfadm` command can be run at a time. Running multiple instances of `vxsfadm` at a time is not supported.

Creating point-in-time copies during database structural changes is not supported (2496178)

SFDB tools do not support creating point-in-time copies while structural changes to the database are in progress, such as adding or dropping tablespaces and adding or dropping data files.

However, once a point-in-time copy is taken, you can create a clone at any time, regardless of the status of the database.

Oracle Data Guard in an Oracle RAC environment

SFDB tools cannot be used with RAC standby databases. SFDB tools can still be used with the primary database, even in a Data Guard Oracle RAC environment.

Known issues

This chapter includes the following topics:

- [Issues related to installation and upgrade](#)
- [Storage Foundation known issues](#)
- [Replication known issues](#)
- [Cluster Server known issues](#)
- [Storage Foundation and High Availability known issues](#)
- [Storage Foundation Cluster File System High Availability known issues](#)
- [Storage Foundation for Oracle RAC known issues](#)
- [Storage Foundation for Databases \(SFDB\) tools known issues](#)

Issues related to installation and upgrade

This section describes the known issues during installation and upgrade. These known issues apply to the following products:

- Veritas InfoScale Foundation
- Veritas InfoScale Storage
- Veritas InfoScale Availability
- Veritas InfoScale Enterprise

Switch fencing in enable or disable mode may not take effect if VCS is not reconfigured [3798127]

When you choose not to reconfigure Veritas Cluster Server (VCS), and set the fencing in enable or disable mode, it may not take effect. This is because the fencing mode switch relies on VCS reconfiguration.

Workaround: If you want to switch the fencing mode, when the installer shows "Do you want to re-configure VCS?", enter y to reconfigure VCS .

After the upgrade to version 7.4, the installer may fail to stop the Asynchronous Monitoring Framework (AMF) process [3781993]

After upgrading from old version product, when the product is stopped by CPI, the AMF process does not stop.

Workaround: There is no workaround for this issue.

LLT may fail to start after upgrade on Solaris 11 [3770835]

On Solaris 11, after you upgrade Storage Foundation for Oracle RAC/VCS/SFHA/SFCFSA to version 7.4, there may occur the error "It failed to start on <server_name>".

Workaround: To solve the issue:

- 1 After upgrade, restart the system.
- 2 Execute the following command:

```
# /opt/VRTS/install/installer -start
```

On SunOS, drivers may not be loaded after a reboot [3798849]

When the product installer stops the processes, it uses the `rem_drv` command to ensure the driver not be loaded back by the operating system. However, after a reboot, the system will not be able to load those drivers back, which makes our product unavailable after a reboot.

Workaround:

On SunOS, if drivers such as `vxdump`, `vxio` are not loaded after a reboot, you need to manually execute the following command to start your product:

```
/opt/VRTS/install/installer -start
```

On Oracle Solaris, drivers may not be loaded after stop and then reboot [3763550]

When the installer stops the processes, it uses the `rem_drv` command to prevent the drivers from being loaded back by the operating system. However, OS cannot load those drivers back after reboot, such as `vxdmp` and `vxio`.

Workaround: Start your product manually:

```
# /opt/VRTS/install/installer -start
```

During an upgrade process, the AMF_START or AMF_STOP variable values may be inconsistent [3763790]

If the value of `AMF_START` or `AMF_STOP` variables in the driver configuration file is '0' before an upgrade, then after the upgrade is complete, the installer changes the value to 1. Simultaneously, the installer also starts the Asynchronous Monitoring Framework (AMF) process.

Workaround: To resolve the issue, stop the AMF process and change the `AMF_START` or `AMF_STOP` value to 0.

Uninstallation fails on global zone on Solaris 11 if product packages are installed on both global zone and local zone [3762814]

If you have a product installed on both local zone and global zone, if you uninstall the product on global zone, the packages fails to be uninstalled.

Workaround: Log into the local zone and uninstall the packages of product from the local zone first.

On Solaris 11, when you install the operating system together with SFHA products using Automated Installer, the local installer scripts do not get generated (3640805)

On Solaris 11, when you use Automated Installer (AI) to install the Solaris 11 operating system together with SFHA products, the local installer scripts fail to get generated.

Workaround:

On the target system(s), execute the following script:

```
/opt/VRTSsfcp/bin/run-once
```

Stopping the installer during an upgrade and then resuming the upgrade might freeze the service groups (2574731)

The service groups freeze due to upgrading using the product installer if you stopped the installer after the installer already stopped some of the processes and then resumed the upgrade.

Workaround: You must unfreeze the service groups manually after the upgrade completes.

To unfreeze the service groups manually

- 1 List all the frozen service groups

```
# hagrps -list Frozen=1
```

- 2 Unfreeze all the frozen service groups:

```
# haconf -makerw
# hagrps -unfreeze service_group -persistent
# haconf -dump -makero
```

Installing VRTSvlic package during live upgrade on Solaris system non-global zones displays error messages [3623525]

While installing VRTSvlic package during live upgrade on Solaris system with non-global zones following error messages are displayed:

```
cp: cannot create /a/sbin/vxlicinst: Read-only file system
cp: cannot create /a/sbin/vxlicrep: Read-only file system
cp: cannot create /a/sbin/vxlictest: Read-only file system
```

Workaround: This message can be ignored. The vxlicinst, vxlicrep, vxlictest utilities are present in /opt/VRTSvlic/sbin/ inside a non-global zone.

VCS installation with CPI fails when a non-global zone is in installed state and zone root is not mounted on the node (2731178)

The CPI tries to boot a zone in installed state during installation/ or uninstallation. The boot fails if the underlying storage for zone root is not imported and mounted onto the node, causing the installation or uninstallation to fail.

Workaround: Make sure that the non-global zones are in running or configured state when CPI is invoked for installation or uninstallation.

Log messages are displayed when VRTSvcs is uninstalled on Solaris 11 [2919986]

The following message is displayed when you uninstall VRTSvcs package on Solaris 11 OS.

The following unexpected or editable files and directories were salvaged while executing the requested package operation; they have been moved to the displayed location in the image:

```
var/VRTSvcs/log -> /var/pkg/lost+found/var/VRTSvcs/log-20111216T122049Z
var/VRTSvcs/lock -> /var/pkg/lost+found/var/VRTSvcs/lock-20111216T122049Z
var/VRTSvcs -> /var/pkg/lost+found/var/VRTSvcs-20111216T122049Z
etc/VRTSvcs/conf/config
->/var/pkg/lost+found/etc/VRTSvcs/conf/config-20111216T122049Z
```

You can safely ignore this message as this is an expected behavior of IPS packaging. The files mentioned in the above message are not part of the package. As a result, uninstallation moves them to `/var/pkg/lost+found` directory.

Cluster goes into STALE_ADMIN_WAIT state during upgrade from VCS 5.1 to 6.1 or later [2850921]

While performing a manual upgrade from VCS 5.1 to VCS 6.1 or later, cluster goes in `STALE_ADMIN_WAIT` state if there is an entry of `DB2udbTypes.cf` in `main.cf`.

Installation of VRTSvcs package in VCS 5.1 creates a symbolic link for `Db2udbTypes.cf` file inside `/etc/VRTSvcs/conf/config` directory which points to `/etc/VRTSagents/ha/conf/Db2udb/Db2udbTypes.cf`. During manual upgrade, the VRTSvcs package for VCS 5.1 gets removed, which in turn removes the symbolic link for file `Db2udbTypes.cf` inside `/etc/VRTSvcs/conf/config` directory. After the complete installation of VRTSvcs for VCS 6.1 or later versions, because of absence of file `Db2udbTypes.cf` inside `/etc/VRTSvcs/conf/config`, cluster goes into `STALE ADMIN WAIT` state.

Workaround: Manually copy `DB2udbTypes.cf` from `/etc/VRTSagents/ha/conf/Db2udb` directory to the `/etc/VRTSvcs/conf/config` directory after the manual upgrade before starting HAD.

Flash Archive installation not supported if the target system's root disk is encapsulated

Veritas does not support SFCFSHA, SFHA, SF Oracle RAC, or SF Sybase CE installation using Flash Archive if the target system's root disk is encapsulated.

Make sure that the target system's root disk is unencapsulated before starting the installation.

The Configure Sybase ASE CE Instance in VCS option creates duplicate service groups for Sybase binary mount points (2560188)

The CPI installer does not check to see if Sybase binary mount points are already configured on systems, nor does it give an error message. It creates a duplicate service group for Sybase binary mount points.

This issue will be resolved in a later release.

The Installer fails to unload GAB module while installation of SF packages [3560458]

The Installer succeeds to upgrade SF package from 6.1.1 or 6.0.5 to 6.2.1 or later, but GAB module (for 6.1.1 or 6.0.5) fails to unload and remains in loaded state.

The issue is seen with the recent updates of Solaris 11U1 (SRU 8) or Solaris 11U2OS 11U1 (SRU 8). During un-installation of SFCFSHA, SFHA, SF Oracle RAC, SF Sybase CE or VCS packages, unloading of GAB fails.

Workaround: Restart the system. Restarting the system will unload the module successfully.

On Solaris 11 non-default ODM mount options will not be preserved across package upgrade (2745100)

On Solaris 11, before the package upgrade if Oracle Disk Manager (ODM) is mounted with non-default mount options such as nocluster, nosmartsync etc, these mount options will not get preserved after package upgrade.

There is no workaround at this time.

Upgrade fails because there is zone installed on the VxFS file system which is offline. The packages in the zone are not updated. (3319753)

If the zone installed on VxFS file system is under VCS control, and the VxFS file system is in offline state, the upgrade fails because it's not able to update the packages in the zones.

Workaround:

Check the status of the mounted file system which has the zones on it. If the file system is offline, you need to first bring it online, then do the upgrade, so that the packages in the local zone can be updated.

If you choose to upgrade nodes without zones first, the rolling upgrade or phased upgrade is not blocked in the beginning, but fails later (3319961)

If you choose to upgrade nodes without zones first, the rolling upgrade or phased upgrade is not blocked in the beginning, but fails later when you start to upgrade the nodes that have zones installed.

This issue occurs in the following scenarios:

- A zone is installed on a Cluster File System (CFS) on one of the nodes.
- A node is installed on a Veritas File System (VxFS) on one of the nodes, and node is under Cluster Server (VCS) control.

Workaround:

- 1 Before you upgrade, uninstall the zones on the nodes which have zones installed. Enter:.

```
zoneadm -z zonename uninstall
```

- 2 Run the installer to run the upgrade.
- 3 After the upgrade completes, reinstall the zones.

Upgrades from previous SF Oracle RAC versions may fail on Solaris systems (3256400)

The `vxio` and `vxdump` modules may fail to stop on Solaris systems during upgrades from previous SF Oracle RAC versions. As a result, the upgrade fails to complete successfully.

Workaround: If `vxio` and `vxdump` fail to stop and no other issues are seen during upgrade, continue with the upgrade and restart the system when the product installer prompts. After the reboot, use the installer to start the product again by entering:

```
# /opt/VRTS/install/installer -start
```

Note: Do not use the response file to upgrade in this situation.

After a locale change restart the vxconfig daemon (2417547, 2116264)

You need to restart the vxconfig daemon you change the locale of nodes that use it. The vxconfig daemon starts at boot. If you have changed locale, you need to restart the daemon.

Workaround: Refer to the *Storage Foundation Cluster File System High Availability Administrator's Guide* for the section, "vxconfigd daemon recovery."

Verification of Oracle binaries incorrectly reports as failed during Oracle Grid Infrastructure installation

The verification of Oracle binaries may incorrectly report as failed during the Oracle Grid Infrastructure installation using the SF Oracle RAC installer. The message is erroneously reported due to a break in passwordless SSH communication. The SSH communication fails because execution of the `root.sh` script changes the owner of the operating system root directory to the grid user directory.

Live upgrade of the InfoScale product may detect the wrong product or ask for the license repeatedly (3870685)

If the product of previous release is installed on Solaris 11.1 or Solaris 11.2 and the OS upgrades to Solaris 11.3 through live upgrade, installer may detect the wrong product or ask for the license repeatedly during the upgrade even though the license is provided. As a result, a chroot command core dump issue occurs when the target BE is Solaris 11.3 SPARC or Solaris 11.3 x86.

Workaround:

Currently there is no resolution provided by installer. You can check with Oracle for possible workarounds.

In RHEV environment, if you stop the SF service and then start it by installer, the permission on dmpnode will get lost (3870111)

If you stop the SF service and then start it again by installer, the permission on dmpnode will get lost.

For example:

You check the permission on dmpnode at first:

```
# ls -l /dev/vx/dmp/huawei-hvs88t0_0  
  
brwxrwxrwx 1 vdsd kvm 201, 1728 Jan 22 02:08  
/dev/vx/dmp/huawei-hvs88t0_0
```

You stop the SF service by installer:

```
# /opt/VRTS/install/installer -stop
```

You check the permission on dmpnode:

```
# ls -l /dev/vx/dmp/huawei-hvs88t0_0  
  
brwxrwxrwx 1 vdsd kvm 201, 1728 Jan 22 02:08  
/dev/vx/dmp/huawei-hvs88t0_0
```

You start the SF service by installer:

```
# /opt/VRTS/install/installer -start
```

When you check the permission on dmpnode, the `user/group/permission` gets lost:

```
# ls -l /dev/vx/dmp/huawei-hvs88t0_0  
  
brw----- 1 root root 201, 1712 Jan 22 02:29  
/dev/vx/dmp/huawei-hvs88t0_0
```

Workaround:

You can restore the permission through `chmod`.

The installer fails to upgrade the product packages on Solaris 11 during an upgrade to InfoScale 7.4 (3896530)

When you try to upgrade to InfoScale 7.4 on Solaris 11, the installer fails to upgrade the product packages on the nodes which have Symantec publisher.

Workaround:

Use `pkg unset-publisher Symantec` command to remove the Symantec publisher, then use the installer to upgrade to InfoScale 7.4.

Resource faults during Rolling upgrade due to perl changes (3930605)

During rolling upgrade in 7.3.1, resources can go into FAULTED state due to Perl changes.

During rolling upgrade, after Phase I upgrade on the first sub cluster is complete and Phase I upgrade on the second sub cluster is about to start, if the `VCS Service Groups` is switched back to the first sub cluster, resources go into PARTIAL , STARTING, or FAULTED state.

Workaround: Complete the upgrade process, the resources become online after the upgrade.

Storage Foundation known issues

This section describes the known issues in this release of Storage Foundation (SF). These known issues apply to the following products:

- Veritas InfoScale Foundation
- Veritas InfoScale Storage
- Veritas InfoScale Enterprise

Dynamic Multi-Pathing known issues

This section describes the known issues in this release of Dynamic Multi-Pathing (DMP).

Vxconfigd may core dump after suppressing paths of a PowerPath device (3869111)

If you try to suppress the paths of a PowerPath device by using the `vxddmpadm exclude path/dmptnodename` command or using the `vxdiskadm` utility, the `vxconfigd` daemon may core dump.

Workaround:

To solve it, choose one of the following methods:

- Restart the system.
- Exclude the PowerPath devices using the controller option, e.g.

```
vxddmpadm exclude ctrl=emcp
```

- Using the `vxdiskadm` utility by suppressing all paths through a controller from VxVM's view.

Veritas Volume Manager known issues

Core dump issue after restoration of disk group backup (3909046)

After you restore a disk group backup using the `vxconfigrestore` command, it is possible that some configuration copies remain in a disabled state. As a result, VxVM generates a core dump when you view the list of disk groups after the restore operation.

Stack Trace:

```
#0 0x00000033a3432625 in raise () from /lib64/libc.so.6
#1 0x00000033a3433e05 in abort () from /lib64/libc.so.6
#2 0x00000033a342b74e in __assert_fail_base () from /lib64/libc.so.6
#3 0x00000033a342b810 in __assert_fail () from /lib64/libc.so.6
#4 0x00000000005060f1 in req_dg_get_info_common (clnt=0x1af1750,
dg=0x7fc330004bb0) at
dg.c:3261
#5 0x00000000005059f5 in req_dg_get_info_name (clnt=0x1af1750,
req=0x1b03f78)
at dg.c:3057
#6 0x00000000005b2e2d in vold_process_request (arg=0x18f5f20) at
request.c:1997
#7 0x00000033a3807a51 in start_thread () from /lib64/libpthread.so.0
#8 0x00000033a34e896d in clone () from /lib64/libc.so.6
```

Workaround: Restart the VxVM configuration daemon.

```
# vxconfigd -kr reset
```

Failed verifydata operation leaves residual cache objects that cannot be removed (3370667)

When you use the `verify data` command, and type

```
# vradmin -g dgname verifydata rvgname IPaddress cachesize=size
```

the command may fail and leave residual cache objects that cannot be removed.

Workaround:

To solve this problem, choose different ways based on different residual cache objects.

To explicitly clean up the cache object that is associated to SO snapshots:

1. List the SO snapshots that are created on a cache object by typing:

```
# vxcache -g dname listvol volumenam
```

2. Unmount the listed snapshots.
3. Remove the snapshot volume. Type:

```
# vxedit -g dname -fr rm volumenam
```

It also removes the cache object.

To clean up the cache object that is not associated to the snapshot volume but associated to the cache volume:

1. Stop the cache object by typing:

```
# vxcache -g dname stop cacheobject_name
```

2. Remove the cache object. Type:

```
# vxedit -g dname -rf rm cacheobject_name
```

It also removes the cache volume.

LUNs claimed but not in use by VxVM may report “Device Busy” when it is accessed outside VxVM (3667574)

When a LUN claimed by Veritas Volume Manager (VxVM) is accessed, the open on the device gets cached for performance improvement. Due to this, some OS utilities which require exclusive access reports `Device Busy`.

Workaround:

To solve this issue, either exclude these LUNs from the VxVM view or disable them by typing `vxddmpadm disable dmpnodename=<> CLI`.

For more details, refer to the tech note:

https://www.veritas.com/support/en_US/article.TECH227660.

Unable to set master on the secondary site in VVR environment if any pending I/O's are on the secondary site (3874873)

There is deadlock situation with the cluster reconfiguration and the network disconnection (serialization) on RVG object. Wherein, the reconfiguration quiesces the disk level I/O's and it expects the replica object to be disconnected. The Rlink cannot be disconnected unless the underlying I/O's are completed and the reconfig thread quiesces these I/Os at disk level.

Workaround:

Pause the Rlink on the primary site and then set master on the secondary slave node.

vxdisksetup -if fails on PowerPath disks of sizes 1T to 2T [3752250]

vxdisksetup -if fails on PowerPath disks of sizes 1T to 2T with the following message:

```
VxVM vxdisksetup ERROR V-5-2-4006 Disk emcpower48 contains auto:\
none DA record emcpower48s2
```

Workaround:

- 1 Format the disk to EFI label:

```
# format
```

- 2 Remove the formatted disk from VxVM control:

```
# vxdisk rm emcpower48s2
```

- 3 Scan the disk again:

```
# vxdisk scandisks
```

The disk should show up as emcpower48, without the s2 suffix.

- 4 Set up the disk:

```
# vxdisksetup -if emcpower48
```

VRAS `verifydata` command fails without cleaning up the snapshots created [3558199]

The `vradmin verifydata` and the `vradmin syncrvg` commands leave behind residues if terminated abnormally. These residues can be snapshot volumes or mount points.

Workaround: Remove the snapshot volumes and unmount the mount points manually.

Root disk encapsulation fails for root volume and swap volume configured on thin LUNs (3538594)

Root disk encapsulation fails if the root disk configuration on a thin LUN includes volumes such as `var`, `usr`, or `home`, in addition to the root volumes and the swap volumes. Root disk encapsulation is not supported in this configuration.

Workaround:

There is no workaround.

The `vxdisk resize` command does not claim the correct LUN size on Solaris 11 during expansion of the LUN from array side (2858900)

The `vxdisk resize` command fails on Solaris 11 during expansion of the LUN from array side. The `vxdisk resize` command does not claim correct LUN size on Solaris 11 during expansion of the LUN from array side. This is because of Oracle bug -19603615. On Solaris 11, the `vxdisk resize` command may exit without any error, returning incorrect LUN size or failing with similar error as follows:

```
bash# vxdisk -g testdg resize disk01 length=8g
VxVM vxdisk ERROR V-5-1-8643 Device disk01: resize failed:\
Operation would block
```

Workaround:

There is no workaround available which can work in all the configuration. In some specific configurations, the following workaround works:

After expansion of LUN from array side, run `format -d` command and then run `vxdisk resize` command.

SmartIO VxVM cache invalidated after relay layout operation (3492350)

If a relay layout operation is done on a volume that has SmartIO VxVM caching enabled, the contents of the cache for the volume may be invalidated.

Workaround:

This behavior is expected. There is no workaround.

Disk greater than 1TB goes into error state [3761474, 3269099]

If a path of a device having multiple paths is labelled with the EFI format using an operating system command such as `format`, the `vxdisk list` command output shows the device in error state.

Workaround:

This issue is a Solaris OS issue. There is no workaround for this issue.

Importing an exported zpool can fail when DMP native support is on (3133500)

On Solaris, when the tunable `dmp_native_support` is set to `on`, importing an exported zpool using the command `zpool import poolname` can fail with following error:

```
Assertion failed: rn->rn_nozpool == B_FALSE, file
../common/libzfs_import.C,
line 1084, function zpool_open_func
Abort (core dumped)
```

Workaround:

Import the zpool using the following command, specifying the DMP device directory:

```
# zpool import -d /dev/vx/dmp poolname
```

Server panic after losing connectivity to the voting disk (2787766)

This issue occurs on A/P arrays. If the voting disk loses connectivity to the primary paths, DMP takes some time to analyze the error and fail over the paths. During this time, the `cssd` reports a timeout and panics. When using Oracle ASM over DMP devices, set the `disktimeout` parameter to an appropriate value. This parameter indicates the maximum time allowed for a voting file I/O to complete. If this time is exceeded, the voting disk is marked as offline.

The default of `disktimeout` is 200. If the value of the tunable is less than this value, reset the value to the default value.

Workaround:

To set the `disktimeout` to 200:

```
$CRS_HOME/bin/crsctl set css disktimeout 200 [-force] test
```

Performance impact when a large number of disks are reconnected (2802698)

If the storage connectivity is lost to part of the storage, the disk group configuration copy is rebalanced to the disks that have connectivity. For example, if the storage for an entire enclosure is removed from a disk group with multiple enclosures. The rebalancing process takes time, during which time the `vxconfigd` daemon is busy and does not respond to commands.

device.map must be up to date before doing root disk encapsulation (2202047)

If you perform root disk encapsulation while the `device.map` file is not up to date, the `vxdiskadm` command displays the following error:

```
VxVM vxencap INFO V-5-2-5327 Missing file: /boot/grub/device.map
```

Workaround: Before you perform root disk encapsulation, run the the following command to regenerate the `device.map` file:

```
# grub-install --recheck /dev/sdb
```

Veritas Volume Manager (VxVM) might report false serial split brain under certain scenarios (1834513)

VxVM might detect and report a false serial split brain when all of the following conditions are met:

- One or more arrays that provide the shared storage for the cluster are being powered off
- At the same time when the arrays are being powered off, an operation that requires an internal transaction is initiated (such as VxVM configuration commands)

In such a scenario, disk group import will fail with a split brain error and the `vxsplitlines` output will show 0 or 1 pools.

Workaround:

To recover from this situation

- 1 Retrieve the disk media identifier (dm_id) from the configuration copy:

```
# /etc/vx/diag.d/vxprivutil dumpconfig device-path
```

The dm_id is also the serial split brain id (ssbid)

- 2 Use the dm_id in the following command to recover from the situation:

```
# /etc/vx/diag.d/vxprivutil set device-path ssbid=dm_id
```

Suppressing the primary path of an encapsulated SAN boot disk from Veritas Volume Manager causes the system reboot to fail (1933631)

If you suppress the primary path of an array from VxVM control and then reboot the system, the system boot fails.

If you have an encapsulated SAN boot device with multiple primary paths, the issue occurs when you suppress the first primary path. When you configure a SAN boot device, the primary path is set as a boot device. In general, the first path of the SAN boot device corresponds to the first configured path during SAN boot. Even if another primary path is configured as a boot device, suppressing the first device from VxVM causes the boot to fail.

Workaround:

When the boot device is suppressed from VxVM, change the OS boot device sequencing accordingly.

For Solaris SPARC system, use the `eeeprom boot-device` command to set the boot device sequencing.

For Solaris x86-64 systems, use the `eeeprom bootpath` command to set the boot device sequencing.

After changing the preferred path from the array side, the secondary path becomes active (2490012)

For EVA arrays, DMP requires that the prefer bit is static. If the prefer bit is not static, issues like the following may occur. After changing the prefer path of LUN from the array side, and performing a disk discovery (`vxdisk scandisks`) from the host, the secondary path becomes active for the LUN.

Workaround:

To work around this issue

- 1 Set the pref bit for the LUN.
- 2 Perform disk discovery again:

```
# vxdisk scandisks
```

Disk group import of BCV LUNs using -o updateid and -ouseclonedev options is not supported if the disk group has mirrored volumes with DCO or has snapshots (2831658)

VxVM uses guid stored in configuration to uniquely identify all objects. The data change object (DCO) volume stores the guid of mirrors and snapshots. If the disk group is imported with `-o updateid` and `-o useclonedev`, it changes the guid of objects in VxVM configuration database and the guids stored in the DCO volume are not updated. The operations involving DCO cannot find objects with the stored guid. This could lead to failure of certain operations involving DCO or could lead to unexpected behavior.

Workaround:

No workaround available.

After devices that are managed by EMC PowerPath lose access to storage, Veritas Volume Manager commands are delayed (2757198)

In an environment which includes devices that are managed by EMC PowerPath, a storage loss causes Veritas Volume Manager commands to be delayed. In the event of storage loss, VxVM sends SCSI inquiry to each LUN path to check the health of path, which are delayed by the presence of EMC PowerPath.

Workaround:

There is no workaround available.

vxresize does not work with layered volumes that have multiple plexes at the top level (3301991)

If a layered volume has multiple plexes at the top level, `vxresize` does not work. For example, if you add a mirror to a concat-mirror volume for a third mirror snapshot. The `vxresize` operation fails with the following message:

```
VxVM vxassist ERROR V-5-1-2528 Volume volname built on layered volumes
have multiple plexes
```

VxVM vxresize ERROR V-5-1-4703 Problem running vxassist command for volume volname, in diskgroup dgroup

Workaround:

To resize the volume

- 1 After adding the mirror to the volume, take a snapshot using the plex.
- 2 Grow the volume and snapshot volume with `vxresize`
- 3 Reattach the snapshot volume to the source volume.

In a clustered configuration with Oracle ASM and DMP and AP/F array, when all the storage is removed from one node in the cluster, the Oracle DB is unmounted from other nodes of the cluster (3237696)

In a clustered configuration with Oracle ASM and DMP and AP/F array, when you remove all the storage from one node in the cluster, I/O is expected to fail on this node. Due to an issue with the Oracle ASM configuration, the Oracle database is unmounted from other nodes of the cluster. This issue is not seen if you delay the I/O failure from DMP. The Oracle database works fine on other node.

Workaround:

Increase the `dmp_lun_retry_timeout` tunable value to 300 with following command.

```
# vxddmpadm settune dmp_lun_retry_timeout=300
```

When all Primary/Optimized paths between the server and the storage array are disconnected, ASM disk group dismounts and the Oracle database may go down (3289311)

The Oracle database shows an I/O error on the control file, but there was no I/O error seen on any DMP device. When all Primary/Optimized paths are disconnected, DMP fails over to other available paths but the failover takes time. In the meantime, the application (ASM/Oracle database) times out the I/O.

The ASM alert log file displays messages such as the following:

```
Errors in file /u01/app/oracle/diag/rdbms/orcl/orcl2/trace/orcl2_ckpt_6955.trc:
ORA-00221: error on write to control file
ORA-00206: error in writing (block 4, # blocks 1) of control file
ORA-00202: control file: '+DATA_P6/ORCL/CONTROLFILE/current.261.826783133'
ORA-15081: failed to submit an I/O operation to a disk
ORA-15081: failed to submit an I/O operation to a disk
```

```
Wed Oct 09 14:16:07 2013
WARNING: group 2 dismounted: failed to read virtual extent 0 of file 261
Wed Oct 09 14:16:07 2013
USER (ospid: 6955): terminating the instance due to error 221
Wed Oct 09 14:16:07 2013
WARNING: requested mirror side 2 of virtual extent 0 logical extent 1 offset
16384
is not allocated; I/O request failed
WARNING: requested mirror side 3 of virtual extent 0 logical extent 2 offset
16384
is not allocated; I/O request failed
```

The above issue may occur when the server is configured as follows:

DB: Oracle 12c

Volume Manager: ASM

Multi-pathing Solutions: DMP

OS: Solaris

Disk Array : HP EVA in ALUA mode

Workaround:

The following workaround can reduce the probability of this issue, and when you see this issue, you could use Oracle commands to start the database manually.

Increase the application time out and make the following changes to reduce the time taken to mark the path as offline:

- In the `/kernel/drv/fp.conf` file, add `fp_offline_ticker=15`.
- In the `/kernel/drv/fcp.conf` file, add `fcp_offline_delay=10`.

Running the `vxdisk disk set clone=off` command on imported clone disk group luns results in a mix of clone and non-clone disks (3338075)

If you do not specify a disk group name, the `vxdisk set` operation works on the `dmname` rather than the `daname`. If a `dmname` is the same as an existing `daname`, the `vxdisk set` operation reflects on the `dm` name.

Workaround: Use the following command syntax to set the attributes:

```
vxdisk -g diskgroup_name set dmname clone=off
```

For example:

```
vxdisk -g dg1 set eva4k6k0_12 clone=off
```

The administrator must explicitly enable and disable support for a clone device created from an existing root pool (3110589)

A non-rpool is a clone of the existing root pool. When native support is enabled, DMP does not touch the clone root pool because the clone may or may not have the VxVM package.

Workaround: To add or remove DMP support for a clone boot device, the administrator must boot through the clone and turn on/off `dmp_native_support`.

Restarting the vxconfigd daemon on the slave node after a disk is removed from all nodes may cause the disk groups to be disabled on the slave node (3591019)

The issue occurs if the storage connectivity of a disk is removed from all the nodes of the cluster and the `vxconfigd` daemon is restarted on the slave node before the disk is detached from the slave. All the disk groups are in the `dgddisabled` state on the slave node, but show as `enabled` on the other nodes.

If the disk was detached before the `vxconfigd` daemon is restarted, the issue does not occur.

In a Flexible Storage Sharing (FSS) environment, removing the storage connectivity on a node that contributes DAS storage to a shared disk group results in global connectivity loss because the storage is not connected elsewhere.

Workaround:

To prevent this issue:

Before restarting the `vxconfigd` daemon, if a disk in a shared disk group has lost connectivity to all nodes in the cluster, make sure that the disk is in the `detached` state. If a disk needs to be detached, use the following command:

```
# vxdisk check diskname
```

To resolve the issue after it has occurred:

If `vxconfigd` is restarted before the disks got detached, remove the node from the cluster and rejoin the node to the cluster.

Failback to primary paths does not occur if the node that initiated the failover leaves the cluster (1856723)

When CVM is configured on non-A/A storage, if a node loses access to the storage through all the primary paths, then all the nodes in the cluster switches to the secondary paths. If the node which raised the protocol leaves the cluster and if all

the rest of the nodes in the cluster are seeing the primary paths as healthy, then failback to primary paths never happens.

Issues if the storage connectivity to data disks is lost on a CVM slave node while vxconfigd was not running on the node (2562889)

If storage connectivity to data disks is lost on a CVM slave node while vxconfigd was not running on the node, this may result in following issues when vxconfigd comes up on this node:

- The shared disk groups on the disconnected storage are marked as dgdisabled on the slave node only.
- The shared disk groups are available to rest of the cluster nodes but no transactions, such as VxVM configuration changes, are possible on any shared disk group.
- Attempts to deport such shared disk groups will fail.

Workaround:

Do one of the following:

- Remove the faulty slave node out of CVM cluster, restore storage connectivity, and rejoin the node to the cluster.
- Restart vxconfigd on the CVM master node.

The vxcdsconvert utility is supported only on the master node (2616422)

The vxcdsconvert utility should be run only from the master node, not from the slave nodes of the cluster.

Re-enabling connectivity if the disks are in local failed (lfailed) state (2425977)

In a Cluster Volume Manager (CVM) cluster, you can disable connectivity to the disks at the controller or enclosure level with the vxmpadm disable command. In this case, CVM may place the disks into the lfailed state. When you restore connectivity with the vxmpadm enable command, CVM may not automatically clear the lfailed state. After enabling the controller or enclosure, you must run disk discovery to clear the locally failed state.

To run disk discovery

- ◆ Run the following command:

```
# vxdisk scandisks
```

Issues with the disk state on the CVM slave node when vxconfigd is restarted on all nodes (2615680)

When a CVM master node and a slave node have lost storage access, and `vxconfigd` is restarted on all nodes, the disk state on the CVM slave node shows as invalid.

Plex synchronization is not completed after resuming synchronization on a new master when the original master lost connectivity (2788077)

When you run `vxrecover -o force`, it recovers only one subvolume and it cannot detect that the rest of the volume needs recovery.

When you run the `vxassist mirror` command, you run the `vxplex att` command serially on each subvolume. If the failure happens before you start the `attach` operation (need to mark the concerned plex as the `attach` operation is in progress), `vxrecover` will not redo the `attach` operation because it cannot find any record of the `attach` operation in progress.

Workaround:

Run the following command on each subvolume to manually recover the complete volume:

```
# usr/lib/vxvm/type/fsgen/vxplex -U fsgen -g diskgroup \  
-o force useopt att volume plex
```

A master node is not capable of doing recovery if it cannot access the disks belonging to any of the plexes of a volume (2764153)

A master node with missing disks is not capable of doing recovery, as it does not have access to the disks belonging to any of the plexes of a volume.

Workaround:

If other nodes have access to the storage, they can do the recovery. Switch the master role to some other node with better storage connectivity.

CVM fails to start if the first node joining the cluster has no connectivity to the storage (2787713)

If the first node joining the cluster has no connectivity to disks, the import of shared disk groups fails. Other nodes that join the cluster later assume that the auto-import of disk groups is already done as part of the existing cluster processing.

Workaround:

Perform a master switch to the node that has connectivity to the disks. Then import the disk groups manually.

CVMVolDg agent may fail to deport CVM disk group when CVMDeportOnOffline is set to 1

When CVMDeportOnOffline is set to 1, the CVM disk group is deported based on the order in which the CVMVolDg resources are taken offline. If the CVMVolDg resources in the disk group contain a mixed setting of 1 and 0 for the `CVMDeportOnOffline` attribute, the disk group is deported only if the attribute value is 1 for the last CVMVolDg resource taken offline. If the attribute value is 0 for the last CVMVolDg resource taken offline, the disk group is not deported.

Workaround: If multiple CVMVolDg resources are configured for a shared disk group and the disk group is required to be deported during offline, set the value of the `CVMDeportOnOffline` attribute to 1 for all of the resources.

The vxsnap print command shows incorrect value for percentage dirty [2360780]

The `vxsnap print` command can display the percentage of regions that differ between snapshots, shown as the %dirty. In SF 6.0, if this command is run while the volumes are online and being actively used, the shown %dirty may lag from actual percentage dirty for instant snap data cache object (DCO) volumes. That is, the command output may show less %dirty than actual.

For Solaris 11.1 or later, uninstalling DMP or disabling DMP native support requires steps to enable booting from alternate root pools (3178642)

For Solaris 11.1 or later, after you uninstall the VxVM package or after you turn off DMP native support, you may see this issue. After reboot, the root pool containing the active boot environment is migrated to the OS device but alternate root pools continue to show DMP device. The status of the alternate root pools and their DMP devices is shown as "UNAVAIL".

```
pool: crpool
state: UNAVAIL
status: One or more devices are unavailable in response to persistent
errors. There are insufficient replicas for the pool to continue
functioning.
action: Destroy and re-create the pool from a backup source. Manually
marking the device repaired using 'zpool clear' or 'fmadm repaired'
may allow some data to be recovered.
Run 'zpool status -v' to see device specific details.
scan: none requested
config:
```

NAME	STATE	READ	WRITE	CKSUM
crpool	UNAVAIL	0	0	0
emc_clariion1_82s0	UNAVAIL	0	0	0

The tunable parameter `dmp_native_support` only unconfigures DMP for the single root pool containing the active boot environment. If the setup has any alternate root pools, for which DMP native support was enabled, then the alternate root pools continue to show the DMP device. If the alternate root pool is configured in the current boot environment and DMP support is removed, the DMP devices required for ZFS are not found. The DMP devices and the root pools display the state as "UNAVAIL".

Workaround:

Even though the status of alternate root pool is "UNAVAIL", the system is bootable using the disk containing the alternate root pool. Reboot the system with the disk containing the alternate root pool. The system comes up with the root pool using the DMP device.

For Solaris 11.1 or later, after enabling DMP native support for ZFS, only the current boot environment is bootable (3157394)

After enabling DMP native support for ZFS on Solaris 11.1 or later, only the current boot environment (BE) is bootable. Any alternate BEs in the same root pool are not bootable. This situation occurs because the DMP native support configures the ZFS root pool so that only DMP can import the root pool. If you attempt to boot the system from the alternate BE, the system panics with the following message:

```
NOTICE: zfs_parse_bootfs: error 19
Cannot mount root on rpool/193 fstype zfs

panic[cpu0]/thread=10012000: vfs_mountroot: cannot mount root
```

```
Warning - stack not written to the dumpbuf
000000001000fa00 genunix:main+17c (1, 100dc958, 12d5c00, 124702c, 0, 10828000)
%l0-3: 0000000010010000 0000000000000000 00000000100dc800 0000000000000000
%l4-7: 0000000010012000 0000000000000000 000000001038f7c0 000000000104c800
```

Workaround:

To enable booting from another BE, configure the ZFS root pool so that it can be imported without DMP.

To configure ZFS root pool to enable booting from all the BEs

- 1 At the OBP PROM, run the following command to list all the BEs:

```
ok> boot -L
```

- 2 Use the following command to boot from the BE for which DMP native support for ZFS is enabled.

```
ok> boot -Z rpool/ROOT/BE_name
```

- 3 After booting through new BE, disable the DMP native support using the following command:

```
# vxdmpadm settune dmp_native_support=off
```

The system is now bootable from any BEs in the ZFS root pool.

When dmp_native_support is set to on, commands hang for a long time on SAN failures (3084656)

When dmp_native_support is set to on, on SAN failures, commands that do I/O operations to the root file system or I/O to disks that contain the root pool may hang for about 1-5 minutes. The commands include commands like "zpool status", or telnet initiated to connect the system. The hang is seen because the drivers below the DMP layer take more time to report the I/O failure when some of the paths to the disk containing the root pool are disconnected. This situation should not lead to any root pool data corruption.

Workaround:

This hang cannot be avoided but the hang time can be reduced by tuning the following parameters

To tune the parameters

- 1 In the `/kernel/drv/fp.conf` file, set

```
fp_offline_ticker=15
```

- 2 In the `/kernel/drv/fcp.conf` file, set

```
fcp_offline_dely=10
```

- 3 Reboot the sytem to apply the changes.

These steps reduce the hang time to a maximum of 1 minute.

vxdisk export operation fails if length of hostprefix and device name exceeds 30 characters (3543668)

If the combined length of the hostprefix and the device name exceeds 30 characters, the vxdisk export operation fails with the following error message:

```
VxVM vxdisk ERROR V-5-1-18318 Device c6t50060E8005655501d86s2: Name too
long for export. Length of Hostprefix + Disk accessname should not exceed
30 characters. Please see vxdctl(1M) man page for information on setting
user-specified hostprefix.
```

Workaround:

Use the enclosure-based naming (EBN) scheme instead of the operating system naming (OSN) scheme. OSN naming typically contains more characters and is not as intuitive. If the EBN name combined with the hostprefix exceeds 30 characters, you can manually set the hostprefix to a smaller size using the `vxdctl set hostprefix=value` command, where *value* is the new hostprefix.

Systems may panic after GPT disk resize operation (3930664)

After you resize the GPT disks using the following command, you may experience a system panic issue, `# vxdisk resize <disk_name> length=<new_size>`. This issue occurs if your deployment setup includes GPT disks partition.

No workaround available to resolve this issue. So you must not resize GPT disks, and to recover the system, wait for the system to restart.

Veritas File System known issues

This section describes the known issues in this release of Veritas File System (VxFS).

The VxFS file system with local scope enabled may hang if two or more nodes are restarted simultaneously (3944891)

In a scenario where the VxFS file system has local scope enabled, restarting two or more nodes simultaneously causes a deadlock which leads the file system to hang.

Workaround:

Reboot all the nodes.

Docker does not recognize VxFS backend file system

When VxFS is used as backing filesystem to run the docker daemon, the following error is displayed:

```
Backing Filesystem: unknown
```

The link for this issues in Github is: <https://github.com/docker/docker/issues/14847>

Workaround:

VxFS is recognized as backing filesystem in the Docker upstream.

Warning message sometimes appear in the console during system startup (2354829)

During system startup, following messages sometimes appear in system console:

```
WARNING: couldn't allocate SDT table for module vxfs
WARNING: couldn't allocate FBT table for module vxfs
Loading smf(5) service descriptions: 2/2
```

These warnings indicate that the SDT and FBT DTrace probes might not be available for the VxFS module. The VxFS module still loads and works correctly. Dtrace SDT/FBT has limits on the size of module that it can support. Since the VxFS module exceeds the size that Dtrace can support, SDT and FBT Dtrace probes might not work for VxFS.

Workaround: There is no workaround for this issue.

vxresize may fail when you shrink a file system with the "blocks are currently in use" error (3762935)

The vxresize shrink operation may fail when active I/Os are in progress on the file system which is being shrunk to a size closer to its current usage. You see a message similar to the following example:

```
UX:vxfs fsadm: ERROR: V-3-20343: cannot shrink /dev/vx/rdisk/dg1/voll -
blocks are currently in use. VxVM vxresize ERROR V-5-1-7514 Problem
running fsadm command for volume voll, in diskgroup dg1
```

Workaround: Re-run the shrink operation after stopping the I/Os.

On Solaris11U2, /dev/odm may show 'Device busy' status when the system mounts ODM [3661567]

If the system tries to mount Oracle Disk Manager (ODM) in a mode which is not supported by the installed license, the later ODM mount may not succeed and shows /dev/odm device busy error.

Workaround: There are two ways to resolve it.

Remove /dev/odm mount point and recreate it. Or reboot the system and then mount /dev/odm.

Delayed allocation may be turned off automatically when one of the volumes in a multi-volume file system nears 100%(2438368)

Delayed allocation may be turned off automatically when one of the volumes in a multi-volume file system is in almost full usage, even if other volumes in the file system have free space.

Workaround: After sufficient space is freed from the volume, the delayed allocation automatically resumes.

The file system deduplication operation fails with the error message "DEDUP_ERROR Error renaming X checkpoint to Y checkpoint on filesystem Z error 16" (3348534)

The file system deduplication operation fails with the error message "DEDUP_ERROR Error renaming X checkpoint to Y checkpoint on filesystem Z error 16", due to the failure in unmounting the checkpoint.

Workaround: Retry the deduplication operation to resolve the problem.

Oracle Disk Manager (ODM) may fail to start after upgrade to 7.4 on Solaris 11 [3739102]

ODM may fail to start after upgrade to 7.4 on Solaris 11.

Workaround: Manually start ODM service by entering:

```
# /lib/svc/method/odm start
```

On the cluster file system, clone dispose may fail [3754906]

In case of clone dispose on cluster, if Veritas File System (VxFS) fails to unlink clones, the specific fset is marked as bad incore and the fullfsck flag is marked on the file system.

Workaround: Run full fsck on the file system, and it will complete the extop processing required for the clone removal.

VRTSvxfs verification reports error after upgrading to 7.4 [3463479]

Upgraded to 7.4, the VRTSvxfs package cannot pass the verification check with the pkg verify VRTSvxfs command. You can see error messages similar to the following:

```
# pkg verify VRTSvxfs
PACKAGE                                STATUS
pkg://Veritas/VRTSvxfs                 ERROR
    driver: vxfs
        etc/name_to_major: 'vxfs' entry not present
```

Workaround: Use the following command to fix this issue:

```
# pkg fix VRTSvxfs
```

spfile created on VxFS and ODM may contain uninitialized blocks at the end (3760262)

spfile created on VxFS and ODM may contain uninitialized blocks at the end due to space allocation with file system block size alignment. This is harmless and does not cause any problem to Oracle startup.

Taking a FileSnap over NFS multiple times with the same target name can result in the 'File exists' error (2353352)

The "File exists" error occurs as a result of the caching behavior of the NFS client. Because the link operation is successful, the NFS client assumes that a file with the specified target name, such as `file2::snap:vxfs:`, was created.. As a result, the NFS client caches a file with this name.

Workaround: Remove the target file after a snapshot is created. This forces the NFS client to remove the name from the cache. For example:

```
# ln file1 file2::snap:vxfs:
# rm file2::snap:vxfs:
```


On the online cache device you should not perform the `mkfs` operation, because any subsequent `fsccache` operation panics (3643800)

When the `mkfs` operation is performed on a volume already in use for SmartIO, caching can lead to unexpected results when the subsequent `sfcache` operations are performed.

Workaround: Workaround is not available.

Deduplication can fail with error 110 (3741016)

In some cases, data deduplication fails with a message similar to the following example:

Saving	Status	Node	Type	Filesystem
00%	FAILED	node01	MANUAL	/data/fs1
2011/10/26 01:38:58 End full scan with error				

In addition, the deduplication log contains an error similar to the following example:

```
2011/10/26 01:35:09 DEDUP_ERROR AddBlock failed. Error = 110
```

These errors indicate that the deduplication process is running low on space and needs more free space to complete.

Workaround: Make more space available on the file system.

A restored volume snapshot may be inconsistent with the data in the SmartIO VxFS cache (3760219)

The data in a volume snapshot may have data that is inconsistent with the VxFS level SmartIO cache. When the volume snapshot is restored and mounted, then before using that file system you should purge the corresponding cache data. Or, disable the caching for that file system.

Workaround:

Purge the file system data from the SmartIO cache after restoring the volume snapshot.

```
# sfcache purge {mount_point|fsuuid}
```

When in-place and relocate compression rules are in the same policy file, file relocation is unpredictable (3760242)

You cannot have in-place compress/uncompress rules and relocate compress/uncompress rules in the same policy file. If they are in the same file, file relocation is unpredictable.

Workaround: Create a different policy file for each policy, and enforce the policy as per the required sequence.

The file system may hang when it has compression enabled (3331276)

In a VxFS file system that has compression enabled, a deadlock in page fault handler can lead to the file system hang.

Workaround:

There is no workaround for this issue.

Replication known issues

This section describes the replication known issues in this release of Veritas InfoScale Storage and Veritas InfoScale Enterprise.

The secondary vradmind may appear hung and the vradmind commands may fail (3940842,3944301)

The secondary vradmind daemon may appear to be hung and the vradmind commands may fail, even if the replication is ON.

This issue is observed in cases where heavy I/Os are in progress.

To address the heavy I/Os, the secondary vradmind continues to perform certain IOCTLs and hence fails to respond to the heartbeat from the primary vradmind. As a result, the connection between the two vradmind is lost. This behavior continues during the next connection attempts too, due to which the vradmind commands fail and the secondary vradmind appears hung.

Workaround:

On both the sites, modify the `/etc/vx/vras/vras_env` file to resolve the issue. Make the following edits:

- Uncomment the following commands:

```
#export VRAS_ENABLE_STATS
#VRAS_ENABLE_STATS=on
```

and

```
#export IPM_HEARTBEAT_TIMEOUT  
#IPM_HEARTBEAT_TIMEOUT=30
```

- Set the VRAS_ENABLE_STATS to OFF
- Increase the IPM_HEARTBEAT_TIMEOUT value

After you modify the files, restart the vradmind daemon on both the sites:

```
/usr/sbin/vxstart_vvr stop  
  
/usr/sbin/vxstart_vvr start
```

Data corruption may occur if you perform a rolling upgrade of InfoScale Storage or InfoScale Enterprise from 7.3.1 or earlier to 7.4 or later during replication (3951527)

A rolling upgrade to InfoScale 7.4 or later is not supported in a Volume Replicator (VVR) environment. With InfoScale 7.3.1 or earlier, you could pause replication, perform a rolling upgrade, and then safely resume replication. However, to upgrade to 7.4 or later, you must first stop any ongoing replication, perform a full upgrade of the product and the disk groups on both the sites, and then start replication.

Workaround:

To upgrade to InfoScale 7.4 or later in a VVR environment

1. Stop replication.
2. Perform a full upgrade on the primary and the secondary sites to the same InfoScale version.
3. Upgrade the disk groups on both the sites.
4. Start replication.

After the product upgrade on secondary site, replication may fail to resume with "Secondary SRL missing" error [3931763]

When you attempt to resume the replication after the product upgrade on secondary site is complete, the replication may fail to resume with a configuration error "Secondary SRL missing".

This issue occurs because even after the product upgrade is complete, the Storage Replicator Log (SRL) volume remains disassociated from the Replicated Volume Group (RVG).

During a product upgrade, the installer pauses the replication and performs several tasks that include dissociation and association of SRL volume. Due to some internal error, the installer fails to reassociate the SRL volume to the RVG. As a result, when you attempt to resume the replication from the primary site to the upgraded secondary site, it fails to start with a "Secondary SRL missing" error.

Workaround: Perform the following steps to restart the replication

1. On the upgraded site, associate the SRL to RVG

```
#vxvol -g DiskGroup_name aslog RVG_name SRL_name
```

2. Start RVG

```
# vxrvlg -g DiskGroup_name -f start RVG_name
```

3. Stop replication at primary site

```
# vradmin -g DiskGroup_name -f stoprep RVG_name
```

4. Start replication at primary site

```
# vradmin -g DiskGroup_name -a startrep RVG_name  
Secondary_hostname
```

vradmin repstatus command reports secondary host as “unreachable”(3896588)

The `vradmin repstatus` command output incorrectly reports all secondary hosts as unreachable if even one of the secondary hosts is unreachable in CVR/VVR multi-secondary environments.

Workaround: Run the following command to obtain the correct status:

```
# vradmin -g dg_name printrvlg rvlg_name
```

RVGPrimary agent operation to start replication between the original Primary and the bunker fails during failback (2036605)

The RVGPrimary agent initiated operation to start replication between the original Primary and the bunker fails during failback – when migrating back to the original Primary after disaster recovery – with the error message:

```
VxVM VVR vxrlink ERROR V-5-1-5282 Error getting information from  
remote host. Internal Error.
```

The issue applies to global clustering with a bunker configuration, where the bunker replication is configured using storage protocol. It occurs when the Primary comes

back even before the bunker disk group is imported on the bunker host to initialize the bunker replay by the RVGPrimary agent in the Secondary cluster.

Workaround:

To resolve this issue

- 1 Before failback, make sure that bunker replay is either completed or aborted.
- 2 After failback, deport and import the bunker disk group on the original Primary.
- 3 Try the start replication operation from outside of VCS control.

A snapshot volume created on the Secondary, containing a VxFS file system may not mount in read-write mode and performing a read-write mount of the VxFS file systems on the new Primary after a global clustering site failover may fail [3761497]

Issue 1:

When the `vradmin ibc` command is used to take a snapshot of a replicated data volume containing a VxFS file system on the Secondary, mounting the snapshot volume in read-write mode may fail with the following error:

```
UX:vxfs mount: ERROR: V-3-21268: /dev/vx/dsk/dg/snapshot_volume  
is corrupted. needs checking
```

This happens because the file system may not be quiesced before running the `vradmin ibc` command and therefore, the snapshot volume containing the file system may not be fully consistent.

Issue 2:

After a global clustering site failover, mounting a replicated data volume containing a VxFS file system on the new Primary site in read-write mode may fail with the following error:

```
UX:vxfs mount: ERROR: V-3-21268: /dev/vx/dsk/dg/data_volume  
is corrupted. needs checking
```

This usually happens because the file system was not quiesced on the original Primary site prior to the global clustering site failover and therefore, the file systems on the new Primary site may not be fully consistent.

Workaround: The following workarounds resolve these issues.

For issue 1, run the `fsck` command on the snapshot volume on the Secondary, to restore the consistency of the file system residing on the snapshot.

For example:

```
# fsck -F vxfs /dev/vx/dsk/dg/snapshot_volume
```

For issue 2, run the `fsck` command on the replicated data volumes on the new Primary site, to restore the consistency of the file system residing on the data volume.

For example:

```
# fsck -F vxfs /dev/vx/dsk/dg/data_volume
```

In an IPv6-only environment RVG, data volumes or SRL names cannot contain a colon (1672410, 1672417)

Issue: After upgrading VVR to an IPv6-only environment in release 6.0 or later, `vradmin` commands may not work when a colon is specified in the RVG, data volume(s) and/or SRL name. It is also possible that after upgrading VVR to an IPv6-only environment, `vradmin createpri` may dump core when provided with RVG, volume and/or SRL names containing a colon in it.

Workaround: Make sure that colons are not specified in the volume, SRL, and RVG names in the VVR configuration

vxassist layout removes the DCM (145413)

If you perform a layout that adds a column to a striped volume that has a DCM, the DCM is removed. There is no message indicating that this has happened. To replace the DCM, enter the following:

```
#vxassist -g diskgroup addlog vol logtype=dcn
```

vradmin functionality may not work after a master switch operation [2158679]

In certain situations, if you switch the master role, `vradmin` functionality may not work. The following message displays:

```
VxVM VVR vxrlink ERROR V-5-1-15861 Command is not supported for  
command shipping. Operation must be executed on master
```

Workaround:

To restore vradmind functionality after a master switch operation

- 1 Restart `vradmind` on all cluster nodes. Enter the following:

```
# /etc/init.d/vras-vradmind.sh stop
# /etc/init.d/vras-vradmind.sh start
```

- 2 Re-enter the command that failed.

Cannot relayout data volumes in an RVG from concat to striped-mirror (2129601)

This issue occurs when you try a relayout operation on a data volume which is associated to an RVG, and the target layout is a striped-mirror.

Workaround:**To relayout a data volume in an RVG from concat to striped-mirror**

- 1 Pause or stop the applications.
- 2 Wait for the RLINKs to be up to date. Enter the following:

```
# vxrlink -g diskgroup status rlink
```

- 3 Stop the affected RVG. Enter the following:

```
# vxrvrg -g diskgroup stop rvrg
```

- 4 Disassociate the volumes from the RVG. Enter the following:

```
# vxvol -g diskgroup dis vol
```

- 5 Relayout the volumes to striped-mirror. Enter the following:

```
# vxassist -g diskgroup relayout vol layout=stripe-mirror
```

- 6 Associate the data volumes to the RVG. Enter the following:

```
# vxvol -g diskgroup assoc rvrg vol
```

- 7 Start the RVG. Enter the following:

```
# vxrvrg -g diskgroup start rvrg
```

- 8 Resume or start the applications.

vradmin verifydata operation fails when replicating between versions 5.1 and 6.0 or later (2360713)

When replicating in a cross-version VVR environment consisting of hosts running Storage Foundation 5.1 and hosts running Storage Foundation 6.0 or later, the `vradmin verifydata` command fails with the following error:

```
VxVM VVR vxrsync ERROR V-5-52-2222 [from host]: VxVM in.vxrsyncd
ERROR V-5-36-2125 Server volume access error during [assign volids]
volume path: [/dev/vx/dsk/dg/snapshot_volume] reason: [this could be
because a target volume is disabled or an rlink associated with a
target volume is not detached during sync operation].
```

Workaround: There are two workarounds for this issue.

- Upgrade the hosts running Storage Foundation 5.1 to Storage Foundation 6.0 or later and re-run the `vradmin verifydata` command.
- Follow the offline verification procedure in the "Verifying the data on the Secondary" section of the *Storage Foundation and High Availability Solutions Replication Administrator's Guide*. This process requires ensuring that the secondary is up-to-date, pausing replication, and running the `vradmin syncrvg` command with the `-verify` option.

vradmin verifydata may report differences in a cross-endian environment (2834424)

When replicating between two nodes in a cross-platform environment, and performing an autosync or replication, the `vradmin verifydata` command may report differences. This is due to different endianness between the platforms. However, the file system on the secondary node will be consistent and up to date.

vradmin verifydata operation fails if the RVG contains a volume set (2808902)

In a VVR environment, the `vradmin verifydata` command fails with the following error if the replicated volume group (RVG) contains any volume set:

```
Message from Primary:
VxVM VVR vxrsync ERROR V-5-52-2009 Could not open device
/dev/vx/dsk/vvrdg/<volname> due to: stat of raw character volume path
failed
```


Bunker replay does not occur with volume sets (3329970)

There are issues with bunker replication using Volume Replicator (VVR) with volume sets. Do not upgrade to Storage Foundation HA 7.4 if you have configured or plan to configure bunker replication using VVR with volume sets.

Workaround:

Contact Veritas Technical Support for a patch that enables you to use this configuration.

SmartIO does not support write-back caching mode for volumes configured for replication by Volume Replicator (3313920)

SmartIO does not support write-back caching mode for volumes that are configured for replication by Volume Replicator (VVR).

Workaround:

If you have configured volumes for replication by VVR, do not enable write-back caching

During moderate to heavy I/O, the vradmin verifydata command may falsely report differences in data (3270067)

While an application is online at the Volume Replicator primary site, the `vradmin verifydata` command may fail. The command output shows the differences between the source data volume and the target data volume.

Workaround:

The reason for this error is that the cache object that is used for the verification might be under allocated. You might need to allocate more space for the shared cache object. For guidelines on shared cache object allocation, see the section "Creating a shared cache object" in the *Storage Foundation Administrator's Guide*.

The vradmin repstatus command does not show that the SmartSync feature is running [3343141]

In a Volume Replicator (VVR) environment, after you start the initial synchronization with the `vradmin -a startrep` command with file system mounted on the primary data volumes, the `vradmin repstatus` command does not show that the SmartSync feature is running. This is an only issue with the output of the `vradmin repstatus` command.

Workaround:

To confirm that SmartSync is running, enter:

```
vxrlink status rlink
```

While vradmin commands are running, vradmind may temporarily lose heartbeats (3347656, 3724338)

This issue may occasionally occur when you use `vradmin` commands to administer Volume Replicator (VVR). While the `vradmin` commands run, `vradmind` may temporarily lose heartbeats, and the commands terminate with the following error message:

```
VxVM VVR vradmin ERROR V-5-52-803 Lost connection to host host;  
terminating command execution.
```

Workaround:

To resolve this issue:

- 1 Depending on the application I/O workload and the network environment, uncomment and increase the value of the `IPM_HEARTBEAT_TIMEOUT` variable in the `/etc/vx/vras/vras_env` on all the hosts of the replicated data set (RDS) to a higher value. The following example increases the timeout value to 120 seconds:

```
export IPM_HEARTBEAT_TIMEOUT  
IPM_HEARTBEAT_TIMEOUT=120
```

- 2 Restart `vradmind` on all the hosts of the RDS to put the new `IPM_HEARTBEAT_TIMEOUT` value into affect. Enter the following on all the hosts of the RDS:

```
# /etc/init.d/vras-vradmind.sh stop  
# /etc/init.d/vras-vradmind.sh start
```

Write I/Os on the primary logowner may take a long time to complete (2622536)

Under a heavy I/O load, write I/Os on the Volume Replicator (VVR) primary logowner take a long time to complete.

Workaround:

There is no workaround for this issue.

DCM logs on a disassociated layered data volume results in configuration changes or CVM node reconfiguration issues (3582509)

If you have configured layered data volumes under an RVG that has DCM protection enabled and at a later point disassociate the data volume from the RVG, you must manually remove the DCM logs from the volume. Leaving DCM logs on a layered data volume after it has been disassociated from the RVG, may result configuration changes, or the CVM node reconfiguration to not work properly.

Workaround:

If the disk group has a layered volume, remove DCM logs after disassociating the volumes from the RVG.

After performing a CVM master switch on the secondary node, both rlinks detach (3642855)

If the VVR logowner (master) node on the secondary site goes down during initial synchronization, then during the RVG recovery (initiated on any secondary side node as a result of node crash), the replication links detach with the following error:

```
WARNING: VxVM VVR vxio V-5-0-187 Incorrect magic number or unexpected
upid (1) rvg rvg1
WARNING: VxVM VVR vxio V-5-0-287 rvg rvg1, SRL srl1: Inconsistent log
- detaching all rlinks.
```

Workaround:

Restart replication using the autosync operation.

The RVGPrimary agent may fail to bring the application service group online on the new Primary site because of a previous primary-elect operation not being run or not completing successfully (3761555, 2043831)

In a primary-elect configuration, the RVGPrimary agent may fail to bring the application service groups online on the new Primary site, due to the existence of previously-created instant snapshots. This may happen if you do not run the `ElectPrimary` command to elect the new Primary or if the previous `ElectPrimary` command did not complete successfully.

Workaround: Destroy the instant snapshots manually using the `vxrvrg -g dg -P snap_prefix snapdestroy rvg` command. Clear the application service group and bring it back online manually.

A snapshot volume created on the Secondary, containing a VxFS file system may not mount in read-write mode and performing a read-write mount of the VxFS file systems on the new Primary after a global clustering site failover may fail (1558257)

Issue 1:

When the `vradmin ibc` command is used to take a snapshot of a replicated data volume containing a VxFS file system on the Secondary, mounting the snapshot volume in read-write mode may fail with the following error:

```
UX:vxfs mount: ERROR: V-3-21268: /dev/vx/dsk/dg/snapshot_volume  
is corrupted. needs checking
```

This happens because the file system may not be quiesced before running the `vradmin ibc` command and therefore, the snapshot volume containing the file system may not be fully consistent.

Issue 2:

After a global clustering site failover, mounting a replicated data volume containing a VxFS file system on the new Primary site in read-write mode may fail with the following error:

```
UX:vxfs mount: ERROR: V-3-21268: /dev/vx/dsk/dg/data_volume  
is corrupted. needs checking
```

This usually happens because the file system was not quiesced on the original Primary site prior to the global clustering site failover and therefore, the file systems on the new Primary site may not be fully consistent.

Workaround: The following workarounds resolve these issues.

For issue 1, run the `fsck` command on the snapshot volume on the Secondary, to restore the consistency of the file system residing on the snapshot.

For example:

```
# fsck -F vxfs /dev/vx/dsk/dg/snapshot_volume
```

For issue 2, run the `fsck` command on the replicated data volumes on the new Primary site, to restore the consistency of the file system residing on the data volume.

For example:

```
# fsck -F vxfs /dev/vx/dsk/dg/data_volume
```

DCM plex becomes inaccessible and goes into DISABLED(SPARSE) state in case of node failure. (3931775)

In FSS environment, when a CVR is configured on primary and secondary site, DCM log plex is created by default on each volume. This log gets created with mirrors across local disks of single node, instead of mirrors across multiple nodes due to `vxassist` command limitation. This limitation restricts mirror, stripe separation, or confinement to allocate the log plexes to associate with the volume. Hence, DCM plex becomes inaccessible and goes into DISABLED (SPARSE) in case of node failure.

Workaround

1. While creating volumes, create and associate DCM logs manually from different nodes using the following command:

```
# vxassist -g <disk_group_name> addlog <volume_name> logtype=dcm  
    <local_disks_across_different_nodes>
```

2. Create RVG on the data volume.

Cluster Server known issues

This section describes the known issues in this release of Cluster Server (VCS). These known issues apply to the following products:

- Veritas InfoScale Availability
- Veritas InfoScale Enterprise

Operational issues for VCS

This section describes the Operational known issues for VCS.

Using TLS v1.0 for VCS may make the configuration susceptible to a network attack

Using TLS v1.0 for VCS may make the configuration susceptible to a POODLE type of attack (3944529)

TLS is capable of using a multitude of ciphers to create the public and private key pairs. For example, TLS v1.0 uses either the RC4 stream cipher, or a block cipher in CBC mode. RC4 is known to have biases and the block cipher in CBC mode is vulnerable to the POODLE attack. TLS v1.0, if configured to use the same cipher suites as SSLv3, includes a means by which a TLS implementation can downgrade the connection to SSL v3.0, thus weakening security. A POODLE type of attack

could also be launched directly at TLS without negotiating a downgrade. For more information about a POODLE attack, visit:

<https://blog.qualys.com/ssllabs/2014/12/08/poodle-bites-tls>

Workaround:

By default, InfoScale uses TLS v1.0 for public communication. However, you may want to disable TLS v1.0 because it is susceptible to a POODLE type of attack. To disable TLS v1.0, set `AllowTLSv1=0` in the `VRTSatlocal.conf` file.

The `hastop -all` command on VCS cluster node with AlternatelIO resource and StorageSG having service groups may leave the node in LEAVING state

On a VCS cluster node with AlternatelIO resource configured and StorageSG attribute contain service groups with Zpool, VxVM or CVMVoIDG resources, `hastop -local` or `hastop -all` commands may leave the node in "LEAVING" state.

This issue is caused by lack of dependency between service group containing LDom resource and service groups containing storage resources exported to logical domain in alternate I/O domain scenarios. In this scenario VCS may attempt to stop the storage service groups before stopping logical domain which is using the resources.

Workaround: Stop the LDom service group before issuing `hastop -local` or `hastop -all` commands.

Missing characters in system messages [2334245]

You may see missing characters, especially in long system messages in response to certain commands.

Workaround: No workaround.

CP server does not allow adding and removing HTTPS virtual IP or ports when it is running [3322154]

CP server does not support adding and removing HTTPS virtual IPs or ports while the CP server is running.

Workaround: No workaround. If you want to add a new virtual IP for HTTPS, you must follow the entire manual procedure for generating HTTPS certificate for the CP server (`server.crt`), as documented in the *Cluster Server Configuration and Upgrade Guide*.

CP server does not support IPv6 communication with HTTPS protocol [3209475]

CP server does not support IPv6 communication when using the HTTPS protocol. This implies that in VCS, CP servers listening on HTTPS can only use IPv4. As a result, VCS fencing clients can also use only IPv4.

Workaround: No workaround.

System encounters multiple VCS resource timeouts and agent core dumps [3424429]

The system encounters multiple VCS resource timeouts and agent core dumps without any specific reason.

The issue pertains to a hardware errata with the Intel Xeon CPUs where a processor can go into a low power sleep mode, but takes a long time to wake up. This can cause erratic scheduling behavior, leading to unexpected delays, expired timers, or occasional freezes. For more information, see the Oracle document:

<https://support.oracle.com/epmos/faces/BugDisplay?id=15659645>

Workaround: Add the following lines to the `/etc/system` file and reboot the system:

```
set idle_cpu_prefer_mwait = 0
set idle_cpu_no_deep_c = 1
```

Some VCS components do not work on the systems where a firewall is configured to block TCP traffic [3545338]

The following issues may occur if you install and configure VCS on systems where a firewall is installed:

- If you set up Disaster Recovery using the Global Cluster Option (GCO), the status of the remote cluster (cluster at the secondary site) shows as "initing".
- If you configure fencing to use CP server, fencing client fails to register with the CP server.
- Setting up trust relationships between servers fails.

Workaround:

- Ensure that the required ports and services are not blocked by the firewall. Refer to the *Cluster Server Configuration and Upgrade Guide* for the list of ports and services used by VCS.
- Configure the firewall policy such that the TCP ports required by VCS are not blocked. Refer to your respective firewall or OS vendor documents for the required configuration.

Issues related to the VCS engine

This section describes the known issues about the VCS engine.

Extremely high CPU utilization may cause HAD to fail to heartbeat to GAB [1744854]

When CPU utilization is very close to 100%, HAD may fail to heartbeat to GAB.

Missing host names in engine_A.log file (1919953)

The GUI does not read the `engine_A.log` file. It reads the `engine_A.ldf` file, gets the message id from it, and then queries for the message from the `bmc` file of the appropriate locale (Japanese or English). The `bmc` file does not have system names present and so they are read as missing.

The `hacf -cmdtoctf` command generates a broken `main.cf` file [1919951]

The `hacf -cmdtoctf` command used with the `-dest` option removes the include statements from the types files.

Workaround: Add include statements in the `main.cf` files that are generated using the `hacf -cmdtoctf` command.

Character corruption observed when executing the `uuidconfig.pl -clus -display -use_llthost` command [2350517]

If password-less `ssh/rsh` is not set, the use of `uuidconfig.pl` command in non-English locale may print garbled characters instead of a non-English string representing the Password prompt.

Workaround: No workaround.

Trigger does not get executed when there is more than one leading or trailing slash in the `triggerpath` [2368061]

The path specified in `TriggerPath` attribute must not contain more than one leading or trailing `'/'` character.

Workaround: Remove the extra leading or trailing `'/'` characters from the path.

Service group is not auto started on the node having incorrect value of EngineRestarted [2653688]

When HAD is restarted by `hashadow` process, the value of `EngineRestarted` attribute is temporarily set to 1 till all service groups are probed. Once all service groups are probed, the value is reset. If HAD on another node is started at roughly the same time, then it is possible that it does not reset the value of `EngineRestarted` attribute. Therefore, service group is not auto started on the new node due to mismatch in the value of `EngineRestarted` attribute.

Workaround: Restart VCS on the node where `EngineRestarted` is set to 1.

Group is not brought online if top level resource is disabled [2486476]

If the top level resource which does not have any parent dependency is disabled then the other resources do not come online and the following message is displayed:

```
VCS NOTICE V-16-1-50036 There are no enabled
resources in the group cvm to online
```

Workaround: Online the child resources of the topmost resource which is disabled.

NFS resource goes offline unexpectedly and reports errors when restarted [2490331]

VCS does not perform resource operations, such that if an agent process is restarted multiple times by HAD, only one of the agent process is valid and the remaining processes get aborted, without exiting or being stopped externally. Even though the agent process is running, HAD does not recognize it and hence does not perform any resource operations.

Workaround: Terminate the agent process.

Parent group does not come online on a node where child group is online [2489053]

This happens if the `AutostartList` of parent group does not contain the node entry where the child group is online.

Workaround: Bring the parent group online by specifying the name of the system then use the `hargp -online [parent group] -any` command to bring the parent group online.

Cannot modify temp attribute when VCS is in LEAVING state [2407850]

An `ha` command to modify a temp attribute is rejected if the local node is in a LEAVING state.

Workaround: Execute the command from another node or make the configuration read-write enabled.

Oracle service group faults on secondary site during failover in a disaster recovery scenario [2653704]

Oracle service group fails to go online in the DR site when disaster strikes the primary site. This happens if the AutoFailover attribute on the Service Group is set to 1 and when the corresponding service group's FireDrill is online in the DR site. Firedrill Service group may remain ONLINE on the DR site.

Workaround: If the service group containing the Oracle (or any database) resource faults after attempting automatic DR failover while FireDrill is online in the DR site, manually offline the FireDrill Service Group. Subsequently, attempt the online of the Oracle Service Group in the DR site.

Service group may fail to come online after a flush and a force flush operation [2616779]

A service group may fail to come online after flush and force flush operations are executed on a service group where offline operation was not successful.

Workaround: If the offline operation is not successful then use the force flush commands instead of the normal flush operation. If a normal flush operation is already executed then to start the service group use `-any` option.

Elevated TargetCount prevents the online of a service group with `hagrp -online -sys` command [2871892]

When you initiate an offline of a service group and before the offline is complete, if you initiate a forced flush, the offline of the service group which was initiated earlier is treated as a fault. As start bits of the resources are already cleared, service group goes to OFFLINE|FAULTED state but TargetCount remains elevated.

Workaround: No workaround.

Auto failover does not happen in case of two successive primary and secondary cluster failures [2858187]

In case of three clusters (clus1, clus2, clus3) in a GCO with steward not configured, if clus1 loses connection with clus2, it sends the inquiry to clus3 to check the state of clus2 one of the following condition persists:

1. If it is able to confirm that clus2 is down, it will mark clus2 as FAULTED.
2. If it is not able to send the inquiry to clus3, it will assume that a network disconnect might have happened and mark clus2 as UNKNOWN

In second case, automatic failover does not take place even if the ClusterFailoverPolicy is set to Auto. You need to manually failover the global service groups.

Workaround: Configure steward at a geographically distinct location from the clusters to which the above stated condition is applicable.

GCO clusters remain in INIT state [2848006]

GCO clusters remain in INIT state after configuring GCO due to :

- Trust between two clusters is not properly set if clusters are secure.
- Firewall is not correctly configured to allow WAC port (14155).

Workaround: Make sure that above two conditions are rectified. Refer to *Cluster Server Administrator's Guide* for information on setting up Trust relationships between two clusters.

The `ha` commands may fail for non-root user if cluster is secure [2847998]

The `ha` commands fail to work for one of the following reasons:

- If you first use a non-root user without a home directory and then create a home directory for the same user.
- If you configure security on a cluster and then un-configure and reconfigure it.

Workaround

- 1 Delete `/var/VRTSat/profile/<user_name>`,
- 2 Delete `/home/user_name/.VRTSat`.
- 3 Delete `/var/VRTSat_lhc/<cred_file>` file which same non-root user owns.
- 4 Run `ha` command with same non-root user (this will pass).

Startup trust failure messages in system logs [2721512]

If you configure a cluster with security enabled, there might be some messages logged in system message logs related to Veritas authentication. These messages can be ignored and have no effect on functionality.

Workaround: No workaround.

Running `-delete -keys` for any scalar attribute causes core dump [3065357]

Running `-delete -keys` for any scalar attribute is not a valid operation and must not be used. However, any accidental or deliberate use of this command may cause engine to core dump.

Workaround: No workaround.

Veritas InfoScale enters into `admin_wait` state when Cluster Statistics is enabled with load and capacity defined [3199210]

Veritas InfoScale enters into `admin_wait` state when started locally if:

1. Statistics attribute value is set to Enabled, which is its default value.
2. Group Load and System Capacity values are defined in units in `main.cf`.

Workaround:

1. Stop Veritas InfoScale on all nodes in the cluster.
2. Perform any one of the following steps:
 - Edit the `main.cf` on one of the nodes in the cluster and set the Statistics attribute to Disabled or MeterHostOnly.
 - Remove the Group Load and System Capacity values from the `main.cf`.
3. Run `hacf -verify` on the node to verify that the configuration is valid.
4. Start Veritas InfoScale on the node and then on the rest of the nodes in the cluster.

Agent reports incorrect state if VCS is not set to start automatically and `utmp` file is empty before VCS is started [3326504]

If you have not configured VCS to start automatically after a reboot and have emptied the `utmp` file before starting VCS manually with the `hastart` command, some agents might report an incorrect state.

The `utmp` file (file name may differ on different operating systems) is used to maintain a record of the restarts done for a particular machine. The `checkboot` utility used by `hastart` command uses the functions provided by the OS which in turn use the `utmp` file to find if a system has been restarted so that the temporary files for various agents can be deleted before agent startup. If OS functions do not return correct value, High Availability Daemon (HAD) starts without deleting the stale agent files. This might result in some agents reporting incorrect state.

Workaround: If a user wishes to delete the `utmp` file this should be done only when VCS is already running or the customer should delete the temporary files in `/var/VRTSvcs/lock/volatile/` manually before starting VCS.

VCS crashes if feature tracking file is corrupt [3603291]

VCS keeps a track of some specific features used in the VCS cluster. For example, if a Global service group is brought online then the feature is logged in a specific feature tracking file. If the file however is corrupt, then VCS may dump core when attempting to write data to the file.

Workaround: Delete the corrupt feature tracking file (`/var/vx/vftrk/vcs`) and restart VCS.

RemoteGroup agent and non-root users may fail to authenticate after a secure upgrade [3649457]

On upgrading a secure cluster to 6.2 or later release, the following issues may occur with unable to open a secure connection error:

- The RemoteGroup agent may fail to authenticate with remote cluster.
- Non-root users may fail to authenticate.

Workaround

- 1 Set `LC_ALL=C` on all nodes before upgrade or perform the following steps after the upgrade on all nodes of the cluster:
 - Stop HAD.
 - Set `LC_ALL=C`.
 - Start HAD using `hastart`.
- 2 Reset `LC_ALL` attribute to the previous value once the non-root users are validated.

If you disable security before upgrading VCS to version 7.0.1 or later on secured clusters, the security certificates will not be upgraded to 2048 bit SHA2 [3812313]

The default security certificates installed with VCS 7.0 and the earlier versions are 1024 bit SHA1. If you disable security before upgrading VCS to version 7.0.1 or later on secured clusters, the installer will upgrade VCS but will not upgrade the security certificates. Therefore, merely enabling security after the VCS upgrade to 7.0.1 or later does not upgrade the security to 2048 bit SHA2 certificates.

Workaround:

When you upgrade VCS to version 7.0.1 or later releases, run the `installer -security` command and select the `reconfigure` option to upgrade the security certificates to 2048 bit SHA2.

Note: On Solaris 11 x64, you will not hit this issue if you upgrade from VCS 7.0 to 7.0.1, because VCS 7.0 on Solaris 11 x64 has 2048 bit SHA2 certificates.

Java console and CLI do not allow adding VCS user names starting with ‘_’ character (3870470)

When a user adds a new user name, VCS checks if first character of the user name is part of the set of allowed characters. The ‘_’ character is not part of the permitted set. So the user name starting with ‘_’ is considered invalid.

Workaround: Use another user name which starts with a character permitted by VCS.

Issues related to the bundled agents

This section describes the known issues of the bundled agents.

Entry points that run inside a zone are not cancelled cleanly [1179694]

Cancelling entry points results in the cancellation of only the `zlogin` process. The script entry points that run inside a zone are forked off using the `zlogin` command. However, the `zlogin` command forks off an `sh` command, which runs in the context of the Solaris zone. This shell process and its family do not inherit the group id of the `zlogin` process, and instead get a new group id. Thus, it is difficult for the agent framework to trace the children or grand-children of the shell process, which translates to the cancellation of only the `zlogin` process.

Workaround: Oracle must provide an API or a mechanism to kill all the children of the `zlogin` process that was started to run the entry point script in the local-zone.

Solaris mount agent fails to mount Linux NFS exported directory

The Solaris mount agent mounts the mount directories. At this point, if it tries to mount a Linux NFS exported directory, the mount fails showing the following error:

```
nfs mount: mount: <MountPoint>: Not owner
```

This is due to system NFS default version mismatch between Solaris and Linux.

The workaround for this is to configure `MountOpt` attribute in mount resource and set `vers=3` for it.

Example

```
root@north $ mount -F nfs south:/test /logo/
nfs mount: mount: /logo: Not owner
root@north $
Mount nfsmount (
    MountPoint = "/logo"
    BlockDevice = "south:/test"
    FSType = nfs
    MountOpt = "vers=3"
)
```

The zpool command runs into a loop if all storage paths from a node are disabled

The Solaris Zpool agent runs `zpool` commands to import and export zpools. If all paths to the storage are disabled, the `zpool` command does not respond. Instead, the `zpool` export command goes into a loop and attempts to export the zpool. This continues till the storage paths are restored and `zpool` is cleared. As a result, the offline and clean procedures of Zpool Agent fail and the service group cannot fail over to the other node.

Workaround: You must restore the storage paths and run the `zpool` clear command for all the pending commands to succeed. This will cause the service group to fail over to another node.

Zone remains stuck in down state if tried to halt with file system mounted from global zone [2326105]

If zone halts without unmounting the file system, the zone goes to down state and does not halt with the `zoneadm` commands.

Workaround: Unmount the file system manually from global zone and then halt the zone. For VxFS, use following commands to unmount the file system from global zone.

To unmount when VxFSMountLock is 1

```
umount -o mntunlock=VCS <zone root path>/<Mount Point>
```

To forcefully unmount when VxFSMountLock is 1:

```
# umount -f -o mntunlock=VCS <zone root path>/<Mount Point>
```

To unmount when VxFSMountLock is 0:

```
# umount <zone root path>/<Mount Point>
```

To forcefully unmount when VxFSMountLock is 0:

```
# umount -f <zone root path>/<Mount Point>
```

To halt the zone, use following command:

```
# zoneadm -z <zone_name> halt
```

Process and ProcessOnOnly agent rejects attribute values with white spaces [2303513]

Process and ProcessOnOnly agent does not accept Arguments attribute values that are separated by multiple whitespaces. The Arguments attribute specifies the set of arguments for a process. If a script controls the process, the script is passed as an argument. You must separate multiple arguments by using a single whitespace. A string cannot accommodate more than one space between arguments, or allow leading or trailing whitespace characters. This attribute must not exceed 80 characters.

Workaround: You should use only single whitespace to separate the argument attribute values. Make sure you avoid multiple whitespaces between the argument attribute values or trailing whitespace characters.

The zpool commands hang and remain in memory till reboot if storage connectivity is lost [2368017]

If the FailMode attribute of `zpool` is set to continue or wait and the underlying storage is not available, the `zpool` commands hang and remain in memory until the next reboot.

This happens when storage connectivity to the disk is lost, the `zpool` commands hang and they cannot be stopped or killed. The `zpool` commands run by the monitor entry point remains in the memory.

Workaround: There is no recommended workaround for this issue.

Offline of zone resource may fail if `zoneadm` is invoked simultaneously [2353541]

Offline of zone EP uses `zoneadm` command to offline a zone. Therefore, if `zoneadm` is invoked simultaneously for multiple zones, the command may fail. This is due to Oracle bug 6757506 that causes a race condition between multiple instances of `zoneadm` command and displays the following message:

```
zoneadm: failed to get zone name: Invalid argument
```

Workaround: No workaround.

Password changed while using `hazonesetup` script does not apply to all zones [2332349]

If you use the same user name for multiple zones, updating password for one zone does not updated the password of other zones.

Workaround: While updating password for VCS user which is used for multiple zones, update password for all the zones.

RemoteGroup agent does not failover in case of network cable pull [2588807]

A RemoteGroup resource with ControlMode set to OnOff may not fail over to another node in the cluster in case of network cable pull. The state of the RemoteGroup resource becomes UNKNOWN if it is unable to connect to a remote cluster.

Workaround:

- Connect to the remote cluster and try taking offline the RemoteGroup resource.
- If connection to the remote cluster is not possible and you want to bring down the local service group, change the ControlMode option of the RemoteGroup resource to MonitorOnly. Then try taking offline the RemoteGroup resource.

Once the resource is offline, change the ControlMode option of the resource to OnOff.

CoordPoint agent remains in faulted state [2852872]

The CoordPoint agent remains in faulted state because it detects `rfsm` to be in replaying state.

Workaround: After HAD has stopped, reconfigure fencing.

Prevention of Concurrency Violation (PCV) is not supported for applications running in a container [2536037]

For an application running in a container, VCS uses a similar functionality as if that resource is not registered to IMF. Hence, there is no IMF control to take a resource offline. When the same resource goes online on multiple nodes, agent detects and reports to engine. Engine uses the offline monitor to take the resource offline. Hence, even though there is a time lag before the detection of the same resource coming online on multiple nodes at the same time, VCS takes the resource offline.

PCV does not function for an application running inside a local Zone on Solaris

Workaround: No workaround.

Share resource goes offline unexpectedly causing service group failover [1939398]

Share resource goes offline unexpectedly and causes a failover when NFSRestart resource goes offline and UseSMF attribute is set to 1 (one).

When NFSRestart resource goes offline, NFS daemons are stopped. When UseSMF attribute is set to 1, the exported file systems become unavailable, hence Share resource unexpectedly goes offline.

Workaround: Set the value of ToleranceLimit of Share resource to a value more than 1.

Mount agent does not support all scenarios of loopback mounts

For a mount point under VCS control, you can create loop back mounts for the mount point. For example, mount point `/mntpt` is mounted on `/a` as loop back mount and `/a` is mounted on `/b` as loop back mount, then offline and online of the mount resource fails.

Workaround: Mount the mount point `/mntpt` on `/b` as loop back mount.

Invalid Netmask value may display code errors [2583313]

If you specify invalid Netmask value for the IP resource attribute, you may see the code errors similar to the following when you try to online the resource.

```
=====
Illegal hexadecimal digit 'x' ignored at
/opt/VRTSperl/lib/site_perl/5.12.2/Net/Netmask.pm line 78.
ifconfig: <Netmask_value>: bad address
=====
```

Workaround: Make sure you specify a valid Netmask value.

Zone root configured on ZFS with ForceAttach attribute enabled causes zone boot failure (2695415)

On Solaris 11 system, attaching zone with `-F` option may result in zone boot failure if zone root is configured on ZFS.

Workaround: Change the ForceAttach attribute of Zone resource from 1 to 0. With this configuration, you are recommended to keep the default value of DetachZonePath as 1.

Error message is seen for Apache resource when zone is in transient state [2703707]

If the Apache resource is probed when the zone is getting started, the following error message is logged:

```
Argument "VCS ERROR V-16-1-10600 Cannot connect to VCS engine\n"
isn't numeric in numeric ge (>=) at /opt/VRTSvcs/bin/Apache/Apache.pm
line 452.
VCS ERROR V-16-1-10600 Cannot connect to VCS engine
LogInt(halog call failed):TAG:E:20314 <Apache::ArgsValid> SecondLevel
MonitorTimeOut must be less than MonitorTimeOut.
```

Workaround: You can ignore this message. When the zone is started completely, the `halog` command does not fail and Apache agent monitor runs successfully.

Monitor falsely reports NIC resource as offline when zone is shutting down (2683680)

If a NIC resource is configured for an Exclusive IP zone, the NIC resource is monitored inside the zone when the zone is functional. If the NIC monitor program is invoked when the zone is shutting down, the monitor program may falsely report

the NIC resource as offline. This may happen if some of the networking services are offline but the zone is not completely shut down. Such reports can be avoided if you override and set the `ToleranceLimit` value to a non-zero value.

Workaround: When a NIC resource is configured for an Exclusive IP zone, you are recommended to set the `ToleranceLimit` attribute to a non-zero value.

Calculate the `ToleranceLimit` value as follows:

Time taken by a zone to completely shut down must be less than or equal to NIC resource's `MonitorInterval` value + (`MonitorInterval` value x `ToleranceLimit` value).

For example, if a zone take 90 seconds to shut down and the `MonitorInterval` for NIC agent is set to 60 seconds (default value), set the `ToleranceLimit` value to 1.

Apache resource does not come online if the directory containing Apache pid file gets deleted when a node or zone restarts (2680661)

The directory in which Apache http server creates `PidFile` may get deleted when a node or zone restarts. Typically the `PidFile` is located at `/var/run/apache2/httpd.pid`. When the zone reboots, the `/var/run/apache2` directory may get removed and hence the http server startup may fail.

Workaround: Make sure that Apache http server writes the `PidFile` to an accessible location. You can update the `PidFile` location in the Apache http configuration file (For example: `/etc/apache2/httpd.conf`).

Online of LDom resource may fail due to incompatibility of LDom configuration file with host OVM version (2814991)

If you have a cluster running LDom with different OVM versions on the hosts, then the LDom configuration file generated on one host may display error messages when it is imported on the other host with a different OVM version. Thus, the online of LDom resource may also fail.

For example, if you have a cluster running LDom with OVM versions 2.2 on one and OVM 2.1 on the other node, the using XML configuration generated on the host with OVM 2.2 may display errors when the configuration is imported on the host with OVM 2.1. Thus, the online of LDom resource fails.

The following error message is displayed:

```
ldm add-domain failed with error Failed to add device
/ldom1/ldom1 as ld1_disk1@primary-vds0 because this device
is already exported on LDom primary. Volume ld1_disk1
already exists in vds primary-vds0.
```

Workaround: If the CfgFile attribute is specified, ensure that the XML configuration generated is compatible with the OVM version installed on the nodes.

Online of IP or IPMultiNICB resource may fail if its IP address specified does not fit within the values specified in the allowed-address property (2729505)

While configuring an IP or IPMultiNICB resource to be run in a zone, if the IP address specified for the resource does not match the values specified in the **allowed-address** property of the zone configuration, then the online of IP resource may fail. This behavior is seen only on Solaris 11 platform.

Workaround: Ensure that the IP address is added to **allowed-address** property of the zone configuration.

Application resource running in a container with PidFiles attribute reports offline on upgrade to VCS 6.0 or later [2850927]

Application resource configured to run in a container configured with PidFiles attribute reports state as offline after upgrade to Veritas InfoScale 6.0 or later versions.

When you upgrade Veritas InfoScale from lower versions to 6.0 or later, if application resources are configured to run in a container with monitoring method set to PidFiles, then upgrade may cause the state of the resources to be reported as offline. This is due to changes introduced in the Application agent where if the resource is configured to run in a container and has PidFiles configured for monitoring the resource then the value expected for this attribute is the pathname of the PID file relative to the zone root.

In releases prior to Veritas InfoScale 7.4, the value expected for the attribute was the pathname of the PID file including the zone root.

For example, a configuration extract of an application resource configured in Veritas InfoScale 5.0MP3 to run in a container would appear as follows:

```
Application apptest (
  User = root
  StartProgram = "/ApplicationTest/app_test_start"
  StopProgram = "/ApplicationTest/app_test_stop"
  PidFiles = {
    "/zones/testzone/root/var/tmp/apptest.pid" }
  ContainerName = testzone
)
```

Whereas, the same resource if configured in Veritas InfoScale 6.0 and later releases would be configured as follows:

```
Application apptest (
  User = root
  StartProgram = "/ApplicationTest/app_test_start"
  StopProgram = "/ApplicationTest/app_test_stop"
  PidFiles = {
    "/var/tmp/apptest.pid" }
)
```

Note: The container information is set at the service group level.

Workaround: Modify the PidFiles pathname to be relative to the zone root as shown in the latter part of the example.

```
# hares -modify apptest PidFiles /var/tmp/apptest.pid
```

NIC resource may fault during group offline or failover on Solaris 11 [2754172]

When NIC resource is configured with exclusive IP zone, NIC resource may fault during group offline or failover. This issue is observed as zone takes long time in shutdown on Solaris 11. If NIC monitor is invoked during this window, NIC agent may treat this as fault.

Workaround: Increase ToleranceLimit for NIC resource when it is configured for exclusive IP zone.

NFS client reports error when server is brought down using shutdown command [2872741]

On Solaris 11, when the VCS cluster node having the NFS share service group is brought down using `shutdown` command, NFS clients may report "Stale NFS file handle" error. During shutdown, the SMF service `svc:/network/shares un-shares` all the shared paths before taking down the virtual IP. Thus, the NFS clients accessing this path get stale file handle error.

Workaround: Before you shutdown the VCS cluster node, disable the `svc:/network/shares` SMF service, so that only VCS controls the un-sharing of the shared paths during the shutdown operation.

NFS client reports I/O error because of network split brain [3257399]

When network split brain occurs, the failing node may take some time to panic. As a result, the service group on the failover node may fail to come online as some of the resources (such as IP resource) are still online on the failing node. The disk group on the failing node may also get disabled but IP resource on the same node continues to be online.

Workaround: Configure the preonline trigger for the service groups containing DiskGroup resource with reservation on each system in the service group:

- 1 Copy the preonline_ipc trigger from
`/opt/VRTSvcs/bin/sample_triggers/VRTSvcs to`
`/opt/VRTSvcs/bin/triggers/preonline/ as T0preonline_ipc:`

`# cp /opt/VRTSvcs/bin/sample_triggers/VRTSvcs/preonline_ipc`
`/opt/VRTSvcs/bin/triggers/preonline/T0preonline_ipc`

- 2 Enable the preonline trigger for the service group.

```
# hagr -modify <group_name> TriggersEnabled  
PREONLINE -sys <node_name>
```

Mount resource does not support spaces in the MountPoint and BlockDevice attribute values [3335304]

Mount resource does not handle intermediate spaces in the configured MountPoint or BlockDevice attribute values.

Workaround: No workaround.

IP Agent fails to detect the online state for the resource in an exclusive-IP zone [3592683]

IP Agent does not detect the online state for the resource inside an exclusive-IP zone monitoring an IPv6 address if the link-local address is down.

Workaround: Bring the link-local address of the device up for the IP agent to detect the IPv6 address state properly.

SFCache Agent fails to enable caching if cache area is offline [3644424]

SFCache agent cannot enable caching if cache area associate with this particular object is in offline state. User need to manually online the cache area to make sure that caching can be enabled/disabled.

Workaround: Online the cache area using `sfcache` command

```
# sfcache online <cache_area_name>
```

RemoteGroup agent may stop working on upgrading the remote cluster in secure mode [3648886]

RemoteGroup agent may report the resource state as UNKNOWN if the remote cluster is upgraded to VCS 6.2 or later in secure mode.

Workaround: Restart the RemoteGroup agent.

(Solaris 11 x64) Application does not come online after the ESX server crashes or is isolated [3838654]

In a VCS cluster, if an ESX server on which the virtual machine is online with applications, crashes or isolates, then the protected online application may not failover to another virtual machine in the cluster. The server failure can be due to a delay in VMware API calls to the isolated ESX server or vCenter server.

Workaround: Increase the OnlineWaitLimit value of the VMwareDisks resource type from 300 to 500.

(Solaris 11 x64) Application may not failover when a cable is pulled off from the ESX host [3842833]

In case a storage cable is pulled off from the ESX host, the applications running inside its virtual machines may not failover to another virtual machine. As VMware vSphere APIs do not allow disk detach operation in case of a storage cable pull scenario, the VMwareDisks resource may fail to go offline.

Workaround: For the affected virtual machine, manually detach the disk from the vSphere console.

(Solaris 11 x64) Disk may not be visible on VM even after the VMwareDisks resource is online [3838644]

The disks that are attached during VMwareDisks resource online operation may not be visible to the VM user through OS commands. Due to Solaris operating

system behavior, the hot plugged disks may not be visible immediately through the commands and hence user is unable to locate those disks.

Workaround: Run the command `devfsadm -Cv` on the virtual machine to rescan devices.

(Solaris 11 x64) Virtual machine may hang when the VMwareDisks resource is trying to come online [3849480]

If VMwareDisks resource attempts to attach a disk to a virtual machine that is already attached to some other virtual machine, then the attach operation may hang, causing the virtual machine to hang. As a result, the VM may miss LLT heartbeats, and may get isolated in the network.

Workaround: Ensure that the disks used by a virtual machine are not attached or used by any other virtual machine outside the VCS cluster.

On Oracle Solaris x86 SRU 11.3.15.4.0, NetBios agent fails to come online (3915231)

On Oracle Solaris x86 SRU 11.3.15.4.0, NetBios agent fails to come online because `nmblookup` utility generates core.

Workaround: Upgrade to the next available SRU version.

On Oracle Solaris SPARC SRU 11.3.17.0.5.0 and SRU 11.3.20.5.0, SambaShare resource does not go offline (3915234, 3931030)

On Oracle Solaris SPARC SRU 11.3.17.0.5.0 and SRU 11.3.20.5.0, SambaShare resource does not go offline because the `smbclient` utility is ambiguous. The utility generates different outputs for the same input and, in some scenarios, generates core.

Workaround: Use an older SRU version.

SambaServer agent does not come online after upgrading to Oracle Solaris x86 SRU 11.3.15.4.0 (3915235)

SambaServer agent does not come online after upgrading to Oracle Solaris x86 SRU 11.3.15.4.0 because:

- PidFile attribute is not populated
- Default configuration file is not loaded

Workaround: Specify the PidFile attribute for the SambaServer agent.

Issues related to the VCS database agents

This section describes the known issues about VCS database agents.

ASMDG agent does not go offline if the management DB is running on the same (3856460)

If an offline is fired on the node on which Flex ASM is running and the same node has Management DB running on it, then the same would not go offline.

Workaround: Use commands to migrate the Management DB to another node before getting the Flex ASM offline. You can run the following commands to check if the Management DB is running on a node:

```
# /oracle/12102/app/gridhome/bin/srvctl status mgmtldb -verbose
Database is enabled
Instance -MGMTDB is running on node vcslx017. Instance status: Open.
```

Run the following commands to migrate the Management DB to another node:

```
# /oracle/12102/app/gridhome/bin/srvctl relocate mgmtldb -node vcslx018
```

ASMDG on a particular does not go offline if its instances are being used by other database instances (3856450)

If you initiate an offline of the ASMDG group on a node which has its ASMInstance being used by one of more DB z resources from the cluster, then the offline would fail and a fault would get reported on both the ASM and DB level.

Workaround: Run the following SQL command to check the ASM DG running on the node:

```
SQL> select INST_ID, GROUP_NUMBER, INSTANCE_NAME,
DB_NAME, INSTANCE_NAME||':'||DB_NAME client_id from gv$asm_client;
```

INST_ID	GROUP_NUMBER	INSTANCE_NAME	DB_NAME	CLIENT_ID
3	2	oradb2	oradb	oradb2:oradb
3	2	oradb3	oradb	oradb3:oradb
3	2	+ASM3	+ASM	+ASM3:+ASM
3	1	+ASM3	+ASM	+ASM3:+ASM

1	2	oradb1	oradb	oradb1:oradb
1	1	-MGMTDB	_mgmtdb	-MGMTDB:_mgmtdb
1	1	+ASM1	+ASM	+ASM1:+ASM
4	2	oradb4	oradb	oradb4:oradb

8 rows selected.

In the above table:

- oradb1 is using the ASMInstance 1
- oradb2 and oradb3 are using ASMInstance 3
- oradb4 is using ASMInstance 4

Use the following SQL to relocate the ASMPool to another node:

```
SQL> alter system relocate client 'oradb4:oradb';
System altered.
```

If the command does not work, please refer Oracle documentation for further information on relocating the client.

Sometimes ASMDG reports as offline instead of faulted (3856454)

Sometimes, you may observe that the agent reports the ASMDG state for the node where the ASM instance is down as offline instead of as faulted, even when the cardinality is violated. This occurs in scenarios in which the ASM instance is abruptly shut down.

Workaround: No workaround.

Netlsnr agent monitoring can't detect tnslnsr running on Solaris if the entire process name exceeds 79 characters [3784547]

If the Oracle listener process is configured with a long name, consequently the tnslnsr process starts with a name longer than 79 characters. As a result, the proc structure doesn't show the full name of the Oracle listener process, and fails the Netlsnr agent monitoring.

Workaround: Configure shorter path or listener name, which does not exceed 79 characters.

The ASMinstAgent does not support having pfile/spfile for the ASM Instance on the ASM diskgroups

The ASMinstAgent does not support having pfile/spfile for the ASM Instance on the ASM diskgroups.

Workaround:

Have a copy of the pfile/spfile in the default \$GRID_HOME/dbs directory to make sure that this would be picked up during the ASM Instance startup.

VCS agent for ASM: Health check monitoring is not supported for ASMinst agent

The ASMinst agent does not support health check monitoring.

Workaround: Set the MonitorOption attribute to 0.

NOFAILOVER action specified for certain Oracle errors

The High Availability agent for Oracle provides enhanced handling of Oracle errors encountered during detailed monitoring. The agent uses the reference file oraerror.dat, which consists of a list of Oracle errors and the actions to be taken.

See the *Cluster Server Configuration and Upgrade Guide* for a description of the actions.

Currently, the reference file specifies the NOFAILOVER action when the following Oracle errors are encountered:

ORA-00061, ORA-02726, ORA-6108, ORA-06114

The NOFAILOVER action means that the agent sets the resource's state to OFFLINE and freezes the service group. You may stop the agent, edit the oraerror.dat file, and change the NOFAILOVER action to another action that is appropriate for your environment. The changes go into effect when you restart the agent.

ASMInstance resource monitoring offline resource configured with OHASD as application resource logs error messages in VCS logs [2846945]

When the Oracle High Availability Services Daemon (OHASD) is configured as an application resource to be monitored under VCS and if this resource is offline on

the failover node then the ASMInstance resource in the offline monitor logs the following error messages in the VCS logs:

```
ASMInst:asminst:monitor:Cluster Synchronization Service  
process is not running.
```

Workaround: Configure the application in a separate parallel service group and ensure that the resource is online.

Oracle agent fails to offline pluggable database (PDB) resource with PDB in backup mode [3592142]

If the PDB is in backup mode and if you attempt to offline the corresponding PDB resource, this will cause PDB resource to go into “Unable to Offline” state.

Workaround: Manually remove the PDB from the backup mode before attempting to take the PDB resource offline.

Clean succeeds for PDB even as PDB status is UNABLE to OFFLINE [3609351]

Oracle does not allow any operation on a PDB when the PDB is in backup mode. This is an expected behavior of Oracle. Therefore, a shutdown fails when it is initiated on a PDB in backup mode and returns an UNABLE TO OFFLINE status for the PDB. If PDB is removed from the backup mode using the SQL script, the agent framework is unable to change the UNABLE TO OFFLINE status of the PDB as clean is called. Since Oracle does not differentiate between clean and offline for PDB, clean succeeds for the PDB in spite of being in UNABLE TO OFFLINE state.

Workaround: No workaround.

Second level monitoring fails if user and table names are identical [3594962]

If the table inside CDB has same name as the user name, second level monitoring fails and Oracle agent fails to update the table. For example, if user name is `c##pdbuser1` and table is created as `c##pdbuser1.vcs`, then Oracle agent is unable to update it.

Workaround: Avoid having identical user and CDB table names.

Monitor entry point times out for Oracle PDB resources when CDB is moved to suspended state in Oracle 12.1.0.2 [3643582]

In Oracle-12.1.0.2.0, when CDB is in SUSPENDED mode, then the SQL command for PDB view (`v$pdb$`) hangs. Due to this, the monitor entry point in PDB gets timed out and there is no issue found in oracle-12.1.0.1.0 .

Workaround: No workaround.

Oracle agent fails to online and monitor Oracle instance if threaded_execution parameter is set to true [3644425]

In Oracle 12c, the threaded execution feature is enabled. The multithreaded Oracle Database model enables Oracle processes to execute as operating system threads in separate address spaces. If Oracle Database 12c is installed, the database runs in the process mode. If you set a parameter to run the database in threaded mode, some background processes on UNIX and Linux run with each process containing one thread, whereas the remaining Oracle processes run as threads within the processes.

When you enable this parameter, Oracle agent is unable to check smon (mandatory process check) and lgwr (optional process check) processes which were traditionally used for monitoring and which now run as threads.

Workaround: Disable the threaded execution feature as it is not supported on Oracle 12C.

Issues related to the agent framework

This section describes the known issues about the agent framework.

Agent framework cannot handle leading and trailing spaces for the dependent attribute (2027896)

Agent framework does not allow spaces in the target resource attribute name of the dependent resource.

Workaround: Do not provide leading and trailing spaces in the target resource attribute name of the dependent resource.

The agent framework does not detect if service threads hang inside an entry point [1442255]

In rare cases, the agent framework does not detect if all service threads hang inside a C entry point. In this case it may not cancel them successfully.

Workaround: If the service threads of the agent are hung, send a kill signal to restart the agent. Use the following command: `kill -9 hung_agent's_pid`. The `haagent -stop` command does not work in this situation.

IMF related error messages while bringing a resource online and offline [2553917]

For a resource registered with AMF, if you run `hagrp -offline` or `hagrp -online` explicitly or through a collective process to offline or online the resource respectively, the IMF displays error messages in either case.

The errors displayed is an expected behavior and it does not affect the IMF functionality in any manner.

Workaround: No workaround.

Delayed response to VCS commands observed on nodes with several resources and system has high CPU usage or high swap usage [3208239]

You may experience a delay of several minutes in the VCS response to commands if you configure large number of resources for monitoring on a VCS node and if the CPU usage is close to 100 percent or swap usage is very high.

Some of the commands are mentioned below:

- `# hares -online`
- `# hares -offline`
- `# hagrp -online`
- `# hagrp -offline`
- `# hares -switch`

The delay occurs as the related VCS agent does not get enough CPU bandwidth to process your command. The agent may also be busy processing large number of pending internal commands (such as periodic monitoring of each resource).

Workaround: Change the values of some VCS agent type attributes which are facing the issue and restore the original attribute values after the system returns to the normal CPU load.

- 1 Back up the original values of attributes such as MonitorInterval, OfflineMonitorInterval, and MonitorFreq of IMF attribute.
- 2 If the agent does not support Intelligent Monitoring Framework (IMF), increase the value of MonitorInterval and OfflineMonitorInterval attributes.

```
# haconf -makerw
# hatype -modify <TypeName> MonitorInterval <value>
# hatype -modify <TypeName> OfflineMonitorInterval <value>
# haconf -dump -makero
```

Where <TypeName> is the name of the agent with which you are facing delays and <value> is any numerical value appropriate for your environment.

- 3 If the agent supports IMF, increase the value of MonitorFreq attribute of IMF.

```
# haconf -makerw
# hatype -modify <TypeName> IMF -update MonitorFreq <value>
# haconf -dump -makero
```

Where <value> is any numerical value appropriate for your environment.

- 4 Wait for several minutes to ensure that VCS has executed all pending commands, and then execute any new VCS command.
- 5 If the delay persists, repeat step 2 or 3 as appropriate.
- 6 If the CPU usage returns to normal limits, revert the attribute changes to the backed up values to avoid the delay in detecting the resource fault.

CFSMount agent may fail to heartbeat with VCS engine and logs an error message in the engine log on systems with high memory load [3060779]

On a system with high memory load, CFSMount agent may fail to heartbeat with VCS engine resulting into V-16-1-53030 error message in the engine log.

VCS engine must receive periodic heartbeat from CFSMount agent to ensure that it is running properly on the system. The heartbeat is decided by AgentReplyTimeout attribute. Due to high CPU usage or memory workload (for example, swap usage greater than 85%), agent may not get enough CPU cycles to schedule. This causes heartbeat loss with VCS engine and as a result VCS engine terminates the agent and starts the new agent. This can be identified with the following error message in the engine log:


```
V-16-1-53030 Termination request sent to CFSSMount  
agent process with pid %d
```

Workaround: Increase the AgentReplyTimeout value and see if CFSSMount agent becomes stable. If this does not resolve the issue then try the following workaround. Set value of attribute NumThreads to 1 for CFSSMount agent by running following command:

```
# hatype -modify CFSSMount NumThreads 1
```

Even after the above command if CFSSMount agent keeps on terminating, report this to Veritas support team.

Logs from the script executed other than the agent entry point goes into the engine logs [3547329]

The agent logs of C-based and script-based entry points get logged in the agent log when the attribute value of LogViaHalog is set to 1 (one). To restore to the older logging behavior in which C-based entry point logs were logged in agent logs and script-based entry point logs were logged in engine logs, you can set the LogViaHalog value as 0 (zero). However, it is observed that some C-based entry point logs continue to appear in the engine logs even when LogViaHalog is set to 1 (one). This issue is observed on all the database agents.

Workaround: No workaround.

VCS fails to process the `hares -add` command resource if the resource is deleted and subsequently added just after the VCS process or the agent's process starts (3813979)

When VCS or the agent processes start, the agent processes the initial snapshots from the engine before probing the resource. During the processing of the snapshots, VCS fails to process the `hares -add` command, thereby skipping the resource addition operation and subsequently failing to probe the resource.

Workaround: This behavior is by the current design of the agent framework.

Issues related to Intelligent Monitoring Framework (IMF)

This section describes the known issues of Intelligent Monitoring Framework (IMF).

Registration error while creating a Firedrill setup [2564350]

While creating the Firedrill setup using the `Firedrill setup` utility, VCS encounters the following error:

```
AMF amfregister ERROR V-292-2-167
Cannot register mount offline event
```

During Firedrill operations, VCS may log error messages related to IMF registration failure in the engine log. This happens because in the firedrill service group, there is a second CFSMount resource monitoring the same MountPoint through IMF. Both the resources try to register for online/offline events on the same MountPoint and as a result, registration of one fails.

Workaround: No workaround.

IMF does not fault zones if zones are in ready or down state [2290883]

IMF does not fault zones if zones are in ready or down state.

IMF does not detect if zones are in ready or down state. In Ready state, there are no services running inside the running zones.

Workaround: Offline the zones and then restart.

IMF does not detect the zone state when the zone goes into a maintenance state [2535733]

IMF does not detect the change in state. However, the change in state is detected by Zone monitor in the next cycle.

Workaround: No workaround.

IMF does not provide notification for a registered disk group if it is imported using a different name (2730774)

If a disk group resource is registered with the AMF and the disk group is then imported using a different name, AMF does not recognize the renamed disk group and hence does not provide notification to DiskGroup agent. Therefore, the DiskGroup agent keeps reporting the disk group resource as offline.

Workaround: Make sure that while importing a disk group, the disk group name matches the one registered with the AMF.

Direct execution of `linkamf` displays syntax error [2858163]

Bash cannot interpret Perl when executed directly.

Workaround: Run `linkamf` as follows:

```
# /opt/VRTSperl/bin/perl /opt/VRTSamf/imf/linkamf <destination-directory>
```

Error messages displayed during reboot cycles [2847950]

During some reboot cycles, the following message might get logged in the engine log:

```
AMF libvxamf ERROR V-292-2-149 Cannot unregister event: no rid -1 found
AMF libvxamf ERROR V-292-2-306 Unable to unregister all events (errno:405)
```

This does not have any effect on the functionality of IMF.

Workaround: No workaround.

Error message displayed when ProPCV prevents a process from coming ONLINE to prevent concurrency violation does not have I18N support [2848011]

The following message is seen when ProPCV prevents a process from coming ONLINE to prevent concurrency violation. The message is displayed in English and does not have I18N support.

```
Concurrency Violation detected by VCS AMF.
Process <process-details> will be prevented from startup.
```

Workaround: No Workaround.

AMF displays StartProgram name multiple times on the console without a VCS error code or logs [2872064]

When VCS AMF prevents a process from starting, it displays a message on the console and in syslog. The message contains the signature of the process that was prevented from starting. In some cases, this signature might not match the signature visible in the PS output. For example, the name of the shell script that was prevented from executing will be printed twice.

Workaround: No workaround.

VCS engine shows error for cancellation of reaper when Apache agent is disabled [3043533]

When `haimfconfig` script is used to disable IMF for one or more agents, the VCS engine logs the following message in the engine log:

```
AMF imf_getnotification ERROR V-292-2-193  
Notification(s) canceled for this reaper.
```

This is an expected behavior and not an issue.

Workaround: No workaround.

Terminating the `imfd` daemon orphans the `vxnotify` process [2728787]

If you terminate `imfd` daemon using the `kill -9` command, the `vxnotify` process created by `imfd` does not exit automatically but gets orphaned. However, if you stop `imfd` daemon with the `amfconfig -D` command, the corresponding `vxnotify` process is terminated.

Workaround: The correct way to stop any daemon is to gracefully stop it with the appropriate command (which is `amfconfig -D` command in this case), or to terminate the daemon using Session-ID. Session-ID is the -PID (negative PID) of the daemon.

For example:

```
# kill -9 -27824
```

Stopping the daemon gracefully stops all the child processes spawned by the daemon. However, using `kill -9 pid` to terminate a daemon is not a recommended option to stop a daemon, and subsequently you must kill other child processes of the daemon manually.

Agent cannot become IMF-aware with agent directory and agent file configured [2858160]

Agent cannot become IMF-aware if Agent Directory and Agent File are configured for that agent.

Workaround: No workaround.

ProPCV fails to prevent a script from running if it is run with relative path [3617014]

If the absolute path is registered with AMF for prevention and the script is run with the relative path, AMF fails to prevent the script from running.

Workaround: No workaround.

Issues related to global clusters

This section describes the known issues about global clusters.

The engine log file receives too many log messages on the secure site in global cluster environments [1919933]

When the WAC process runs in secure mode on one site, and the other site does not use secure mode, the engine log file on the secure site gets logs every five seconds.

Workaround: The two WAC processes in global clusters must always be started in either secure or non-secure mode. The secure and non-secure WAC connections will flood the engine log file with the above messages.

Application group attempts to come online on primary site before fire drill service group goes offline on the secondary site (2107386)

The application service group comes online on the primary site while the fire drill service group attempts to go offline at the same time, causing the application group to fault.

Workaround: Ensure that the fire drill service group is completely offline on the secondary site before the application service group comes online on the primary site.

Issues related to the Cluster Manager (Java Console)

This section describes the known issues about Cluster Server Manager (Java Console).

Some Cluster Manager features fail to work in a firewall setup [1392406]

In certain environments with firewall configurations between the Cluster Manager and the VCS cluster, the Cluster Manager fails with the following error message:

```
V-16-10-13 Could not create CmdClient. Command Server  
may not be running on this system.
```

Workaround: You must open port 14150 on all the cluster nodes.

VCS Cluster Configuration wizard issues

IPv6 verification fails while configuring generic application using VCS Cluster Configuration wizard [3614680]

The VCS Cluster Configuration wizard fails to check whether IPv6 IP is already plumbed while configuring a generic application through the Virtual IP page. The wizard does neither displays a warning if IPv6 IP is already plumbed elsewhere nor indicates whether it is reachable through a ping.

Workaround: Manually ensure that IPv6 is not plumbed elsewhere on the network before configuring the generic application through the wizard.

InfoScale Enterprise: Unable to configure clusters through the VCS Cluster Configuration wizard (3911694)

Unable to configure clusters through the VCS Cluster Configuration wizard in InfoScale Enterprise because the `configure_cluster.response` file that is created has the product type as `Availability` instead of `Enterprise`.

Workaround: In InfoScale Enterprise, configure clusters through the CPI.

LLT known issues

This section covers the known issues related to LLT in this release.

LLT port stats sometimes shows recvcnt larger than recvbytes (1907228)

With each received packet, LLT increments the following variables:

- `recvcnt` (increment by one for every packet)
- `recvbytes` (increment by size of packet for every packet)

Both these variables are integers. With constant traffic, `recvbytes` hits and rolls over `MAX_INT` quickly. This can cause the value of `recvbytes` to be less than the value of `recvcnt`.

This does not impact the LLT functionality.

Cannot configure LLT if full device path is not used in the lltab file (2858159)

(Oracle Solaris 11) On virtual machines ensure that you use the full path of the devices corresponding to the links in `lltab`. For example, use `/dev/net/net1` instead of `/dev/net/net:1` in the `lltab` file, otherwise you cannot configure LLT.

Fast link failure detection is not supported on Solaris 11 (2954267)

Fast link failure detection is not supported on Solaris 11 operating system because the operating system cannot provide notification calls to LLT when a link failure occurs. If the operating system kernel notifies LLT about the link failure, LLT can detect a link failure much earlier than the regular link failure detection cycle. As Solaris 11 does not notify LLT about link failures, failure detection cannot happen before the regular detection cycle.

Workaround: None

I/O fencing known issues

This section describes the known issues in this release of I/O fencing.

CP server repetitively logs unavailable IP addresses (2530864)

If coordination point server (CP server) fails to listen on any of the IP addresses that are mentioned in the `vxcps.conf` file or that are dynamically added using the command line, then CP server logs an error at regular intervals to indicate the failure. The logging continues until the IP address is bound to successfully.

```
CPS ERROR V-97-51-103 Could not create socket for host
10.209.79.60 on port 14250
CPS ERROR V-97-1400-791 Coordination point server could not
open listening port = [10.209.79.60]:14250
Check if port is already in use.
```

Workaround: Remove the offending IP address from the listening IP addresses list using the `rm_port` action of the `cpsadm` command.

For details, see the Configuration and Upgrade Guide for the appropriate InfoScale component.

Fencing port b is visible for few seconds even if cluster nodes have not registered with CP server (2415619)

Even if the cluster nodes have no registration on the CP server and if you provide coordination point server (CP server) information in the `vxfenmode` file of the cluster nodes, and then start fencing, the fencing port b is visible for a few seconds and then disappears.

Workaround: Manually add the cluster information to the CP server to resolve this issue. Alternatively, you can use installer as the installer adds cluster information to the CP server during configuration.

The `cpsadm` command fails if LLT is not configured on the application cluster (2583685)

The `cpsadm` command fails to communicate with the coordination point server (CP server) if LLT is not configured on the application cluster node where you run the `cpsadm` command. You may see errors similar to the following:

```
# cpsadm -s 10.209.125.200 -a ping_cps
CPS ERROR V-97-1400-729 Please ensure a valid nodeid using
environment variable
CPS_NODEID
CPS ERROR V-97-1400-777 Client unable to communicate with CPS.
```

However, if you run the `cpsadm` command on the CP server, this issue does not arise even if LLT is not configured on the node that hosts CP server. The `cpsadm` command on the CP server node always assumes the LLT node ID as 0 if LLT is not configured.

According to the protocol between the CP server and the application cluster, when you run the `cpsadm` on an application cluster node, `cpsadm` needs to send the LLT node ID of the local node to the CP server. But if LLT is unconfigured temporarily, or if the node is a single-node VCS configuration where LLT is not configured, then the `cpsadm` command cannot retrieve the LLT node ID. In such situations, the `cpsadm` command fails.

Workaround: Set the value of the `CPS_NODEID` environment variable to 255. The `cpsadm` command reads the `CPS_NODEID` variable and proceeds if the command is unable to get LLT node ID from LLT.

When I/O fencing is not up, the `svcs` command shows VxFEN as online (2492874)

Solaris SMF marks the service status based on the exit code of the start method for that service. The VxFEN start method executes the `vxfen-startup` script in the background and exits with code 0. Hence, if the `vxfen-startup` script subsequently exits with failure then this change is not propagated to SMF. This behavior causes the `svcs` command to show incorrect status for VxFEN.

Workaround: Use the `vxfenadm` command to verify that I/O fencing is running.

In absence of cluster details in CP server, VxFEN fails with pre-existing split-brain message (2433060)

When you start server-based I/O fencing, the node may not join the cluster and prints error messages in logs similar to the following:

In the `/var/VRTSvcs/log/vxfen/vxfen.log` file:

```
VXFEN vxfenconfig ERROR V-11-2-1043  
Detected a preexisting split brain. Unable to join cluster.
```

In the `/var/VRTSvcs/log/vxfen/vxfen.log` file:

```
operation failed.  
CPS ERROR V-97-1400-446 Un-authorized user cpsclient@sys1,  
domaintype vx; not allowing action
```

The `vxfend` daemon on the application cluster queries the coordination point server (CP server) to check if the cluster members as seen in the GAB membership are registered with the CP server. If the application cluster fails to contact the CP server due to some reason, then fencing cannot determine the registrations on the CP server and conservatively assumes a pre-existing split-brain.

Workaround: Before you attempt to start VxFEN on the application cluster, ensure that the cluster details such as cluster name, UUID, nodes, and privileges are added to the CP server.

The `vxfsnwap` utility does not detect failure of coordination points validation due to an RSH limitation (2531561)

The `vxfsnwap` utility runs the `vxfenconfig -o modify` command over RSH or SSH on each cluster node for validation of coordination points. If you run the `vxfsnwap` command using RSH (with the `-n` option), then RSH does not detect the failure of validation of coordination points on a node. From this point, `vxfsnwap` proceeds as if the validation was successful on all the nodes. But, it fails at a later stage when it tries to commit the new coordination points to the VxFEN driver. After the failure, it rolls back the entire operation, and exits cleanly with a non-zero error code. If you run `vxfsnwap` using SSH (without the `-n` option), then SSH detects the failure of validation of coordination of points correctly and rolls back the entire operation immediately.

Workaround: Use the `vxfsnwap` utility with SSH (without the `-n` option).

Fencing does not come up on one of the nodes after a reboot (2573599)

If VxFEN unconfiguration has not finished its processing in the kernel and in the meantime if you attempt to start VxFEN, you may see the following error in the `/var/VRTSvcS/log/vxfen/vxfen.log` file:

```
VXFEN vxfenconfig ERROR V-11-2-1007 Vxfen already configured
```

However, the output of the `gabconfig -a` command does not list port b. The `vxfenadm -d` command displays the following error:

```
VXFEN vxfenadm ERROR V-11-2-1115 Local node is not a member of cluster!
```

Workaround: Start VxFEN again after some time.

Hostname and username are case sensitive in CP server (2846392)

The hostname and username on the CP server are case sensitive. The hostname and username used by fencing to communicate with CP server must be in same case as present in CP server database, else fencing fails to start.

Workaround: Make sure that the same case is used in the hostname and username on the CP server.

Server-based fencing comes up incorrectly if default port is not mentioned (2403453)

When you configure fencing in customized mode and do not provide default port, fencing comes up. However, the `vxfenconfig -l` command output does not list the port numbers.

Workaround: Retain the "port_https=<port_value>" setting in the `/etc/vxfenmode` file, when using customized fencing with at least one CP server. The default port value is 443.

Fencing may show the RFSM state as replaying for some nodes in the cluster (2555191)

Fencing based on coordination point clients in Campus cluster environment may show the RFSM state as replaying for some nodes in the cluster.

Workaround:

Restart fencing on the node that shows RFSM state as replaying.

The vxfenswap utility deletes comment lines from the `/etc/vxfemode` file, if you run the utility with hacli option (3318449)

The vxfenswap utility uses RSH, SSH, or hacli protocol to communicate with peer nodes in the cluster. When you use vxfenswap to replace coordination disk(s) in disk-based fencing, vxfenswap copies `/etc/vxfenmode` (local node) to `/etc/vxfenmode` (remote node).

With the hacli option, the utility removes the comment lines from the remote `/etc/vxfenmode` file, but, it retains comments in the local `/etc/vxfenmode` file.

Workaround: Copy the comments manually from local `/etc/vxfenmode` to remote nodes.

The vxfentsthdw utility may not run on systems installed with partial SFHA stack [3333914]

The vxfentsthdw utility runs if the SFHA stack and VCS are fully installed with properly configured SF and VxVM. It also runs if the entire SFHA stack and VCS are not installed. However, partial installs where SF is installed and configured but VCS is not installed is not supported. The utility will display an error with the `-g` or `-c` options.

Workaround: Install VRTSvxfen package, then run the utility from either the install media or from the `/opt/VRTSvcs/vxfen/bin/` location.

When a client node goes down, for reasons such as node panic, I/O fencing does not come up on that client node after node restart (3341322)

This issue happens when one of the following conditions is true:

- Any of the CP servers configured for HTTPS communication goes down.
- The CP server service group in any of the CP servers configured for HTTPS communication goes down.
- Any of the VIPs in any of the CP servers configured for HTTPS communication goes down.

When you restart the client node, fencing configuration starts on the node. The fencing daemon, vxfsd, invokes some of the fencing scripts on the node. Each of these scripts has a timeout value of 120 seconds. If any of these scripts fails, fencing configuration fails on that node.

Some of these scripts use `cpsadm` commands to communicate with CP servers. When the node comes up, `cpsadm` commands try to connect to the CP server using

VIPs for a timeout value of 60 seconds. So, if the multiple `cpsadm` commands that are run within a single script exceed the timeout value, then the total timeout value exceeds 120 seconds, which causes one of the scripts to time out. Hence, I/O fencing does not come up on the client node.

Note that this issue does not occur with IPM-based communication between CP server and client clusters.

Workaround: Fix the CP server.

The `vxfenconfig -l` command output does not list Coordinator disks that are removed using the `vxddmpadm exclude dmpnodename=<dmp_disk/node>` command [3644431]

After you remove a Coordinator disk used by fencing or fencing disk group by running the `vxddmpadm exclude dmpnodename=<dmp_disk/node>` command, the removed disk is not listed in the `vxfenconfig -l` command output.

In case of a split brain, the `vxfen` program cannot use the removed disk as a coordination point in the subsequent fencing race.

Workaround: Run the `vxddmpadm include dmpnodename=<dmp_disk/node>` command to again enable the dmp disk. This disk will show up in subsequent `vxfenconfig -l` output.

Stale `.vxfendargs` file lets hashadow restart `vxfeed` in Sybase mode (2554886)

When I/O fencing is configured in customized mode, `vxfeed`, the user mode daemon of I/O fencing, creates the `/opt/VRTSvcs/lock/.vxfendargs` file. VCS uses this file to restart the `vxfeed` daemon when it gets killed. However, VCS does not use this file when I/O fencing is configured in Sybase mode. This file is not removed from the system when I/O fencing is unconfigured.

If user configures I/O fencing in Sybase mode and an old `/opt/VRTSvcs/lock/.vxfendargs` file is present in the system from an earlier configuration of I/O fencing in customized mode, then VCS attempts to restart the `vxfeed` daemon every time it is killed. This interferes with the functioning of I/O fencing in the Sybase mode.

Workaround: Before you configure I/O fencing in Sybase mode, delete the `/opt/VRTSvcs/lock/.vxfendargs` file if it is present in the system.

CP server configuration fails while setting up secure credentials for CP server hosted on an SFHA cluster (2621029)

When you configure CP server using the `configure_cps.pl` utility, the configuration fails while setting up secure credentials for CP server that is hosted on an SFHA cluster. You may see the following error:

```
Creating softlink to credential directory /etc/VRTScps/db/CPSERVER
on node nodename.
Unable to connect to node nodename using /usr/bin/ssh.
Please configure ssh communication and retry. Exiting.
```

Workaround: You can use any of the following options:

- Before running the `configure_cps.pl` utility, change the default shell for root user to either KSH or bash.
- Perform the following steps after running the `configure_cps.pl` utility on each node of the cluster:

- Manually remove the old credential directory or softlink. For example:

```
# rm -rf /var/VRTSvcs/vcsauth/data/CPSERVER
```

- Create a new soft-link to the shared location of the credential directory:

```
# ln -s path_of_CP_server_credential_directory \
/var/VRTSvcs/vcsauth/data/CPSERVER
```

- Start the CPSSG service group:

```
# hagrps -online CPSSG -any
```

The CoordPoint agent faults after you detach or reattach one or more coordination disks from a storage array (3317123)

After you detach or reattach a coordination disk from a storage array, the CoordPoint agent may fault because it reads an older value stored in the I/O fencing kernel module.

Workaround: Run the `vxfsnwap` utility to refresh the registration keys on the coordination points for both server-based I/O fencing and disk-based I/O fencing. But, even if the registrations keys are not lost, you must run the `vxfsnwap` utility to refresh the coordination point information stored in the I/O fencing kernel module.

For more information on refreshing registration keys on the coordination points for server-based and disk-based I/O fencing, refer to the *Cluster Server Administrator's Guide*.

The upper bound value of FaultTolerance attribute of CoordPoint agent should be less than the majority of the coordination points. (2846389)

The upper bound value of `FaultTolerance` attribute of `CoordPoint` agent should be less than the majority of the coordination points. Currently this value is less than the number of coordination points.

GAB known issues

This section covers the known issues related to GAB in this release.

GAB may fail to stop during a phased upgrade on Oracle Solaris 11 (2858157)

While performing a phased upgrade on Oracle Solaris 11 systems, GAB may fail to stop. However, CPI gives a warning and continues with stopping the stack.

Workaround: Reboot the node after the installer completes the upgrade.

Cannot run pfiles or truss files on gablogd (2292294)

When `pfiles` or `truss` is run on `gablogd`, a signal is issued to `gablogd`. `gablogd` is blocked since it has called an `gab_ioctl` and is waiting for events. As a result, the `pfiles` command hangs.

Workaround: None.

(Oracle Solaris 11) On virtual machines, sometimes the common product installer (CPI) may report that GAB failed to start and may exit (2879262)

GAB startup script may take longer than expected to start up. The delay in start up can cause the CPI to report that GAB failed and exits.

Workaround: Manually start GAB and all dependent services.

During upgrade, GAB kernel module fails to unload [3560458]

On Solaris 11U1(SRU >= 8), if the SFRAC/SFCFSHA component is upgraded from version 6.0 or later to version 7.0, GAB module fails to unload and remains in loaded state with old version.

Workaround:

You can resolve this issue by rebooting the system. Rebooting the system unloads the module successfully.

Storage Foundation and High Availability known issues

This section describes the known issues in this release of Storage Foundation and High Availability (SFHA). These known issues apply to Veritas InfoScale Enterprise.

Cache area is lost after a disk failure (3158482)

SmartIO supports one VxFS cache area and one VxVM cache area. If you create one cache area, and the disk fails, the cache area becomes disabled. If you attempt to create a second cache area of the other type before the cache disk group is enabled, then the first cache area is lost. It cannot be brought online.

For example, first you created a VxFS cache area. The disk failed and the cache area is disabled. Now create the VxVM cache area. While creating VxVM cache area, SmartIO looks for an existing default cache area. Due to the failed disk, the existing cache area cannot be found. So SmartIO creates a VxVM cache area with the same name. Now even if disk containing VxFS cache area comes up, SmartIO cannot access the original cache area. In this scenario, the VxFS cache area is lost. Losing the cache area in this case does not result into any data loss or data inconsistency issues.

Workaround:

Create a new VxFS cache area.

NFS issues with VxFS Storage Checkpoints (2027492)

NFS clients mounting VxFS Storage Checkpoints that are NFS-exported by SFCFSHA cluster nodes using a Virtual IP may receive the following error message upon Virtual IP failover:

NFS clients mounting VxFS Storage Checkpoints that are NFS-exported by SFHA cluster nodes using a Virtual IP may receive the following error message upon Virtual IP failover:

```
Stale NFS file handle
```

This is a result of major numbers of VxFS Storage Checkpoints not necessarily being the same on all SFCFSA cluster nodes.

This is a result of major numbers of VxFS Storage Checkpoints not necessarily being the same on all SFHA cluster nodes.

Workaround: There is no workaround for this issue.

Some SmartTier for Oracle commands do not work correctly in non-POSIX locales (2138030)

Some SmartTier for Oracle commands do not work correctly in non-POSIX locale settings.

Workaround: Set the environment variable `LANG=C` systemwide in the `/etc/profile` file.

In an IPv6 environment, db2icrt and db2idrop commands return a segmentation fault error during instance creation and instance removal (1602444)

When using IBM DB2 `db2icrt` command to create a DB2 database instance on a pure IPv6 environment, the `db2icrt` command returns segmentation fault error message. For example:

```
$ /opt/ibm/db2/V9.5/instance/db2icrt -a server -u db2fen1 db2inst1
/opt/ibm/db2/V9.5/instance/db2iutil: line 4700: 26182 Segmentation fault
$ {DB2DIR?}/instance/db2isrv -addfcm -i ${INSTNAME?}
```

The `db2idrop` command also returns segmentation fault, but the instance is removed successfully after the `db2idrop` command is issued. For example:

```
$ /opt/ibm/db2/V9.5/instance/db2idrop db2inst1
/opt/ibm/db2/V9.5/instance/db2iutil: line 3599: 7350 Segmentation fault
$ {DB2DIR?}/instance/db2isrv -remove -s DB2_${INSTNAME?} 2> /dev/null
```

```
DBI1070I Program db2idrop completed successfully.
```

This happens on DB2 9.1, 9.5, and 9.7.

This issue has been identified as an IBM issue. Once IBM has fixed this issue, then IBM will provide a hotfix for this segmentation problem.

At this time, you can communicate in a dual-stack to avoid the segmentation fault error message until IBM provides a hotfix.

To communicate in a dual-stack environment

- ◆ Add an IPv6 hostname as an IPv4 loopback address to the `/etc/hosts` file. For example:

```
127.0.0.1 swlx20-v6
```

Or

```
127.0.0.1 swlx20-v6.punipv6.com
```

127.0.0.1 is the IPv4 loopback address.

swlx20-v6 and swlx20-v6.punipv6.com are the IPv6 hostnames.

Oracle 11gR1 may not work on pure IPv6 environment (1819585)

There is problem running Oracle 11gR1 on a pure IPv6 environment.

Tools like `dbca` may hang during database creation.

Workaround: There is no workaround for this, as Oracle 11gR1 does not fully support pure IPv6 environment. Oracle 11gR2 release may work on a pure IPv6 enviroment, but it has not been tested or released yet.

Not all the objects are visible in the VOM GUI (1821803)

After upgrading SF stack from 5.0MP3RP2 to 5.1, the volumes are not visible under the Volumes tab and the shared diskgroup is discovered as Private and Deported under the Diskgroup tab in the VOM GUI.

Workaround:

To resolve this known issue

- ◆ On each manage host where `VRTSsfmh 2.1` is installed, run:

```
# /opt/VRTSsfmh/adm/dclisetup.sh -U
```

An error message is received when you perform off-host clone for RAC and the off-host node is not part of the CVM cluster (1834860)

There is a known issue when you try to perform an off-host clone for RAC and the off-host node is not part of the CVM cluster. You may receive a similar error message:

```
Cannot open file /etc/vx/vxdba/rac11g1/.DB_NAME
(No such file or directory).
SFORA vxreptadm ERROR V-81-8847 Cannot get filename from sid
for 'rac11g1', rc=-1.
SFORA vxreptadm ERROR V-81-6550 Could not connect to repository
database.
VxVM vxdg ERROR V-5-1-582 Disk group SNAP_rac11dgl: No such disk
group SFORA
vxsnapadm ERROR V-81-5623 Could not get CVM information for
SNAP_rac11dgl.
SFORA dbed_vmclonedb ERROR V-81-5578 Import SNAP_rac11dgl failed.
```

Workaround: Currently there is no workaround for this known issue. However, if the off-host node is part of the CVM cluster, then off-host clone for RAC works fine.

Also the `dbed_vmclonedb` command does not support `LOCAL_LISTENER` and `REMOTE_LISTENER` in the `init.ora` parameter file of the primary database.

A volume's placement class tags are not visible in the Veritas Enterprise Administrator GUI when creating a dynamic storage tiering placement policy (1880081)

A volume's placement class tags are not visible in the Veritas Enterprise Administrator (VEA) GUI when you are creating a SmartTier placement policy if you do not tag the volume with the placement classes prior to constructing a volume set for the volume.

Workaround: To see the placement class tags in the VEA GUI, you must tag the volumes prior to constructing the volume set. If you already constructed the volume set before tagging the volumes, restart `vxsvc` to make the tags visible in the GUI.

Storage Foundation Cluster File System High Availability known issues

This section describes the known issues in this release of Storage Foundation Cluster File System High Availability (SFCFSHA). These known issues apply to the following products:

- Veritas InfoScale Storage
- Veritas InfoScale Enterprise

In an FSS environment, creation of mirrored volumes may fail for SSD media [3932494]

In an FSS environment where SSD devices are used from Storage Access Layer (SAL), the creation of mirrored volumes may fail if `vxconfigd` is restarted on the master node.

This issue occurs because the `Mediatype` attribute for a device is inconsistently propagated from the kernel during `vxconfigd` startup.

Workaround: Before creating a disk group, set the media type attribute to SSD

```
vxdisk set -f diskname mediatype=ssd
```

Mount command may fail to mount the file system (3913246)

For a file system that was earlier mounted on cluster nodes the first `Mount` command may fail with the following error:

```
UX:vxfs mount.vxfs: ERROR: V-3-28543: Cannot be mounted until it has been cle
Please run "fsck -t vxfs -y /dev/vx/rdsk/<DiskGroup_Name>/<Volume_Name>" befo
Please refer to fsck_vxfs man page for details.
```

At the same time, the following error message may appear in the system log:

```
vxfs: msgcnt 463 msg 021: V-2-21: vx_fs_init - /dev/vx/dsk/<DiskGroup_Name>/
file system validation failure.
```

These are generic messages and the `Mount` command may fail due to multiple reasons.

Workaround:

1. Verify all the per node logs to check for a dirty log, if any.
2. Mount the file system with `delayfsck` mount option.

Notes:

- This workaround is applicable only if a dirty log is found.
- Step 2 is applicable only for disk layout version 11 or later. For disk layout version prior to 11, you must run the `Full fsck` command before the file system is mounted.
- Even if you have mounted the file system with `delayfsck` mount option, you must run the `Full fsck` command at a later point in time. You may plan the application downtime and then run the `Full fsck` command.

After the local node restarts or panics, the FSS service group cannot be online successfully on the local node and the remote node when the local node is up again (3865289)

When all the nodes that are contributing storage to a shared Flexible Storage Sharing (FSS) DG leave the cluster, the CVMVolDG resources and their dependent resources such as CFSSMount will be FAULTED. When the nodes rejoin the cluster, the resources/service groups will still remain in the FAULTED or OFFLINE state.

Workaround:

The FAULT on these resources should be manually CLEARED and the OFFLINED resources or service groups should be manually ONLINED.

- To clear the fault on the resource, use the following command:

```
# hares -clear <res> [-sys <system>]
```

- To bring the individual OFFLINED resource to the ONLINE state, use the following command:

```
# hares -online [-force] <res> -sys <system>
```

- To bring all the OFFLINED resource under a service group to the ONLINE state, use the following command:

```
# hagr -online [-force] <group> -any [-clus <cluster> | -localclus]
```

In the FSS environment, if DG goes to the dgdisable state and deep volume monitoring is disabled, successive node joins fail with error 'Slave failed to create remote disk: retry to add a node failed' (3874730)

In the Flexible Storage Sharing (FSS) environment, if deep monitoring is not enabled for the volume used for the file system, the CVMVolDg agent is able to detect fault and deport the disabled DG. Any new node joining to the cluster fails with error:

```
# /opt/VRTS/bin/vxclustadm -v nodestate
state: out of cluster
reason: Slave failed to create remote disk: retry to add a node failed
```

Workaround:

Enable deep monitoring for the resource using the '-D' option during adding the service group:

```
# cfsmntadm add -D <dgname> <volname> <mountpoint>all=cluster
```

If you have created the service group, use the below command to enable the deep monitoring of volumes:

```
# hares -modify <res_name> CVMVolumeIoTest <vol_list>
```

DG creation fails with error "V-5-1-585 Disk group punedatadg: cannot create: SCSI-3 PR operation failed" on the VSCSI disks (3875044)

If the disks that do not support SCSI3 PR are used to create the shared disk group, the operation fails as the data disk fencing functionality cannot be provided on such disks. The operation fails with error:

```
VxVM vxdg ERROR V-5-1-585 Disk group <DGNAME>: cannot create: SCSI-3
PR operation failed
```

Workaround:

If you still want to allow such disks to be part of shared disk group, disable the data disk fencing functionality in the cluster by running the command on all the nodes in the cluster:

```
# vxdctl scsi3pr off
```

After the disabling process, take caution that it may not protect the disks against the ghost I/Os from nodes that are not part of the cluster.

Write back cache is not supported on the cluster in FSS scenario [3723701]

Write back cache is not supported in FSS scenario on Cluster file system. When the Write back is enabled, for example, node N1 and N2 both have its own SSD and they are using each other's SSD as remote cache. Then it may cause data corruption and the recovery is not possible on cluster.

Workaround: This issue has been fixed.

CVMVOLDg agent is not going into the FAULTED state. [3771283]

In CVMVOLDg monitor script we are not able to parse a variable and hence the volume does not go into the disabled state. This is the reason why the CVMVOLDg agent is not going into the FAULTED state.

Workaround:

Enable CVMVOLIOTEST on the volume for the resource to go into FAULTED state, using the following commands:

```
# haconf -makerw

# hares -modify test_vol_dg CVMVolumeIoTest testvol

# haconf -dump -makero
```

CFS commands might hang when run by non-root (3038283)

The CFS commands might hang when run by non-root.

Workaround

To resolve this issue

- ◆ Use `halogin` command to save the authentication information before running any CFS commands on a non-root session.

When you run the `halogin` command, VCS stores encrypted authentication information in the user's home directory.

The fsappadm subfilemove command moves all extents of a file (3258678)

This issue occurs under following conditions:

- You run the `fsppadm subfilemove` command from a cluster file system (CFS) secondary node.

- You specify a range of extents for relocation to a target tier.

If the extent size is greater than or equal to 32768, the `fsppadm subfilemove` command moves all extents of the specified table to the target tier. The expectation is to move a specified range of extents.

Workaround:

- ◆ On the CFS primary node, determine the primary node using one of the following commands:

```
# fsclustadm showprimary mountpoint

# fsclustadm idtoname nodeid
```

Certain I/O errors during clone deletion may lead to system panic. (3331273)

Certain I/O errors during clone deletion may lead to system panic.

Workaround:

There is no workaround for this issue.

Panic due to null pointer de-reference in vx_bmap_lookup() (3038285)

If you use the `fsadm -b` command on a CFS secondary node to resize the file system, it might fail with the following error message printed in the syslog:

```
Reorg of inode with shared extent larger than 32768 blocks
can be done only on the CFS Primary node
```

Workaround: Resize the file system with the `fsadm` command from the primary node of the cluster.

In a CFS cluster, that has multi-volume file system of a small size, the fsadm operation may hang (3348520)

In a CFS cluster, that has multi-volume file system of a small size, the `fsadm` operation may hang, when the free space in the file system is low.

Workaround: There is no workaround for this issue.

Storage Foundation for Oracle RAC known issues

This section describes the known issues in this release of Storage Foundation for Oracle RAC (SFRAC). These known issues apply to Veritas InfoScale Enterprise.

Oracle RAC known issues

This section lists the known issues in Oracle RAC.

Oracle Grid Infrastructure installation may fail with internal driver error

The Oracle Grid Infrastructure installation may fail with the following error:

```
[INS-20702] Unexpected Internal driver error
```

Workaround:

Export the `OUI_ARGS` environment variable, before you run the SF Oracle RAC installation program:

```
export OUI_ARGS=-ignoreInternalDriverError
```

For more information, see the Oracle Metalink document: 970166.1

During installation or system startup, Oracle Grid Infrastructure may fail to start

After successful installation of Oracle RAC 11g Release 2 Grid Infrastructure, while executing the `root.sh` script, `ohasd` may fail to start. Similarly, during system startup, Oracle Grid Infrastructure may fail to start though the VCS engine logs may indicate that the `cssd` resource started Oracle Grid Infrastructure successfully.

The following message may be displayed on running the `strace` command:

```
# /usr/bin/strace -ftt -p pid_of_ohasd.bin
14:05:33.527288 open("/var/tmp/.oracle/npohasd",
O_WRONLY <unfinished ...>
```

For possible causes and workarounds, see the Oracle Metalink document: 1069182.1

Storage Foundation Oracle RAC issues

This section lists the known issues in SF Oracle RAC for this release.

CSSD configuration fails if OCR and voting disk volumes are located on Oracle ASM (3914497)

The Veritas installer fails to configure CSSD if OCR and voting disk volumes are located on Oracle ASM. This is because the installer does not support the configuration of CSSD with OCR and voting disk volumes on Oracle ASM.

Workaround: Configure the CSSD resource manually.

For instructions, see *Section: Installation and upgrade of Oracle RAC* in the *Storage Foundation for Oracle RAC Configuration and Upgrade Guide* document.

ASM disk groups configured with normal or high redundancy are dismounted if the CVM master panics due to network failure in FSS environment or if CVM I/O shipping is enabled (3600155)

Disk-level remote write operations are paused during reconfiguration for longer than the default ASM heartbeat I/O wait time in the following scenarios:

- CVM master node panics
- Private network failure

As a result, the ASM disk groups get dismounted.

Workaround: See to the Oracle metalink document: 1581684.1

PrivNIC and MultiPrivNIC agents not supported with Oracle RAC 11.2.0.2 and later versions

The PrivNIC and MultiPrivNIC agents are not supported with Oracle RAC 11.2.0.2 and later versions.

For more information, see the following Technote:

<http://www.veritas.com/docs/000010309>

CSSD agent forcibly stops Oracle Clusterware if Oracle Clusterware fails to respond (3352269)

On nodes with heavy load, the CSSD agent attempts to check the status of Oracle Clusterware till it reaches the `FaultOnMonitorTimeouts` value. However, Oracle Clusterware fails to respond and the CSSD agent forcibly stops Oracle Clusterware. To prevent the CSSD agent from forcibly stopping Oracle Clusterware, set the value of the `FaultOnMonitorTimeouts` attribute to 0 and use the

`AlertOnMonitorTimeouts` attribute as described in the following procedure.

Perform the following steps to prevent the CSSD agent from forcibly stopping Oracle Clusterware:

- 1 Change the permission on the VCS configuration file to read-write mode:

```
# haconf -makerw
```

- 2 Set the `AlertOnMonitorTimeouts` attribute value to 4 for the CSSD resource:

```
# hatype -display CSSD | grep AlertOnMonitorTimeouts
CSSD AlertOnMonitorTimeouts 0
# hares -override cssd_resname AlertOnMonitorTimeouts
# hatype -modify CSSD AlertOnMonitorTimeouts 4
```

- 3 Set the `FaultOnMonitorTimeouts` attribute value to 0 for the CSSD resource:

```
# hatype -display CSSD | grep FaultOnMonitorTimeouts
CSSD FaultOnMonitorTimeouts 4
# hares -override cssd_resname FaultOnMonitorTimeouts
# hatype -modify CSSD FaultOnMonitorTimeouts 0
```

- 4 Verify the `AlertOnMonitorTimeouts` and `FaultOnMonitorTimeouts` settings:

```
# hatype -display CSSD | egrep \
"AlertOnMonitorTimeouts|FaultOnMonitorTimeouts"
CSSD AlertOnMonitorTimeouts 4
CSSD FaultOnMonitorTimeouts 0
```

- 5 Change the permission on the VCS configuration file to read-only mode:

```
# haconf -dump -makero
```

Intelligent Monitoring Framework (IMF) entry point may fail when IMF detects resource state transition from online to offline for CSSD resource type (3287719)

When IMF detects a state transition from ONLINE to OFFLINE state for a registered online resource, it sends a notification to the CSSD agent. The CSSD agent schedules a monitor to confirm the state transition of the resource. The resources of type CSSD takes more time to go online or offline fully. Therefore, if this immediate monitor finds the resource still in online state, it assumes that the IMF notification is false and attempts to register the resource in online state again.

In such partial state transitions, the agent repeatedly attempts to register the resource until the `RegisterRetryLimit` is reached (default value is 3) or the resource registration is successful. After the resource is completely offline, the next resource registration with IMF will be successful.

Workaround: Increase the value of the `RegisterRetryLimit` attribute if multiple registration attempts fail.

The vxconfigd daemon fails to start after machine reboot (3566713)

The `shutdown -r` command makes sure that the file contents on the OS file system are written properly to the disk before a reboot. The `volboot` file is created in the OS file system, and is used to bring up the `vxconfigd` daemon after the system reboot. If the machine reboots for any reason without proper shutdown, and the `volboot` file contents are not flushed to the disk, `vxconfigd` will not start after the system reboots.

Workaround:

You must rerun the `vxinstall` script to re-create the `volboot` file and to start the `vxconfigd` daemon and other daemons.

Health check monitoring fails with policy-managed databases (3609349)

The health check option of the Cluster Server agent for Oracle fails to determine the status of the Oracle resource in policy-managed database environments. This is because the database SID is dynamically created during the time of the health check as a result of which the correct SID is not available to retrieve the resource status.

Issue with format of the last 8-bit number in private IP addresses (1164506)

The PrivNIC/MultiPrivNIC resources fault if the private IP addresses have a leading 0 in any of the octets that comprise the IP address, for example `X.X.X.01` or `X.X.0X.1` or `X.0X.X.1` or `0X.X.X.1`, where X is an octet of the IP address.

When you configure private IP addresses for Oracle Clusterware, ensure that the IP addresses have a format as displayed in the following two-node example:

- On galaxy: 192.168.12.1
- On nebula: 192.168.12.2

Confirm the correct format by viewing the PrivNIC or MultiPrivNIC resource in the `/etc/VRTSvcs/conf/config/main.cf` file.

CVMVolDg agent may fail to deport CVM disk group

The CVM disk group is deported based on the order in which the CVMVolDg resources are taken offline. If the CVMVolDg resources in the disk group contain a mixed setting of 1 and 0 for the `CVMDeportOnOffline` attribute, the disk group is deported only if the attribute value is 1 for the last CVMVolDg resource taken offline. If the attribute value is 0 for the last CVMVolDg resource taken offline, the disk group is not deported.

Workaround: If multiple CVMVolDg resources are configured for a shared disk group, set the value of the `CVMDeportOnOffline` attribute to 1 for all of the resources.

PrivNIC resource faults in IPMP environments on Solaris 11 systems (2838745)

The PrivNIC resource faults on Solaris 11 systems when private interfaces used by IPMP are configured under PrivNIC resource.

Workaround: Avoid using PrivNIC or MultiPrivNIC agents in IPMP environments.

Warning message displayed on taking cssd resource offline if LANG attribute is set to "eucJP" (2123122)

When you take the cssd resource offline using the `hares -offline cssd` command and the LANG attribute is set to "eucJP", the following message may be observed in the `hamsg engine_A` command output:

```
VCS INFO V-16-2-13716 Could not find message V-16-2-13716
```

You may ignore the message.

Error displayed on removal of VRTSjadba language package (2569224)

Removal of the VRTSjadba language package displays the following error on the screen:

```
Executing postremove script.
Generating BMC map file...
bmcmap ERROR V-33-1000-10001 Unable to create BMC map
```

You may ignore the error.

Veritas Volume Manager can not identify Oracle Automatic Storage Management (ASM) disks (2771637)

Veritas Volume Manager (VxVM) commands can not identify disks that are initialized by ASM. Administrators must use caution when using the VxVM commands to avoid accidental overwriting of the ASM disk data.

vxdisk resize from slave nodes fails with "Command is not supported for command shipping" error (3140314)

When running the `vxdisk resize` command from a slave node for a local disk, the command may fail with the following error message:

```
VxVM vxdisk ERROR V-5-1-15861 Command is not supported for command shipping.
Operation must be executed on master
```

Workaround: Switch the master to the node to which the disk is locally connected and run the `vxdisk resize` on that node.

Oracle Universal Installer fails to start on Solaris 11 systems (2784560)

The Oracle Universal Installer (OUI) fails to start when the SF Oracle RAC installer invokes the OUI for the installation of Oracle Clusterware/Grid Infrastructure software.

Workaround: Install the following packages before installing Oracle Clusterware/Grid Infrastructure.

```
SUNWxwplt
SUNWmfrun
```

For instructions, see the Oracle documentation.

CVM requires the T10 vendor provided ID to be unique (3191807)

For CVM to work, each physical disk should generate a unique identifier (UDID). The generation is based on the T10 vendor provided ID on SCSI-3 vendor product descriptor (VPD) page 0x83. In some cases, the T10 vendor provided ID on SCSI-3 VPD page 0x83 is the same for multiple devices, which violates the SCSI standards. CVM configurations should avoid using such disks.

You can identify the T10 vendor provided ID using the following command:

```
# sq_inq --page=0x83 /dev/diskname
```

On VxVM you can identify the T10 vendor provided ID using the following command:

```
# /etc/vx/diag.d/vxscsiinq -e 1 -p 0x83 /dev/vx/rdmp/diskname
```

You can verify the VxVM generated UDID on the disk using the following command:

```
# vxdisk list diskname | grep udid
```

FSS Disk group creation with 510 exported disks from master fails with Transaction locks timed out error (3311250)

Flexible Storage Sharing (FSS) Disk group creation for local disks that are exported may fail if the number of disks used for disk group creation is greater than 150, with the following error message:

```
VxVM vxdbg ERROR V-5-1-585 Disk group test_dg: cannot create: Transaction
locks timed out
```

A similar error can be seen while adding more than 150 locally exported disks (with vxdbg adddisk) to the FSS disk group, with the following error message:

```
VxVM vxdbg ERROR V-5-1-10127 associating disk-media emc0_0839 with emc0_0839:
Transaction locks timed out
```

Workaround:

Create an FSS disk group using 150 or less locally exported disks and then do an incremental disk addition to the disk group with 150 or less locally exported disks at a time.

vxdisk export operation fails if length of hostprefix and device name exceeds 30 characters (3543668)

If the combined length of the hostprefix and the device name exceeds 30 characters, the vxdisk export operation fails with the following error message:

```
VxVM vxdisk ERROR V-5-1-18318 Device c6t50060E8005655501d86s2: Name too
long for export. Length of Hostprefix + Disk accessname should not exceed
30 characters. Please see vxdctl(1M) man page for information on setting
user-specified hostprefix.
```

Workaround:

Use the enclosure-based naming (EBN) scheme instead of the operating system naming (OSN) scheme. OSN naming typically contains more characters and is not as intuitive. If the EBN name combined with the hostprefix exceeds 30 characters,

you can manually set the hostprefix to a smaller size using the `vxctl set hostprefix=value` command, where *value* is the new hostprefix.

Change in naming scheme is not reflected on nodes in an FSS environment (3589272)

In a Flexible Storage Sharing (FSS) environment, if you change the naming scheme on a node that has local disks, the remote disk names are not reflected with the corresponding name change. If you change the naming scheme on a node where exported disks are present, to reflect the updated remote disk names, you must either export the disks again or restart the node where the remote disks are present

Workaround:

There is no workaround for this issue.

When you upgrade SFRAC version from 6.2.1 to 7.2 on Solaris 11 Update 2, vxglm process fails to stop [3876778]

When you upgrade SFRAC version from 6.2.1 to 7.2 on Solaris 11 Update 2, SFRAC shutdown does not complete successfully as vxglm failed to stop.

Workaround: Restart the system to resolve the failures and then retry. If the issues persist after restart, contact Veritas Technical Support or refer to the *Installation Guide* for further troubleshooting.

Storage Foundation for Databases (SFDB) tools known issues

This section describes the known issues in this release of Storage Foundation for Databases (SFDB) tools.

Clone operations fail for instant mode snapshot (3916053)

For Oracle version 12.2.0.1.0, cloning a container database (CDB) fails for “instant mode” snapshots.

The cloning fails with the following error message:

```
SFDB vxsfadm ERROR V-81-0564 Oracle returned error.
Reason: ORA-01157: cannot identify/lock data file 5 - see DBWR trace file
ORA-01110: data file 5: '/data/DB12R2/pdbseed/system01.dbf'
```

Workaround: There is no workaround for this issue. Alternatively, you can use online or offline mode snapshots.

Sometimes SFDB may report the following error message: SFDB remote or privileged command error (2869262)

While using SFDB tools, if you attempt to run commands, such as `dbed_update` then you may observe the following error:

```
$ /opt/VRTSdbed/bin/dbed_update
No repository found for database faildb, creating new one.
SFDB vxsfadm ERROR V-81-0450 A remote or privileged command could not
be executed on swpa04
```

Reason: This can be caused by the host being unreachable or the `vxdbd` daemon not running on that host.

Action: Verify that the host `swpa04` is reachable. If it is, verify that the `vxdbd` daemon is running using the `/opt/VRTS/bin/vxdbdctrl` status command, and start it using the `/opt/VRTS/bin/vxdbdctrl start` command if it is not running.

Workaround: There is no workaround for this issue.

SFDB commands do not work in IPV6 environment (2619958)

In IPV6 environment, SFDB commands do not work for SF, SFCFSHA, SFHA or SFRAC.

Workaround:

There is no workaround at this point of time.

When you attempt to move all the extents of a table, the `dbdst_obj_move(1M)` command fails with an error (3260289)

When you attempt to move all the extents of a database table, which is spread across multiple mount-points in a single operation, the `dbdst_obj_move(1M)` command fails. The following error is reported:

```
bash-2.05b$ dbdst_obj_move -S sdb -H $ORACLE_HOME -t test3 -c MEDIUM
FSPPADM err : UX:vxfs fsppadm: WARNING: V-3-26543: File handling failure
on /snap_datadb/test03.dbf with message -
SFORA dst_obj_adm ERROR V-81-6414 Internal Error at fsppadm_err
```

Note: To determine if the table is spread across multiple mount-points, run the `dbdst_obj_view(1M)` command

Workaround: In the `dbdst_obj_move(1M)` command, specify the range of extents that belong to a common mount-point. Additionally, if your table is spread across "n" mount-points, then you need to run the `dbdst_obj_move(1M)` command "n" times with a different range of extents.

Attempt to use SmartTier commands fails (2332973)

The attempts to run SmartTier commands such as `dbdst_preset_policy` or `dbdst_file_move` fail with the following error:

```
fspadm: ERROR: V-3-26551: VxFS failure on low level mechanism
with message - Device or resource busy
```

This error occurs if a sub-file SmartTier command such as `dbdst_obj_move` has been previously run on the file system.

Workaround: There is no workaround for this issue. You cannot use file-based SmartTier and sub-file SmartTier simultaneously.

Attempt to use certain names for tiers results in error (2581390)

If you attempt to use certain names for tiers, the following error message is displayed:

```
SFORA dbdst_classify ERROR V-81-6107 Invalid Classname BALANCE
```

This error occurs because the following names are reserved and are not permitted as tier names for SmartTier:

- BALANCE
- CHECKPOINT
- METADATA

Workaround: Use a name for SmartTier classes that is not a reserved name.

Clone operation failure might leave clone database in unexpected state (2512664)

If the clone operation fails, it may leave the clone database in an unexpected state. Retrying the clone operation might not work.

Workaround:

If retrying does not work, perform one the following actions depending on the point-in-time copy method you are using:

- For FlashSnap, resync the snapshot and try the clone operation again.
- For FileSnap and Database Storage Checkpoint, destroy the clone and create the clone again.
- For space-optimized snapshots, destroy the snapshot and create a new snapshot.

Contact Veritas support if retrying using the workaround does not succeed.

Clone command fails if PFILE entries have their values spread across multiple lines (2844247)

If you have a parameter, such as `log_archive_dest_1`, in single line in the `init.ora` file, then `dbed_vmclonedb` works but `dbed_vmcloneb` fails if you put in multiple lines for parameter.

Workaround:Edit the PFILE to arrange the text so that the parameter values are on a single line. If the database uses a spfile and some parameter values are spread across multiple lines, then use the Oracle commands to edit the parameter values such as they fit in a single line.

Data population fails after datafile corruption, rollback, and restore of offline checkpoint (2869259)

Sometimes when a datafile gets corrupted below its reservation size, the rollback may not pass and the file may not be rolled back correctly.

There is no workaround at this point of time.

Flashsnap clone fails under some unusual archivelog configuration on RAC (2846399)

In a RAC environment, when using FlashSnap, the archive log destination to snapshot must be a shared path, and must be the same across all the nodes. Additionally, all nodes must use the same archive log configuration parameter to specify the archive log destination. Configurations similar to the following are not supported:

```
tpcc1.log_archive_dest_1='location=/tpcc_arch'
tpcc2.log_archive_dest_2='location=/tpcc_arch'
tpcc3.log_archive_dest_3='location=/tpcc_arch'
```

Where `tpcc1`, `tpcc2`, and `tpcc3` are the names of the RAC instances and `/tpcc_arch` is the shared archive log destination.

Workaround: To use FlashSnap, modify the above configuration to `*.log_archive_dest_1='location=/tpcc_arch'`. For example,

```
tpcc1.log_archive_dest_1='location=/tpcc_arch'
tpcc2.log_archive_dest_1='location=/tpcc_arch'
tpcc3.log_archive_dest_1='location=/tpcc_arch'
```

vxdbd process is online after Flash archive installation (2869269)

After a Flash archive installation of the SF stack, the `vxdbd` process is up, even if the stack is not configured.

Workaround: You can ignore, or stop the `vxdbd` process using the `/opt/VRTSdbed/common/bin/vxdbdctrl stop` command.

On Solaris 11.1 SPARC, setting up the user-authentication process using the `sfae_auth_op` command fails with an error message (3556996)

The debug logs display the missing `ps` utility as the 'ucb' package was absent in the default operating system installation. Due to which, the user-authentication process fails and the following error message is reported:

```
#/opt/VRTS/bin/sfae_auth_op -o setup
Setting up AT
Starting SFAE AT broker
```

```
SFDB vxsfadm ERROR V-81-0372 AT broker failed to start:
```

Workaround: Install the `pkg:/compatibility/ucb` package such that the `ps` utility is available in `/usr/ucb/ps`.

In the cloned database, the seed PDB remains in the mounted state (3599920)

In Oracle database version 12.1.0.2, when a container database (CDB) is cloned, the **PDB\$SEED** pluggable database (PDB) remains in the mounted state. This behavior is observed because of the missing datafiles in the cloned database for all point-in-time copies.

When you attempt to open the cloned seed database, the following error is reported:

```
"ORA-01173" oracle error.
...
```

```
SFDB vxsfadm ERROR V-81-0564 Oracle returned error.
```

```
Reason: ORA-01122: database file 15 failed verification check
ORA-01110: data file 15: '/tmp/test1/data/sfaedb/newtbs1.dbf'
ORA-01202: wrong incarnation of this file - wrong creation time
...
```

Workaround: There is no workaround for this issue.

Cloning of a container database may fail after a reverse resync commit operation is performed (3509778)

After a reverse resync operation is performed, the cloning of a container database may fail with the following error message:

```
SFDB vxsfadm ERROR V-81-0564 Oracle returned error.
```

```
Reason: ORA-01503: CREATE CONTROLFILE failed
ORA-01189: file is from a different RESETLOGS than previous files
ORA-01110: data file 6: '/tmp/testRecoverdb/data/sfaedb/users01.dbf'
```

Workaround: There is no workaround for this issue.

If one of the PDBs is in the read-write restricted state, then cloning of a CDB fails (3516634)

Cloning a container database (CDB) for point-in-time copies fails if some of the pluggable databases (PDBs) are open in the restricted mode. The failure occurs with the following error message:

```
SFDB vxsfadm ERROR V-81-0564 Oracle returned error.
```

```
Reason: ORA-65106: Pluggable database #3 (PDB1) is in an invalid state.
```

Workaround: There is no workaround for this issue.

Cloning of a CDB fails for point-in-time copies when one of the PDBs is in the read-only mode (3513432)

For Oracle version 12.1.0.1, cloning a container database (CDB) fails if one of the pluggable databases (PDBs) is in the read-only mode. The failure occurs with the following error message:

SFDB vxsfadm ERROR V-81-0564 Oracle returned error.

Reason: ORA-00376: file 9 cannot be read at this time
 ORA-01111: name for data file 9 is unknown - rename to correct file
 ORA-01110: data file 9: '/ora_base/db_home/dbs/MISSING00009'...

Workaround: There is no workaround for this issue.

If a CDB has a tablespace in the read-only mode, then the cloning fails (3512370)

For Oracle version 12.1.0.1, when a container database (CDB) has a tablespace in the read-only mode for all point-in-time copies, cloning of that CDB fails with the following error message:

SFDB vxsfadm ERROR V-81-0564 Oracle returned error.

Reason: ORA-01122: database file 15 failed verification check
 ORA-01110: data file 15: '/tmp/test1/data/sfaedb/newtbs1.dbf'
 ORA-01202: wrong incarnation of this file - wrong creation time
 ...

Workaround: There is no workaround for this issue.

If any SFDB installation with authentication setup is upgraded to 7.4, the commands fail with an error (3644030)

The commands fail with the error message similar to the following:

SFDB vxsfadm ERROR V-81-0450 A remote or privileged command could not be executed on prodhost

Reason: This can be caused by the host being unreachable or the vxdbd daemon not running on that host or because of insufficient privileges.

Action: Verify that the prodhost is reachable. If it is, verify that the vxdbd daemon is enabled and running using the [/opt/VRTS/bin/sfae_config status] command, and enable/start vxdbd using the [/opt/VRTS/bin/sfae_config enable] command if it is not enabled/running. Also make sure you are authorized to run SFAE commands if running in secure mode.

Workaround: Set up the authentication for SFDB again. See *Storage and Availability Management for Oracle Databases* or *Storage and Availability Management for DB2 Databases*.

Error message displayed when you use the `vxsfadm -a oracle -s filesnap -o destroyclone` command (3901533)

The `vxsfadm -a oracle -s filesnap -o destroyclone` command gives with the following error message:

```
Redundant argument in sprintf at
/opt/VRTSdbed/lib/perl/DBED/Msg.pm line 170.
```

Eg:

```
vxsfadm -s filesnap -a oracle -o destroyclone --name file1 --clone_name cln1
Redundant argument in sprintf at /opt/VRTSdbed/lib/perl/DBED/Msg.pm line 170.
Shutting down clone database... Done
Destroying clone... Done
```

Workaround: This message can be ignored. It does not affect the functionality in any manner.