

HCIE-Storage (Written)

Huawei H13-629

Version Demo

Total Demo Questions: 60

Total Premium Questions: 601

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Storage Technology Trends	6
Topic 2, Storage System Fundamentals	115
Topic 3, Storage System Construction	180
Topic 4, Storage System O&M	123
Topic 5, Storage System Disaster Recovery	177
Total	601

QUESTION NO: 1

A physical standby database is kept synchronized with the primary database, through Redo Apply, which recovers the redo data received from the primary database and applies the redo to the physical standby database. Which of the following disaster recovery technologies is working as the above statement?

- A. LogReplecation
- B. DataGuard
- C. GoldenGate
- D. VVR

ANSWER: B

Explanation:

DataGuard is correct because Oracle Data Guard is the disaster-recovery solution that maintains one or more standby databases as synchronized copies of a primary Oracle database. A physical standby database uses the same database structure as the primary database and receives redo generated by primary-database activity. Its Redo Apply service then uses Oracle recovery mechanisms to apply that redo to the physical standby database, keeping it transactionally consistent with the primary database.

This behavior precisely matches the statement's key terms: "physical standby database," transport of redo from the primary database, and "Redo Apply." Data Guard supports protection against site failures by enabling the standby database to take over the primary role when a switchover or failover is required. Depending on the configured protection mode and transport settings, redo can be shipped synchronously or asynchronously to balance data-protection requirements with application-performance and distance considerations.

Oracle documents that Redo Apply maintains a physical standby database by applying redo data received from the primary database. See the [Oracle Data Guard Concepts and Administration guide](#) and Oracle's [Data Guard overview](#).

QUESTION NO: 2

What are the recovery mode drills of the Huawei Replication Director disaster recovery?

(Multiple choice.)

- A. Planned Migration
- B. Failover
- C. Business Recovery
- D. fault recovery

ANSWER: A D

Explanation:

Planned Migration and **fault recovery** are the recovery-mode drills supported by Huawei ReplicationDirector disaster-recovery workflows. Planned migration is used when both production and recovery environments are available and the service movement is intentional. ReplicationDirector orchestrates the ordered transition of protected applications and storage replication so that services can be brought up at the target site in a controlled manner. This mode is particularly important for activities such as site maintenance, infrastructure replacement, and scheduled data-center migration.

Fault recovery is used to validate or perform recovery after a production-side fault. The workflow coordinates the recovery-side resources and application startup order, enabling protected services to be restored from the disaster-recovery copy. After the original production environment is available again, the recovery process can be followed by resynchronization and an orderly return of services as required by the protection design. These modes reflect ReplicationDirector's role as an orchestration layer for storage replication and application-consistent disaster recovery, rather than merely a mechanism for

creating copies. See Huawei's [data-protection product information](#) and the [Huawei Enterprise Support portal](#) for ReplicationDirector documentation and version-specific operation guides.

QUESTION NO: 3

Which of the following are Oracle11g logical backup methods? (Multiple choice)

- A. EXP
- B. RMAN
- C. EXPdp
- D. imp

ANSWER: A C

Explanation:

Oracle Database 11g logical backups are created by exporting database objects and their data into dump files. **EXP** is the original Export utility, commonly called traditional export. It extracts logical definitions and data for selected database objects, such as tables, schemas, or the full database, into an export dump file. This file represents a logical copy rather than a physical copy of database blocks or datafiles.

EXPdp is Oracle Data Pump Export, the newer export mechanism in Oracle 11g. It likewise creates logical dump files, while adding a server-based architecture, parallel execution capabilities, more granular object selection, and improved performance and manageability for many export operations. Oracle documents Data Pump Export as the utility used to unload data and metadata into a set of operating-system dump files. These exported logical backups can subsequently be loaded into an Oracle database through the corresponding import utilities. Oracle's documentation on the traditional Export utility and Data Pump Export is available at [Oracle Database 11g Original Export](#) and [Oracle Database 11g Data Pump Export](#).

QUESTION NO: 4

Production time is Monday to Friday 8:00 to 18:00. There are three LUNs, respectively LUN1, LUN2, LUN3, LUN1 with a size of 500G, all IO.

The average response time needs to be within 5ms, the size of LUN2 is 4TB, the coexistence of hot and cold data (with 12.5% of the hot data), and the size of LUN3 is 10TB. If performance requirements are not high and the engineer plans to use

SmartTier, which of the following description is correct? (Multiple choice)

- A. Use 3-tier Tier, 1TB high performance layer, 4TB performance layer, 10TB capacity layer, and configure LUN1 to allocate high-performance layer capacity first, LUN2 preferentially allocates performance layer capacity and LUN3 preferentially allocates capacity layer capacity
- B. Use 3-tier Tier, 1TB high-performance layer, 4TB performance layer, 10TB capacity layer, configure LUN1 and LUN2 to migrate automatically, and LUN3 priority allocation of low-performance layer capacity
- C. The use of 3-tier Tier, high-performance layer 1TB, performance layer is 4TB, capacity layer is 10TB, while configuring LUN1, LUN2, LUN automatic relocation. The LUN3 is migrated preferentially to the low-performance layer
- D. Set the monitoring period of the Smart tier to 8:00 to 18:00 every day from Monday to Friday. The migration period is not overlapped with the monitoring period

ANSWER: A D

Explanation:

Using 3-tier Tier, 1TB high performance layer, 4TB performance layer, 10TB capacity layer, and configuring LUN1 to allocate high-performance layer capacity first, LUN2 preferentially allocates performance layer capacity and LUN3 preferentially allocates capacity layer capacity is appropriate because the initial tiering policy matches each LUN's stated workload. LUN1 has a sustained I/O requirement and a response-time target, so placing its 500 GB capacity in the high-performance tier provides sufficient headroom while prioritizing low latency. LUN2 is 4 TB and contains 12.5% hot data; initially placing it in the performance tier supplies capacity suited to its overall workload while allowing SmartTier analysis and relocation policies to optimize data placement later. LUN3 is 10 TB with no demanding performance requirement, making the capacity tier the appropriate initial placement. This arrangement also aligns the configured tier capacities with the three LUN capacities and intended service levels.

Setting SmartTier monitoring to run from 08:00 to 18:00 on weekdays captures I/O behavior during the defined production period, so tiering decisions are based on representative workload statistics. Configuring migration for a period that does not overlap monitoring prevents relocation activity from interfering with workload observation and avoids introducing migration I/O during the production monitoring window. Huawei describes SmartTier as a policy-based feature that analyzes data access characteristics and relocates data among performance tiers; scheduling should therefore account for business workload periods. See Huawei's [storage documentation portal](#) and the [OceanStor Dorado product information](#).

QUESTION NO: 5

In the Huawei OceanStor storage array, the write cache high watermark is 80 and the low watermark is 20, and the service type is multiple-order small IO. To achieve better performance, it is advisable to adjust the water marks in which way?

- A. high water to 90
- B. high water to 70
- C. low water line 10
- D. low water line 30

ANSWER: B

Explanation:

For a workload characterized by multiple small write I/O operations, setting the write-cache high watermark to 70 is appropriate. The high watermark is the threshold at which the storage system begins more actively destaging dirty write-cache pages to disks. Reducing this threshold from 80 to 70 starts destaging earlier, leaving more cache capacity available to accept incoming bursts of small writes. This helps prevent write cache from becoming heavily occupied and reduces the likelihood that host writes must wait for cache space to be released.

This tuning is particularly useful for small-I/O workloads because they can create many cache pages and rapidly consume write-cache resources. Earlier, controlled destaging maintains a healthier balance between absorbing new host writes and flushing existing cached data, supporting steadier write performance. The low watermark can remain at 20, so the adjustment changes the point at which destaging is initiated without unnecessarily changing the cache's lower recovery threshold. Huawei documentation and product-support resources for OceanStor configuration and performance guidance are available through the [Huawei Enterprise Storage Support portal](#) and the [OceanStor documentation library](#).

QUESTION NO: 6

Client server resources are limited but there is insufficient network bandwidth. What kind of deduplication strategy should be used if the client data needs to be backed up to the physical tape?

- A. source-side deduplication (backup software)
- B. target-side deduplication (backup software deleted)
- C. Target-side de-duplication (backup software global deduction)
- D. target-side deduplication (backup media deduplication)

ANSWER: D

Explanation:

Target-side deduplication (backup media deduplication) is appropriate because deduplication processing is performed at the backup target rather than on the client server. When client CPU and memory resources are constrained, avoiding source-side hash calculation and duplicate-data lookup reduces the processing burden placed on production hosts. Since sufficient network bandwidth is available, transferring the data to the target before it is deduplicated does not create the usual bandwidth constraint that would favor source-side processing. The backup target can then identify duplicate data blocks and maintain the optimized backup representation while supporting the required tape-oriented backup workflow. This approach keeps client-side backup processing lightweight, centralizes deduplication management, and is well suited to environments in which the backup infrastructure has stronger processing capacity than the protected clients. Huawei's OceanProtect data-protection portfolio is designed around centralized backup storage and efficient data reduction capabilities; see the [Huawei OceanProtect product overview](#) and the [Huawei data protection solutions page](#).

QUESTION NO: 7

Which of the following information about the InfoEqualizer function of Huawei OceanStor 9000 is correct? (Multiple choice)

- A.** When using the InfoEqualizer function, you need to configure DNS. If you use an external DNS server, when configuring the DNS server, there is a setting of "reverse lookup zone name", network ID setting in this tab, you need to input network ID in reverse order.
- B.** When there is a DNS server in the user network, configure forwarding records on that DNS server. When the user network has no DNS server, configure the InfoEqualizer DNS IP address as the DNS server address on the client.
- C.** The root domain name is the domain name InfoEqualizer configured for the cluster. This domain name must be configured as the fully qualified domain name, that is, the domain name that includes the root node.
- D.** The domain name of the partition is configured as the relative domain name, and the root domain name of the InfoEqualizer is automatically taken as the suffix of the domain name, and the full name of the domain name is formed with maximum support of 32 partitions.

ANSWER: B C D

Explanation:

InfoEqualizer uses DNS-based name resolution to direct clients to appropriate OceanStor 9000 cluster resources. Where the customer network already has a DNS server, that DNS server must forward requests for the InfoEqualizer domain to the InfoEqualizer DNS service. If the customer network has no DNS server, clients use the InfoEqualizer DNS IP address directly. This ensures that client requests for the configured service names can be resolved correctly. The InfoEqualizer root domain must be entered as a fully qualified domain name so that it represents the complete DNS namespace used by the cluster. Partition domain names are then configured as relative names; InfoEqualizer appends the configured root domain to construct each partition's full domain name. This model provides a consistent namespace for access separation across partitions, and supports up to 32 partitions. Huawei identifies OceanStor 9000 as a scale-out storage platform and provides product and documentation resources through its [OceanStor 9000 product page](#) and [Enterprise Support portal](#). DNS forwarding and fully qualified domain naming are fundamental to making InfoEqualizer names resolvable from client networks.

QUESTION NO: 8

In the scenario of database protection, when the slave uses the virtual snapshot to verify LUN data consistency on the production end and the DR end and the disaster recovery end LUN database can not be started normally. Which of the following are the possible reasons for this? (multiple choice)

- A.** After a snapshot of the LUN of the redundancy site is created, the snapshot is not activated at the same time
- B.** Disaster recovery data stored in the Cache, not yet written to the hard disk
- C.** When creating a snapshot on the disaster recovery end LUN The production end is copying data to the disaster recovery end

D. The production end and disaster recovery end adopt the disaster recovery solution of incremental LUN copy. The incremental data is abnormally interrupted during the synchronization process and has not been recovered

ANSWER: A C D

Explanation:

“After a snapshot of the LUN of the redundancy site is created, the snapshot is not activated at the same time” is a possible cause because a virtual snapshot must be activated and mapped for host access before the database host can use its point-in-time data. Without activation, the expected snapshot-based verification volume is unavailable or cannot provide the intended data view.

“When creating a snapshot on the disaster recovery end LUN The production end is copying data to the disaster recovery end” can cause the database image to represent an application-inconsistent point in time. Database writes commonly span data files, logs, and metadata; taking the DR-side snapshot while replication is actively applying writes can capture a state that is not suitable for normal database startup unless an appropriate consistency mechanism is used.

“The production end and disaster recovery end adopt the disaster recovery solution of incremental LUN copy. The incremental data is abnormally interrupted during the synchronization process and has not been recovered” is also a valid cause. A recovery LUN that has not completed or recovered its required incremental synchronization may not contain a coherent recoverable database image. Huawei replication features rely on a defined synchronization and consistency state before a secondary copy is used for recovery or verification. See Huawei’s [HyperReplication Feature Guide](#) and [HyperSnap Feature Guide](#).

QUESTION NO: 9

A customer plans to deploy three application systems on an OceanStor V3 storage system. The first is a service system that carries Oracle databases, and it requires high random read and write performance. The second is a mail system used to process daily mails. The third is a video system used to archive video files. Which SmartTier policy settings are correct for creating LUNs for each application system? (Multiple Choice)

A. Service system:

1. Initial capacity allocation: Allocate from the performance layer first
2. Data relocation policy: Lowest relocation

B. Mail system:

1. Initial capacity allocation: Automatic allocation
2. Data relocation policy: Automatic relocation

C. Video system:

1. Initial capacity allocation: Allocate from the capacity layer first
2. Data relocation policy: Automatic relocation

D. Video system:

1. Initial capacity allocation: Allocate from the capacity layer first
2. Data relocation policy: Lowest relocation

ANSWER: A B C

Explanation:

SmartTier should be configured according to each workload’s initial performance requirement and the expected variation in data access frequency. For the service system, allocating capacity from the performance layer first places the Oracle database workload on the highest-performing media at creation time, which is appropriate for sustained random reads and writes. Using the lowest relocation setting reduces the priority and potential service impact of subsequent tier-relocation activity for this latency-sensitive workload. For the mail system, both automatic initial allocation and automatic relocation are appropriate because mail data commonly contains a mix of active and less-active content; SmartTier can place and adjust data according to observed access patterns. For the video archive system, allocating from the capacity layer first is appropriate because archived video is typically large and capacity-oriented. Automatic relocation still lets SmartTier identify any unexpectedly active video extents and move them to a faster tier when access statistics justify it, while cold archive extents remain on economical capacity media. Huawei describes SmartTier as using access-frequency information to

relocate data among storage tiers and optimize the balance of performance and capacity utilization. See Huawei's [OceanStor documentation portal](#) and [Huawei storage product documentation](#).

QUESTION NO: 10

For commonly used SQL Server 2005 and 2008, users can connect to the specified server through the SQL Server Management Studio tool Database Engine. There are two connections in the connection group authentication method. Which client application connection method can provide the required landing User ID and password?

- A. SQL Server Authentication
- B. Windows authentication
- C. Log in as superuser
- D. Other ways when landing

ANSWER: A

Explanation:

SQL Server Authentication is the connection method that requires the user to enter a SQL Server login name (User ID) and its associated password in the Database Engine connection dialog in SQL Server Management Studio. These credentials are maintained by SQL Server itself: a server-level SQL login is created, password policy settings can be applied, and the login is mapped to the appropriate database user and permissions. This permits a client to authenticate by explicitly supplying the SQL Server credentials when establishing a connection to the specified instance.

SQL Server supports this method as part of mixed-mode authentication, which allows SQL Server logins in addition to Windows-based identities. In the SSMS *Connect to Server* dialog, selecting SQL Server Authentication makes the Login and Password fields available, allowing the required credentials to be supplied directly. This is particularly applicable where the connecting user is not relying on an existing Windows domain or local Windows identity. Microsoft describes the server authentication modes and SQL Server logins in its [authentication mode documentation](#) and documents the SSMS Database Engine connection workflow at [Connect and query SQL Server using SSMS](#).

QUESTION NO: 11

Client 1 attempts to read file X from an OceanStor 9000, and node A responds to the read request. Which statements are true about the process of reading file x? (Choose two.)

- A. If the system detects that node B has cached file X, client 1 can directly read the file from node
- B. If the system detects that node B has cached file X, node A can directly read the file from node B.
- C. After processing the file read request initiated by client 1, the system retains file X on both node A and node B.
- D. If the system detects that node A has cached file X, only data of file X is cached on node A but the parity data is not cached on node A.

ANSWER: B D

Explanation:

In an OceanStor 9000 distributed file system, the node that receives a client read request acts as the service node and coordinates access to the requested file. When the required file data is already held in another node's cache, the service node can obtain that cached data directly from the node holding it through the cluster's internal high-speed communication mechanism. Therefore, "If the system detects that node B has cached file X, node A can directly read the file from node B." correctly describes a cache-assisted distributed read.

OceanStor 9000 cache handling distinguishes file data from RAID parity information. Caching file data improves subsequent read performance, while parity is protection metadata used for reconstruction and is not maintained as ordinary file-data cache on the service node. Accordingly, "If the system detects that node A has cached file X, only data of file X is cached on

node A but the parity data is not cached on node A.” is also correct. This design supports efficient distributed reads without treating parity blocks as client-readable cached content. Huawei’s OceanStor 9000 documentation resources are available through the [OceanStor 9000 product support page](#) and the [OceanStor 9000 documentation portal](#).

QUESTION NO: 12

The OceanStor 18000 system used by a client business application has a total of three pools, each of which is a 9-disk RAID5. Someday analyzing the log of business interruption, we found that there are some scenarios in the system that caused the failure. What will cause the host access to the stored business in the following scenario? (Multiple choice.)

- A. In the fully redundant group environment, occurred the master reset.
- B. With the BST (bad track) feature turned off, a BST bad track appears on the LUN.
- C. The POOL where the business is located is allthin LUN and not a snapshot. LUN copy, clone and copy services are configured in the log POOL, which is a space full of print information.
- D. Log in the following print:2014-01-19 08:04:40} {Receive a disk-fault event of the disk (328) of pool (2) from ctrl.}2014-01-19 08:40:35} {Receive a disk-fault event of the disk (613) of pool (3) from ctrl.}

ANSWER: B C

Explanation:

“With the BST (bad track) feature turned off, a BST bad track appears on the LUN” can interrupt host business access because the system is not using its bad-track handling mechanism to isolate or manage an affected media location. When host I/O reaches the bad track, the LUN can return I/O errors or experience failed requests, directly affecting the application’s ability to read or write its storage.

“The POOL where the business is located is allthin LUN and not a snapshot. LUN copy, clone and copy services are configured in the log POOL, which is a space full of print information” describes a thin-provisioned pool reaching full capacity. Thin LUNs consume physical pool capacity when new data is written. Once no allocatable capacity remains, subsequent writes requiring new extents cannot be satisfied. The host application can therefore encounter write failures and a business interruption even when snapshots, clones, and copy services are not consuming pool space. Capacity monitoring and timely pool expansion are consequently essential for thin-provisioned service continuity.

Huawei’s enterprise storage resources describe OceanStor storage management and capacity-provisioning functions: [Huawei Enterprise Storage Support](#) and [Huawei Data Storage Products](#).

QUESTION NO: 13

Which of the following statements about Huawei OceanStor V3 series storage array synchronous remote replication and asynchronous remote replication is incorrect? (Multiple choice)

- A. Synchronous remote replication and asynchronous remote replication both use differential bitmaps.
- B. Synchronous remote replication uses differential logging, whereas asynchronous remote replication uses differential bitmaps.
- C. Synchronous remote replication and asynchronous remote replication both use snapshots.
- D. Troubleshooting a replication task in a remote replication consistency group does not affect tasks in other groups.

ANSWER: A C

Explanation:

The statements “Synchronous remote replication and asynchronous remote replication both use differential bitmaps” and “Synchronous remote replication and asynchronous remote replication both use snapshots” are incorrect because the two replication modes use different mechanisms to maintain and resynchronize data. In synchronous remote replication, host write completion depends on the replication write being committed according to the synchronous replication policy. When the

replication link is unavailable, the system records changed data through differential logging so that the remote copy can be brought current after recovery. It is therefore inaccurate to characterize synchronous replication as using the same differential-bitmap mechanism as asynchronous replication.

Asynchronous remote replication operates in replication periods and identifies the changed data that must be transferred during those periods. Its implementation uses differential bitmaps and snapshot-based processing to establish consistent asynchronous synchronization points. A synchronous replication relationship is not based on periodic snapshot synchronization; it continuously mirrors writes during normal operation. Understanding this distinction is important when planning recovery-point objectives, link bandwidth, and behavior after a replication interruption. Huawei's storage documentation and product material describe remote replication as a data-protection capability with distinct synchronous and asynchronous modes; see the [Huawei technical documentation portal](#) and [Huawei enterprise storage products](#).

QUESTION NO: 14

OceanStor V3 does not start up normally within 15 minutes after the system is powered on. What can you use to log in to the storage system in order to locate this fault? (Multiple choice.)

- A. Sys status view system status
- B. Through the dev status view system status
- C. Through the dev status view device management started
- D. Through dev status view device management started

ANSWER: A C D

Explanation:

When an OceanStor V3 system does not complete startup within the expected time, service personnel can use the maintenance or serial-console login path to inspect the controller's startup state. **Sys status view system status** is applicable because the system-status output provides a direct indication of whether the storage system has reached its expected operating state and helps identify a startup stage at which progress has stopped. **Through dev status view device management started** is also applicable because confirming the DeviceManager startup state helps determine whether the management service has been initialized. This distinction is important during fault localization: a controller may be powered on while management components are still unavailable or have failed to start. The repeated option text, **Through dev status view device management started**, expresses the same valid diagnostic action and is therefore also correct. These checks support initial isolation of an abnormal boot issue before collecting logs or escalating the fault for deeper hardware or software analysis. Huawei's storage support resources and product documentation are available through the [Huawei Storage Support portal](#) and the [OceanStor documentation portal](#).

QUESTION NO: 15

You need to create a new OLAP backup database for batch import of online database data at night, and the database software is SQL server 2008. The database file size is 20TB and the storage array is OceanStor v3, which requires a stable write bandwidth of 1800MB / s.

What kind of allocation set is the most suitable to meet these requirements?

- A. 2TB NL-SAS hard drive, 8D + 2P RAID6 configuration
- B. With 2TB NL-SAS hard disk, configure 8D + 1P RAID5
- C. 600GB 15Krpm SAS hard drive, 8D + 1P RAID 5 configuration
- D. Configure raid 10 with 600GB 15rpm SAS hard disk

ANSWER: C

Explanation:

The **600GB 15Krpm SAS hard drive, 8D + 1P RAID 5 configuration** is the best balance of usable capacity, sustained write performance, and disk efficiency for this OLAP backup-import workload. A 20 TB database requires a substantial number of spindles, and 15,000-rpm SAS drives provide materially higher IOPS and sequential throughput than NL-SAS drives. This is important during the nightly batch-import window, where the specified 1,800 MB/s stable write target must be maintained rather than merely reached in a short burst.

An 8D + 1P RAID 5 layout provides eight data disks per RAID group and one parity disk, yielding efficient usable capacity while spreading sequential write activity across multiple high-performance SAS spindles. With enough such RAID groups to provide the required usable capacity, aggregate throughput can be scaled to meet the bandwidth objective. The layout is especially appropriate for large, predominantly sequential batch writes because full-stripe write behavior uses the RAID group efficiently.

This choice also avoids consuming the approximately 50% raw-capacity overhead associated with mirrored layouts, leaving more capacity available for the large backup database while retaining suitable performance. Huawei positions OceanStor systems for enterprise database workloads; see the [Huawei enterprise storage portfolio](#). Microsoft also documents the historical SQL Server 2008 platform lifecycle at [SQL Server 2008 end of support](#).

QUESTION NO: 16

When designing SQL server database disaster recovery, the calculation of disaster recovery capacity must consider which of the following file size? (multiple choice)

- A. control file
- B. data file
- C. log file
- D. configuration file

ANSWER: B C

Explanation:

For SQL Server disaster-recovery capacity planning, the protected database footprint includes both **data file** capacity and **log file** capacity. Data files contain the database's persisted objects and data pages, so the disaster-recovery target requires sufficient capacity to hold the database data set, including expected growth and the storage overhead required by the chosen replication, backup, or recovery design. SQL Server data files use the primary and secondary data-file types and commonly have .mdf and .ndf extensions.

Log files are also essential because SQL Server records transactions in the transaction log before data changes are committed. Recovery mechanisms use this log to maintain consistency and, where the recovery-point objective requires it, to restore or transmit changes made after an earlier full backup or replica synchronization point. Therefore, disaster-recovery sizing must account for transaction-log storage, log-backup retention or transfer requirements, and peak log generation. Microsoft's documentation describes database files and filegroups at [Database files and filegroups](#), and explains the transaction log's role in recovery at [The transaction log](#).

QUESTION NO: 17

The backup back-end capacity is derived from the initial capacity and data increment.

- A. True
- B. False

ANSWER: A

Explanation:

True is correct. Backup back-end capacity planning begins with the amount of source data that must be protected in the initial backup and then adds the changed-data volume generated over subsequent backup cycles. The initial capacity establishes the baseline amount of protected data, while the data increment represents the additional capacity consumed as changed blocks, files, or objects are retained according to the backup policy. In practice, the estimate is calculated across the required retention period: initial protected data plus accumulated incremental backup data. The final physical capacity selected can then be adjusted for the storage system's data-reduction effectiveness, such as deduplication and compression, as well as replication copies, reserved space, and the applicable protection scheme. This capacity-planning approach ensures that the backup repository can retain the baseline recovery point and the incremental recovery points needed to meet the specified recovery objectives. Huawei's OceanProtect backup-storage portfolio is designed for this type of capacity-efficient backup retention and data-reduction planning. See [Huawei OceanProtect](#) and the [Huawei Enterprise Support documentation portal](#) for product documentation and planning guidance.

QUESTION NO: 18

The OceanStor 18000 system used by a client business application has a total of three pools, each of which is a 9-disk RAID5. Someday analyzing the log of business interruption, we found that there are some scenarios in the system that caused the failure. What will cause the host access to the stored business in the following scenario? (Multiple choice.)

- A. In the fully redundant group environment, occurred the master reset.
- B. With the BST (bad track) feature turned off, a BST bad track appears on the LUN.
- C. The POOL where the business is located is all thin LUN and not a snapshot. LUN copy, clone and copy services are configured in the log POOL, which is a space full of print information.
- D. Log in the following print:
2014-01-19 08:04:40} {Receive a disk-fault event of the disk (328) of pool (2) from ctrl.}
2014-01-19 08:40:35} {Receive a disk-fault event of the disk (613) of pool (3) from ctrl.}

ANSWER: B C

Explanation:

"With the BST (bad track) feature turned off, a BST bad track appears on the LUN." is correct because BST is intended to isolate and manage media bad tracks so that affected areas can be handled without exposing an unrecoverable medium error to the host whenever possible. If this protection mechanism is disabled and a bad track is encountered while processing host I/O, the storage system can be unable to complete reads or writes for the affected logical address. This results in host-side I/O errors and can interrupt an application using that LUN.

"The POOL where the business is located is all thin LUN and not a snapshot. LUN copy, clone and copy services are configured in the log POOL, which is a space full of print information." is also correct in its intended meaning: a thin-provisioned LUN depends on sufficient physical capacity in its storage pool. When the pool is full, new allocation required by host writes cannot be satisfied. The array therefore reports write failures or makes the LUN unavailable for further writes, which can directly interrupt the dependent host service. Huawei documentation describes capacity-pool space monitoring and thin-LUN allocation as essential operational controls; pool capacity must be expanded or reclaimed before exhaustion affects services. See Huawei's [OceanStor product documentation](#) and the [Huawei Storage management documentation](#).

QUESTION NO: 19

When the database is protected and the virtual snapshots are used from the end to verify LUN data consistency on the production end and the disaster recovery end. However, the disaster recovery end of the LUNd database can not be started normally. What is the possible reason? (Multiple choice.)

- A. After a snapshot of the LUN of the redundancy site is created, the snapshot is not activated at the same time.
- B. Disaster recovery storage side of the database resides in the cache has not been written to the hard disk.
- C. When creating a snapshot on the disaster recovery end LUN, the production end is copying data to the disaster recovery end.

D. The production end and disaster recovery end adopt the disaster recovery solution of incremental LUN copy. The incremental data is abnormally interrupted during the synchronization process and has not been repaired.

ANSWER: A B D

Explanation:

“After a snapshot of the LUN of the redundancy site is created, the snapshot is not activated at the same time.” is a possible cause because a snapshot must be activated before its point-in-time data can be accessed and presented for validation or database use. The activation state is therefore essential when a snapshot is part of the recovery-side verification workflow.

“Disaster recovery storage side of the database resides in the cache has not been written to the hard disk.” is also a possible cause. A database requires a recoverable, write-consistent on-disk image, including its data and log-write ordering. If required writes remain pending in cache rather than being persistently committed in the recovery copy at the relevant consistency point, the recovered database image may not be startable.

“The production end and disaster recovery end adopt the disaster recovery solution of incremental LUN copy. The incremental data is abnormally interrupted during the synchronization process and has not been repaired.” is a valid cause because an interrupted incremental-copy chain can leave the recovery LUN without all changed blocks needed for a coherent database image. Recovery should use a verified synchronization point and repaired replication relationship before the target copy is used. Huawei’s storage documentation describes snapshot activation and replication-based protection concepts in its [OceanStor documentation](#) and the [Huawei Enterprise Storage support portal](#).

QUESTION NO: 20

Oracle database normally starts the opening state of the process. Which files must it read? (Multiple choice.)

- A. Parameter file
- B. Control Documents
- C. Archive log file
- D. Data files

ANSWER: A B D

Explanation:

During a normal Oracle startup through the OPEN state, Oracle reads the parameter file, control files, and data files as part of the successive startup phases. First, the instance is started in the NOMOUNT state. At this point, Oracle reads an initialization parameter file—either an SPFILE or a text PFILE—to obtain instance configuration settings such as memory and control-file locations. Next, when the database is mounted, Oracle opens and reads the control files. These files contain the database’s physical structure and checkpoint metadata, enabling Oracle to identify the database files that belong to the database. Finally, when the database is opened, Oracle opens the data files so that the database contents become available to users and applications. Therefore, *Parameter file*, *Control Documents*, and *Data files* are the required selections for the normal startup/open sequence described. Oracle’s administration documentation explains the NOMOUNT, MOUNT, and OPEN startup stages and the files used at each stage. See [Oracle Database Administrator’s Guide: Starting Up and Shutting Down](#) and [Oracle Database Concepts: Physical Database Structure](#).

QUESTION NO: 21

Engineers in strict accordance with the Oceanstor 9000 software installation and operation guide are to complete the pre-inspection and configuration, the software deployment process displayed ome exceptions, the following treatment is appropriate

(multiple choice)

- A. For the "configuration operating system failed" error, check the node / var / log partition and clean up the partition, initialize the system, and then re-run the software deploy

- B. If it is not a failure, restart the deployment after reinstalling the deploy deployment package
- C. After the deployment fails, run the `sh /opt/huawei/deploy/script/clean.sh` command on each node to initialize the system and restart the software deployment
- D. After deployment abnormalities, reboot all nodes, reinstall oceanstor toolkit, and start the software deployment

ANSWER: A C

Explanation:

For the “configuration operating system failed” error, checking the node’s `/var/log` partition is an appropriate first recovery action because deployment requires sufficient local filesystem space for operating-system configuration logs, temporary files, and deployment tasks. After the underlying partition-space issue is cleared, initializing the affected system state and rerunning deployment provides a clean and repeatable recovery path.

After a deployment failure, running `sh /opt/huawei/deploy/script/clean.sh` on every node is also appropriate when the installation procedure requires deployment-state cleanup before a retry. A failed distributed deployment can leave partial configuration, temporary deployment metadata, and inconsistent task state across nodes. Cleaning every participating node restores a consistent baseline, allowing the deployment software to execute its configuration workflow again without inheriting stale state from the failed run. These actions follow the general OceanStor 9000 deployment principle of resolving the reported prerequisite or environment fault first, then clearing failed deployment state before restarting the deployment process. Refer to the [Huawei OceanStor 9000 support page](#) and the [Huawei Enterprise Support documentation portal](#) for the applicable installation documentation and software package guidance.

QUESTION NO: 22

One customer purchased two Huawei OceanStor 9000 storage devices and deployed them in A and B cities respectively and configured the InfoReplicator function. Regarding production site A data reliability, which of the following descriptions is wrong?

- A. Due to system upgrades, you need to make a master-slave switchover for the remote replication pair that has been created, and the following conditions need to be met: 1. remote copy pair health status is normal, and the copy link is normal 2. from the directory has been canceled write protection
- B. If the replication link is down, you can only modify the pair attribute in the cluster where the home directory of production site A resides. After the replication link is restored, the modified pair attribute is automatically updated to the production site from the directory
- C. Set the channel authentication IP address, which is the static front-end service IP address of any node in the remote cluster, and is used to establish a replication area channel connection, this node must be in the copy area
- D. The channel authentication password can be modified. After the password is modified, if the copy link is disconnected, you need to manually re-enter the new password and keep the password. The password of the originating and receiving ends of the channel in the system is consistent so that the duplicated link can be recovered

ANSWER: D

Explanation:

The statement saying that the channel authentication password can be modified and then manually re-entered at both ends as a conditional recovery action is the wrong description. InfoReplicator replication-area channels use authenticated channel configuration to establish and maintain trusted communication between the two clusters. Channel authentication must be managed through the supported configuration and channel-management procedures so that the configured authentication information is valid for the established replication relationship. It is not a normal operational procedure to treat a replication-link interruption after a password change as a situation resolved simply by manually entering a password at the originating and receiving ends. This characterization is misleading for production-site reliability because recovery of remote replication communication must preserve the configured channel relationship and use the documented management workflow, rather than an ad hoc password re-entry process. Huawei documents OceanStor 9000 features and product documentation through its enterprise support portal, including the relevant InfoReplicator guidance: [Huawei OceanStor 9000 Support](#). Huawei’s storage documentation resources are also available from the [Huawei Enterprise Documentation portal](#).

QUESTION NO: 23

Simpana storage strategy is an important parameter is data retention rules to determine the need to back up the contents of the backup time and backup methods. Each copy has an independent data retention period, each copy can be set using_____. (Multiple choice.)

- A. Basic rules of retention
- B. Advanced retention rules
- C. Extended retention rules
- D. Special retention rules

ANSWER: A C

Explanation:

In Simpana (now Commvault), retention is configured independently for each storage-policy copy, allowing organizations to control how long data is preserved according to the purpose of that copy. **Basic rules of retention** are used for the standard retention policy of a copy. They define the principal retention thresholds, commonly based on a combination of retention time and the number of backup cycles, so that protected data remains available for the required operational period.

Extended retention rules supplement the basic policy by retaining selected backup jobs for longer periods. This supports longer-term retention schedules, such as preserving periodic full backups for monthly, quarterly, or annual requirements, while routine backups can age according to the basic retention settings. This per-copy design enables the same protected dataset to have different retention behavior in different copies, for example an operational copy and a compliance-oriented copy.

Commvault documentation describes storage policies and copies as the mechanisms used to control backup-data storage and retention. See the [Commvault Documentation portal](#) and the [Commvault data protection overview](#) for product concepts and current documentation access.

QUESTION NO: 24

In dual active disaster recovery solution architectures, the two FC switches will use the same domain ID after cascading to divide the zone.

- A. True
- B. False

ANSWER: B

Explanation:

False is correct. When Fibre Channel switches are cascaded, they form or extend a single SAN fabric, and each switch participating in that fabric must have a unique domain ID. The domain ID is part of the Fibre Channel addressing scheme: it contributes to the 24-bit FCID assigned to devices that log in through a switch. Unique domain IDs allow the fabric to route frames to the correct switch and preserve an unambiguous address space across the cascaded topology.

In a dual-active disaster-recovery architecture, zoning is configured based on the required host-to-storage connectivity and isolation policy, but cascading the switches does not make identical domain IDs appropriate. During fabric initialization and domain-ID negotiation, duplicate domain IDs create a conflict that must be resolved before the fabric can operate normally. Therefore, the two cascaded switches need distinct domain IDs, regardless of whether the deployment uses dual-active storage access and regardless of the zoning design. This follows standard Fibre Channel fabric behavior used by Huawei SAN deployments. See the Fibre Channel standards resources published by [INCITS Technical Committee T11](#) and Huawei's [Enterprise Support documentation portal](#) for SAN configuration documentation.

QUESTION NO: 25

An Oceanstor V3 system "alert" shows a message: "SPF [0] LOW Rx power alarm because Low Rx pow ALM of SFP0 is deleted/deflected. "The cause of the malfunction is:

- A. Equipment distribution voltage is insufficient, the device into a stream of electricity.
- B. The power module is not enough power to replace the power module.
- C. The voltage of the board is low or the light receiving power is low.
- D. Optical module port card loose connection.

ANSWER: C

Explanation:

The board voltage is low or the received optical power is low. This alarm is associated with the SFP transceiver's receive-side diagnostic monitoring. "LOW Rx power" indicates that the optical power arriving at SFP0 has fallen below the module's configured low-receive threshold, so the storage system raises an alarm to identify an abnormal optical-link condition. The immediate condition being reported is low received optical power; this can be assessed by checking the transceiver's diagnostic values and the physical optical path connected to the port. A low-voltage condition affecting the relevant board or transceiver power environment can also impair normal transceiver operation and is therefore included in the stated fault condition. Huawei OceanStor V3 maintenance documentation provides alarm handling and component-status information through the product documentation portal, including procedures for checking hardware and port health. See the [Huawei OceanStor V3 Series support page](#) and the [Huawei OceanStor V3 documentation](#). The appropriate response is to verify the reported receive-power diagnostic reading and ensure that the board and SFP have normal operating power before further link-level remediation.

QUESTION NO: 26

In order to verify the consistency of the slave LUN data with the source data and to protect slave LUN data, which of the following descriptions is incorrect:

- A. Create a virtual snapshot of the slave LUN and activate it, then map the snapshot LUN to the host for authentication.
- B. After the snapshot LUN is mapped to the host, the host can read and write the snapshot LUN.
- C. After the verification is completed, stop the upper-layer service application, release the LUN mapping, and finally delete the snapshot.
- D. Stop the snapshot before the host can write snapshot LUN.

ANSWER: D

Explanation:

"Stop the snapshot before the host can write snapshot LUN." is the incorrect description. A virtual snapshot must remain activated and available when it is mapped to a host for consistency verification. Activation makes the snapshot LUN accessible while preserving the point-in-time data represented by the snapshot. The host can then use the mapped snapshot LUN for the required verification activities without directly operating on the protected slave LUN.

Stopping a snapshot is not a prerequisite for host write access. On the contrary, stopping or deactivating the snapshot ends its active service state and prevents it from being used as the online mapped copy for the verification process. If a workflow requires writes to the snapshot LUN, those writes are performed while the relevant snapshot service and host mapping are active, subject to the storage system's snapshot capabilities and host access configuration. After verification, the mapping and snapshot can be removed in an orderly manner to release resources. Huawei's storage documentation provides the operational guidance for snapshot activation, host mapping, and lifecycle management; see the [Huawei Enterprise Documentation portal](#) and the [Huawei Enterprise Storage support page](#).

QUESTION NO: 27

The OceanStor 9000 log analysis tool is used to analyze the results of information collected and propose solutions to problems using on-site product maintenance tools. Which of the following statements are wrong regarding the use of log analysis tools? (Multiple choice)

- A.** Log analysis tools can run on any PC, do not need to interoperate with the storage system, the analysis process time according to the amount of the log may be, may it takes a long time
- B.** Log Analysis Tools After starting the analysis, the analysis tasks need to be stopped due to other reasons. Click on the analysis tool to abort this analysis
- C.** Engineer A deleted the original compression of the results of the collection, then the results will be collected to re-catalog the contents of all documents compressed into zip format for the dayChi analysis tools for analysis
- D.** After the analysis is completed, all the problems caused by the log analysis will be presented in the list of problems, and "View Details" contains the approximate cause of the problem and processing advice

ANSWER: C D

Explanation:

The statements beginning "Engineer A deleted the original compression of the results of the collection" and "After the analysis is completed, all the problems caused by the log analysis" are the incorrect statements. A collected log package is an input artifact produced by the prescribed collection procedure. If the original compressed collection result has been deleted, engineers should repeat collection using the supported maintenance workflow rather than manually reconstructing an archive by cataloging and compressing all files. Manual packaging can omit required metadata, alter the expected directory structure, or include files that are not part of the collection result, preventing reliable parsing by the analysis tool.

Log analysis also produces findings based on the signatures, rules, and information available in the collected logs; it does not guarantee that every possible product issue will be identified or that every issue has one definitive root cause and remediation. Results and detail views are diagnostic aids that must be assessed together with the service scenario, alarms, configuration, and engineer verification. Huawei's OceanStor 9000 support resources and documentation are available through the [OceanStor 9000 support page](#) and the [Huawei Enterprise Knowledge Center](#).

QUESTION NO: 28

A customer protects an Oracle database with asynchronous remote replication. The database data files, online redo logs, and control files have write dependencies. Which of the following practices are correct? (Multiple choice.)

- A.** Place write-dependent database LUNs in the same replication consistency group.
- B.** Design the asynchronous synchronization policy according to the business RPO.
- C.** Place data files and online redo logs in unrelated consistency groups to improve recoverability.
- D.** Ignore archived redo logs because they are never required during database recovery.

ANSWER: A B

Explanation:

LUNs containing Oracle data files, online redo logs, and control files that must be recovered together should be placed in the same replication consistency group. The consistency group coordinates synchronization of its member LUNs, producing a recovery point that preserves write-order consistency across the protected database components.

The synchronization policy must be designed from the required RPO. With asynchronous replication, changes can exist at the primary site that have not yet reached the secondary site; therefore, the selected interval, available bandwidth, and workload change rate must be assessed together. For a recoverable database DR design, relevant archived redo logs and recovery files must also be protected according to the recovery procedure. Oracle documents the role of redo and recovery files in the [Oracle Database Backup and Recovery User's Guide](#).

QUESTION NO: 29

What are the design information research activities in a backup solution? (Multiple choice.)

- A. Customer needs collection and refining
- B. Backup capacity calculation
- C. Project background research
- D. Network environment is collected

ANSWER: A C D

Explanation:

Customer needs collection and refining, project background research, and network environment is collected are design-information research activities because they establish the business and technical facts on which a backup design is based. Customer-needs collection identifies the protected services, recovery objectives, retention expectations, compliance requirements, service-level expectations, and operational responsibilities. Refining those needs turns broad requests into design inputs that can be validated with the customer.

Project-background research establishes the scope and context of the deployment, including the existing data-protection approach, current pain points, planned changes, implementation constraints, and stakeholder expectations. This context is needed to ensure that the proposed backup architecture fits the customer's environment and delivery objectives. Collecting the network environment is also essential discovery work: backup traffic paths, bandwidth, connectivity, routing, security boundaries, and available management networks directly affect the placement and communication design of backup components. Huawei's OceanProtect portfolio is designed around protecting data across heterogeneous production environments, making accurate workload and infrastructure discovery a prerequisite to solution design. General contingency-planning guidance likewise emphasizes understanding business processes, system dependencies, and infrastructure before selecting recovery capabilities. See [Huawei OceanProtect Data Protection](#) and [NIST SP 800-34 Rev. 1](#).

QUESTION NO: 30

Many "log file switch" wait events are observed during Oracle performance diagnosis, which ways can be used to optimize the performance? (Multiple Choice)

- A. Increase log cache.
- B. Migrate the log file to the high performance layer.
- C. Increase the number of log groups.
- D. Increase the concurrency and throughput for archiving the log file.

ANSWER: B C D

Explanation:

Frequent Oracle log-file-switch waits indicate that the database cannot reuse the next online redo log when it needs to switch. A practical response is to provide more online redo log groups, giving checkpoint processing and archive processing additional time before a previously used group must be reused. This directly reduces the likelihood that a switch stalls because a required operation has not finished.

Placing redo log files on a high-performance storage tier reduces redo-log write latency and helps checkpoint-related work complete promptly. Redo logs are write-intensive, sequential I/O files, so low-latency, high-IOPS storage is an appropriate placement choice when redo and switch activity is a bottleneck. Improving archive-log concurrency and throughput likewise lets ARCn processes finish archiving filled redo logs faster, preventing switches from waiting for a log group that is still needed for archival recovery.

Oracle recommends sizing and configuring redo logs and log groups to avoid unnecessary switches and ensuring that the I/O subsystem supporting redo processing is sufficiently performant. See Oracle's [Managing the Redo Log](#) documentation and its [wait-event troubleshooting guidance](#).

QUESTION NO: 31

A customer needs to create an OLAP backup database for batch importing online database data at night. The database software is SQL Server 2008. The data file size is 20 TB. The storage system is OceanStor V3. The stable write throughput is required to reach 1800 MB/s. Which configuration best meets customer requirements?

- A. Use 2 TB NL-SAS disks and configure RAID 6 (8D+2P).
- B. Use 2 TB NL-SAS disks and configure RAID 5 (8D+1P).
- C. Use 600 GB 15k rpm SAS disks and configure RAID 5 (8D+1P).
- D. Use 600 GB 15k rpm SAS disks and configure RAID 10.

ANSWER: C

Explanation:

Use 600 GB 15k rpm SAS disks and configure RAID 5 (8D+1P). This configuration provides the appropriate balance of usable capacity, sustained write performance, and protection for a large SQL Server OLAP backup database. The 15k rpm SAS tier is designed for substantially higher I/O capability and lower latency than capacity-oriented NL-SAS disks, which is important when batch imports must consistently sustain 1800 MB/s rather than merely reach a short burst rate. Multiple RAID 5 (8D+1P) disk groups can be created to supply the required 20 TB usable capacity and distribute sequential import writes across enough SAS spindles. In each disk group, eight data disks contribute capacity and aggregate data bandwidth, while one parity disk provides single-disk fault protection with better capacity efficiency than mirrored layouts. This makes the design well suited to the stated workload, which requires a high-throughput nightly loading window while retaining practical capacity efficiency for a 20 TB database. OceanStor V3 systems support disk-domain and RAID-group planning so capacity and performance can be scaled together as the number of disks and groups increases. See Huawei's [OceanStor V3 support documentation](#) and the [Huawei storage RAID configuration documentation](#).

QUESTION NO: 32

A customer uses asynchronous HyperReplication to protect an OceanStor production LUN at a remote site. Which of the following statements are correct? (Multiple choice.)

- A. The synchronization interval should be planned according to the required RPO and service change rate.
- B. An initial synchronization is required to establish the secondary copy.
- C. After a recoverable link interruption, synchronization can continue by transferring changed data when the link is restored.
- D. Every host write is acknowledged only after the remote LUN completes the write.

ANSWER: A B C

Explanation:

Asynchronous HyperReplication transfers changed data to the secondary LUN according to the configured synchronization policy. The replication interval and the available link bandwidth must be designed according to the business RPO and the expected change rate. An initial synchronization is required to establish the secondary copy before normal incremental synchronization can protect changed data.

During normal operation, the secondary LUN is a protected recovery copy and is not used for ordinary host read/write production service. If a replication link interruption occurs, the relationship retains the changed-data information needed for a subsequent incremental synchronization when connectivity is restored, provided that the relationship remains recoverable and sufficient resources are available. Unlike synchronous replication, asynchronous replication does not wait for remote-site

acknowledgement for every host write. Huawei feature documentation can be obtained through the [Huawei Support HyperReplication documentation search](#).

QUESTION NO: 33

Which of the following is a description of the WORM feature of OceanStor V3: (Multiple choice)

- A. When the file is locked, the file protection time can be extended or shortened by manual operation
- B. Locked file system WORM clock is greater than the file atime value, the file will be in an expired state
- C. When the file is in an append state, after the legal clock of the WORM file system exceeds the file expiration time, the file is moved to expired status
- D. When the size of the file is 0 bytes, the lock state can be migrated to the append state

ANSWER: B C D

Explanation:

OceanStor V3 WORM (Write Once Read Many) uses the file system's WORM clock and file timestamps to enforce retention and determine the lifecycle state of protected files. A locked file reaches the expired state when the WORM clock becomes greater than the file's atime value, which represents the applicable retention-expiration time in this state. This clock-based evaluation ensures that expiration is controlled by the storage system rather than by a client attempting to modify ordinary file metadata.

An append-state file can also transition to expired status once the WORM file system's legal clock exceeds the file expiration time. This preserves the WORM retention model while allowing the file to remain in its appendable lifecycle state until its configured retention period has ended. OceanStor WORM additionally permits a zero-byte file in the locked state to be migrated to the append state. This supports the defined state-transition behavior for an empty protected file without compromising retention controls for content that has already been written. Huawei describes OceanStor file-system data-protection capabilities through its [OceanStor storage product information](#) and publishes product documentation through the [Huawei Enterprise Support documentation portal](#).

QUESTION NO: 34

Which of the following statement regarding the Huawei disaster recovery solution on the implementation of disaster recovery exercise is wrong?

- A. The disaster recovery walkthrough automatically stops the production application and unproduces the production to production host mapping on the deviceManager
- B. The application of disaster recovery drill is mainly about the production system to switch to disaster recovery operation, exercise disaster recovery environment is correct
- C. After the disaster recovery drills are executed, the replication relationship will be split, and the disaster recovery end will take over the production system running service
- D. The Drill process can be normal in the production side to stop the business, a quick move to a disaster recovery operation

ANSWER: C

Explanation:

"After the disaster recovery drills are executed, the replication relationship will be split, and the disaster recovery end will take over the production system running service" is the wrong statement because a disaster recovery drill is intended to validate recoverability and the availability of the disaster-recovery environment without making that environment the live production service provider. The central purpose is to confirm that replicated data, host mappings, applications, and operating procedures can support recovery when required, while protecting the production workload from disruption. A drill may involve controlled preparation of secondary resources so that testing can occur, and the exact storage workflow

depends on the replication configuration and product version. However, the disaster-recovery site running a test workload is not equivalent to taking over the active production service. Production takeover is an actual disaster-recovery switchover or takeover operation, performed when production must be moved to the recovery site. That operation has different service-impact and replication-state implications from a drill. Therefore, describing the DR exercise as taking over the production system's running service incorrectly conflates a non-disruptive verification activity with an operational failover. Huawei's storage support resources provide product documentation for replication and DR procedures, including drill and recovery operations: [Huawei Enterprise Support](#) and [Huawei Storage products](#).

QUESTION NO: 35

In the scenario of protecting the database, the slave uses the virtual snapshot method to verify the services of the production end and the DR-side and discover the data of the DR on the disaster recovery end. What are the possible reasons that the library can not start normally? (Multiple choice.)

- A. After a snapshot of the LUN of the redundancy site is created, the snapshot is not activated at the same time.
- B. Disaster recovery side of the data stored in the cache, not yet written to the hard disk.
- C. When creating a snapshot on the disaster recovery end LUN, the production end is copying data to the disaster recovery end.
- D. The production end and disaster recovery end adopt the disaster recovery solution of incremental LUN copy, and the incremental data link is abnormal during the synchronization.

ANSWER: A C D

Explanation:

"After a snapshot of the LUN of the redundancy site is created, the snapshot is not activated at the same time" is a valid cause because the host must be mapped to, and able to access, the activated snapshot object before it can mount the database volumes and start database services. A snapshot that is created but not activated cannot provide the intended independent point-in-time data view to the verification host.

"When creating a snapshot on the disaster recovery end LUN, the production end is copying data to the disaster recovery end" can also prevent a normal database startup. The verification copy must represent a recoverable consistency point. If replication activity is occurring while the copy is being prepared, the resulting data image can lack the required application-consistent relationship among database data files, control files, and logs.

"The production end and disaster recovery end adopt the disaster recovery solution of incremental LUN copy, and the incremental data link is abnormal during the synchronization" is likewise a cause. An interrupted incremental-copy synchronization can leave the recovery-side LUNs without the complete set of changed blocks required for the selected recovery point, so the database may be unable to open normally. Huawei storage documentation describes snapshot-based data copies and replication features as mechanisms that require correct activation and a valid, consistent replication state for recovery use. See [Huawei Enterprise Support](#) and [Huawei Data Storage](#).

QUESTION NO: 36

Which of the following RAID levels can tolerate the simultaneous failure of two member disks in the same RAID group?

- A. RAID 0
- B. RAID 1
- C. RAID 5
- D. RAID 6

ANSWER: D

Explanation:

RAID 6 is correct. RAID 6 uses dual parity information and can tolerate two concurrent disk failures within the same RAID group. This provides stronger protection than RAID 5, which protects against only one disk failure.

The usable capacity and write overhead of RAID 6 must be considered during storage planning. RAID 10 can also provide high availability through mirroring, but its ability to tolerate multiple failures depends on which disks fail; two failures in the same mirror pair cause data loss. Huawei storage planning guidance is available through the [Huawei Enterprise Storage Support portal](#).

QUESTION NO: 37

On the command line of the N8500 (V200R001) clustered NAS system, you run the `NFS> share add rw, no_root_squas / vx / xscmfs` command. Which of the following is correct regarding this?

- A. The client can modify the file system when it accesses it
- B. It allows asynchronous writes to the file system and writes to the disk of the file system when the operation request is not responding
- C. It sets the client's root user does not have root privileges to access the NFS shared file system
- D. NFS shares created can be accessed by all clients

ANSWER: A D

Explanation:

The command creates an NFS share for the specified file-system path using the listed export-access attributes. The `rw` attribute grants read/write access, so an authorized NFS client mounting this export can create, modify, and delete files and directories in accordance with normal file-system permissions. Therefore, "The client can modify the file system when it accesses it" correctly describes the effect of the share configuration.

The `no_root_squash` attribute preserves the UID/GID identity of the client-side root user rather than mapping it to an anonymous identity. Although this setting affects privilege mapping, it does not itself limit the export to a listed host. Because no client-IP address, host name, or client group is supplied in this command, the share is not restricted to particular clients by the command; clients that can reach the NFS service may mount it, subject to system-level network and access controls. Thus, "NFS shares created can be accessed by all clients" is also correct in the context of this unrestricted share command. Huawei storage documentation is available through the [Huawei Storage Support portal](#); the standard semantics of NFS export access and root-squashing behavior are also documented in [exports\(5\)](#).

QUESTION NO: 38

In VSphere 5.0 and later, View Storage Accelerate (VSA) caches virtual machine disk data to EXSi hosts. It uses the content-based read cache (CBRC) in EXSi hosts, which no longer reads the entire operating system from the storage system. It reads the regular data block from the cache, thereby enhancing the performance of the virtual host.

- A. True
- B. False

ANSWER: A

Explanation:

True is correct. VMware View Storage Accelerator, introduced with vSphere 5.0, uses the Content-Based Read Cache (CBRC) capability on ESXi hosts to cache frequently used, identical disk blocks from virtual desktops. The cache is especially valuable where many linked-clone desktops share the same base operating-system image and therefore repeatedly request the same blocks during events such as desktop boot storms, login storms, antivirus scans, and patching activity. Instead of sending every repeated read request to the shared storage array, ESXi can satisfy requests for eligible cached blocks from its local host cache. This reduces read I/O load on the storage system and lowers storage-network traffic while improving desktop responsiveness and the host's ability to support concurrent virtual-desktop activity. The cache is

content based: blocks are identified by their contents, allowing common blocks across virtual machines to be cached once and reused. It does not eliminate storage access for all virtual-machine data, but it substantially reduces repeated reads of common operating-system and application blocks, which is the performance benefit described in the statement. VMware's feature documentation identifies View Storage Accelerator as using CBRC on ESXi hosts; see [VMware Horizon View documentation](#) and [Broadcom vSphere storage documentation](#).

QUESTION NO: 39

When the storage controller BBU is damaged, the controller automatically sets the IO write policy to transparent and is locked to ensure data reliability. Before the fault is recovered, it can be forced to write back.

- A. True
- B. False

ANSWER: B

Explanation:

False is correct. The BBU is a cache-protection component: it provides the power needed to preserve controller cache contents if external power fails. When the BBU is damaged or reports an abnormal condition, the controller cannot safely rely on volatile write cache. To protect host data integrity, the storage system changes the write behavior to write-through (the question's "transparent" policy) and locks the protection-related policy. Write I/O is then committed to persistent media before completion is acknowledged to the host.

This safeguard must remain in effect until the BBU fault is cleared and the cache-protection capability has been restored and verified. A forced return to write-back while cache protection is unavailable would allow acknowledged writes to reside only in volatile cache; an unexpected power interruption could therefore lose those writes. The controller's protective lock is specifically intended to prevent that condition. Huawei's enterprise storage documentation and support resources describe cache protection and controller-component maintenance principles: [Huawei Enterprise Support](#) and [Huawei Enterprise Storage](#).

QUESTION NO: 40

Which of the following statements about the Oceanstor V3 file system deduplication features are correct? (multiple choice)

- A. can only be opened when creating a file object, business can not be opened or closed during data deduplication and compression function
- B. Using online processing to deduplicate data before it is written to disk
- C. When both deduplication and compression are turned on, the data will be deduplicated and then compressed
- D. Use smart cache can accelerate the search performance of fingerprints deleted

ANSWER: B C D

Explanation:

OceanStor V3 file-system data reduction uses inline processing: duplicate-content detection is performed as data is written, so redundant blocks can be identified before they consume backend disk capacity. This approach reduces physical capacity usage while preserving the file system's logical view of the data. Where both functions are enabled, the system performs deduplication before compression. Deduplicating first avoids compressing multiple identical copies of the same data and allows compression to operate only on the remaining unique data, improving overall reduction efficiency. Deduplication depends on fingerprint indexes to locate previously stored matching data. SmartCache, using SSD resources to accelerate hot metadata and cache access, can improve the efficiency of fingerprint lookup operations and therefore benefit deduplication processing performance. These capabilities align with OceanStor V3's file-service data-reduction design, in which inline deduplication, compression sequencing, and SSD-assisted cache acceleration work together to reduce capacity consumption without changing application access to files. Huawei describes OceanStor storage software capabilities and

file-service efficiency features at [Huawei OceanStor Dorado](#) and provides product documentation through the [Huawei Enterprise Storage Support](#) portal.

QUESTION NO: 41

For commonly used SQL server 2005 and 2008, users can connect to a specific service through the SQL server management studio tool database engine. T login there are two ways to authenticate. Which mode is needed for the client application connection?

- A. sql server authentication
- B. Windows authentication
- C. Log in as super user
- D. Other ways when landing

ANSWER: A

Explanation:

sql server authentication is the expected mode for a client application connection in this context. SQL Server Authentication uses a login name and password maintained by SQL Server itself, allowing an application to connect by supplying database credentials in its configured connection information. This is particularly suitable where the application must use a dedicated database login, where the client is not operating under a trusted Windows domain identity, or where the database service needs an account independent of the interactive Windows user.

For SQL Server 2005 and 2008, enabling this type of login requires the server to be configured for mixed authentication mode, which supports both Windows identities and SQL Server-managed logins. After the server setting is enabled, an administrator creates the required SQL Server login, maps it to the relevant database user, and grants only the permissions needed by the application. Microsoft's SQL Server documentation describes SQL Server Authentication as authentication based on a login ID and password supplied to SQL Server, and its authentication-mode documentation explains mixed mode support: [Choose an authentication mode](#) and [Create a login](#).

QUESTION NO: 42

A computer center plans to purchase a set of Huawei Oceanstor 9000 products as a follow-up storage system. The requirements are as follows:

- (1) The minimum storage system needs to provide a total bandwidth of 3-4GB / s
- (2) storage space can reach 1PB, to meet a single file GB storage
- (3) computer center room requires control deployed in a cabinet use
- (4) Internal computer cluster nodes are now all InfiniBand network

Based on the above information, which of the following configuration suggestions do you think is more suitable for this project?

- A. Because computing tasks are designed for large file services, it is more appropriate to deploy with an 8-node OceanStor 9000 C node. The capacity on the use of 4TB SATA, to meet customer needs
- B. Using Oceanstor 9000 C Node combined with the total bandwidth and total capacity requirements, consider the 8-node configuration 4TB SATA full disk, backend 10GETOE + front-end IB networking
- C. From the reliability point of view, it is recommended to use theInfoEqualizer function to access the OceanStor 9000 storage through a domain name. The InfoEqualizer DNS segment needs to be considered and OceanStor 9000 back-end to interoperability
- D. OceanStor 9000 supports front-end 10GE TOE + back-end IB networking,the front-end is easy to log in and manage with the management PC interworking, considering this networking solution can reduce the cost of front-end IB switches.

ANSWER: B

Explanation:

Using Oceanstor 9000 C Node combined with the total bandwidth and total capacity requirements, consider the 8-node configuration 4TB SATA full disk, backend 10GETOE + front-end IB networking is the suitable recommendation because it matches both the required scale and the existing compute-network environment. An eight-node C-node configuration populated with 4 TB SATA disks provides the required petabyte-class raw capacity while retaining the scale-out performance needed for an aggregate throughput target of approximately 3–4 GB/s. The C-node form factor is also intended for high-density, rack-based deployment, which fits the computer-room constraint. Most importantly, the client-facing network is InfiniBand, allowing the OceanStor 9000 front end to connect directly to the existing InfiniBand-based compute cluster. This avoids introducing an unnecessary protocol or network transition between the cluster and storage service. Using 10GE TOE for the back-end network provides the node-to-node storage communication design associated with this deployment approach, while InfiniBand supplies the high-throughput front-end access path for the compute workload. Huawei's OceanStor 9000 is a scale-out storage platform designed to expand capacity and performance by adding nodes; see the [Huawei scale-out storage overview](#) and [Huawei Enterprise Storage support portal](#).

QUESTION NO: 43

Two Huawei N8500 standard nodes opened the CIFS service. Which of the following CIFS parameters allows users to view the other users without seeing the head record?

- A. oplocks
- B. noguest
- C. fs_mode
- D. hide_unreadable

ANSWER: D

Explanation:

The correct setting is **hide_unreadable**. In CIFS/SMB service configuration, this parameter controls whether directory entries for objects that a user is not permitted to read are displayed in directory listings. When it is enabled, the CIFS server evaluates the requesting user's permissions while generating the listing and suppresses entries that the user cannot read. Thus, users can browse the shared directory and see the files or folders available to them without seeing entries for protected content. This behavior is commonly used where a single share serves multiple users or groups but each user must have a restricted view based on the access-control permissions already applied to files and directories. The setting affects visibility in directory enumeration; it relies on the underlying permissions to determine which entries are unreadable for that particular user. The CIFS/SMB implementation documentation describes `hide_unreadable` as hiding files and directories from users who lack read access, which matches the required result. See the [Samba smb.conf parameter reference](#) and Huawei's [OceanStor documentation portal](#) for CIFS service and permission-configuration guidance.

QUESTION NO: 44

A Windows client accesses an OceanStor 9000 file system by using SMB. The client has multiple network adapters, and the storage system also provides multiple SMB network paths. Which SMB feature can use these paths concurrently to improve throughput and provide network fault tolerance?

- A. SMB Multichannel
- B. SMB 1.0 browser service
- C. NFS lock daemon
- D. FTP passive mode

ANSWER: A

Explanation:

SMB Multichannel is correct. SMB Multichannel is an SMB 3 capability that enables an SMB client to establish multiple connections for an SMB session when multiple suitable network interfaces are available. It can aggregate available bandwidth and continue service through remaining connections if one network path fails.

SMB Multichannel requires compatible SMB clients and servers and appropriate network planning. It is not a replacement for storage-side RAID protection or remote replication; it improves the client-to-file-service network path. Microsoft documents the feature in [SMB Multichannel](#).

QUESTION NO: 45

A Linux client mounts an NFS file system exported by an OceanStor storage system. Which of the following statements about NFS mount options are correct? (Multiple choice.)

- A. `rsz` and `wsz` control the maximum NFS read and write request sizes.
- B. The `hard` option continues retrying NFS requests after a temporary server or network interruption.
- C. The `noatime` option can reduce access-time metadata updates.
- D. The `nolock` option is recommended whenever multiple clients need coordinated file locking.

ANSWER: A B C

Explanation:

The `rsz` and `wsz` options specify the maximum data size used by the NFS client for read and write requests. Values supported by the client, server, protocol version, and network can improve throughput by reducing the number of NFS operations needed for large sequential transfers.

The `hard` option causes the NFS client to continue retrying requests after a server or network interruption. This behavior is generally appropriate for data requiring reliable access because applications do not receive an I/O error merely because a temporary NFS connectivity failure occurs. The `noatime` option avoids updating file access-time metadata when that information is not required, which can reduce metadata-write activity. Linux mount-option behavior is described in the [nfs\(5\) manual page](#).

QUESTION NO: 46

True or false, with RedHat 6 using the storage array Huawei S5500T, when the storage array and the host are directly connected, the storage array host port can select a P2P network topology.

- A. Correct
- B. mistakes

ANSWER: A

Explanation:

Correct is the right response. A point-to-point (P2P) Fibre Channel topology is specifically intended for a direct connection between two Fibre Channel ports: in this case, an HBA port on the Red Hat Enterprise Linux 6 host and a host-side Fibre Channel port on the Huawei S5500T. Since there is no intervening Fibre Channel switch or additional device participating in the path, the two ports establish a dedicated FC link, which is the normal P2P use case. The storage port topology must be configured consistently with the physical connection so that link initialization and host discovery can complete correctly. Red Hat Enterprise Linux includes Fibre Channel support through its supported HBA drivers and the Linux SCSI stack; once the direct FC link is established, the host can discover and use the LUNs mapped by the storage system. Red Hat's storage documentation describes Fibre Channel as a supported storage transport and covers its host-side operation in RHEL 6. Huawei's S5500T product documentation and support resources provide the applicable storage configuration guidance. See [Red Hat Enterprise Linux 6 Fibre Channel documentation](#) and [Huawei OceanStor S5500T support resources](#).

QUESTION NO: 47

In Huawei FusionSphere, which are typical application scenarios of the snapshot function?

(Multiple Choice)

- A. Data backup
- B. Data restoration
- C. Data copy
- D. Data deletion
- E. Data backup, data restoration, and data copy

ANSWER: E

Explanation:

Data backup, data restoration, and data copy are all typical snapshot application scenarios in Huawei FusionSphere. A snapshot preserves the state of a virtual machine disk or storage volume at a specific point in time while avoiding a full immediate copy of all source data. This point-in-time baseline can support backup workflows by providing a consistent source from which protected data can be retained or exported. It also supports restoration by allowing data to be returned to the captured state when recovery is required. In addition, snapshots can be used as the basis for data copies, such as creating an isolated copy for testing, development, analysis, or other non-production use without disrupting the active production workload. These related uses reflect the core value of snapshot technology: rapidly preserving a recoverable point in time and reusing that point for protection and copy-based workflows. Huawei documentation for FusionSphere storage and data-protection capabilities describes snapshots as a mechanism for point-in-time data protection and recovery; see [Huawei FusionSphere documentation](#) and [Huawei storage documentation](#).

QUESTION NO: 48

VIPs in Oracle RAC systems are created by VIPCA scripts during the final stages of installing clusterware and are registered as CRS resources and bind to the public network card for each node.

- A. True
- B. False

ANSWER: A

Explanation:

True is correct. In traditional Oracle RAC and Oracle Clusterware installations, the Virtual IP Configuration Assistant (VIPCA) is run as part of the Clusterware configuration process to configure a virtual IP address for each cluster node. VIPCA creates the VIP configuration and registers the associated VIP resources with Oracle Clusterware (CRS), allowing Clusterware to monitor, start, stop, and fail over those addresses when appropriate. A node's VIP is associated with that node's public network, rather than the RAC private interconnect. This design gives clients a quickly detectable connection failure if a node becomes unavailable, rather than requiring them to wait for normal TCP timeout behavior. When the node is healthy, its VIP is hosted through the node's public network interface; Clusterware controls the resource lifecycle and can relocate it as required by the configured failover behavior. The terminology in the statement reflects the pre-SCAN RAC architecture and installation procedures commonly documented for earlier Oracle RAC releases, in which VIPCA was explicitly invoked during Clusterware setup. Oracle's RAC documentation describes VIP addresses as Clusterware-managed resources used for client connectivity and fast failure detection. See the [Oracle Real Application Clusters Administration and Deployment Guide](#) and the [Oracle RAC Administration Guide](#).

QUESTION NO: 49

Which of the following descriptions about the InfoEqualizer features of the OceanStor 9000 storage system are incorrect. (Multiple Selections)

- A. Dynamic domain name, which returns the dynamic IP address of the node for the client, supports the failover function and is recommended for CIFS protocol sharing
- B. Static domain name, which returns the domain name of the node's static external IP address to clients, is recommended for NFS protocol sharing
- C. dynamic external IP address, node failure, the dynamic IP address can be automatically assigned to other working nodes, the client can pass again This IP address is connected to the OceanStor 9000
- D. static external IP address, the node fails, the client can not access the oceanstor 9000 through the IP address

ANSWER: A B

Explanation:

The incorrect descriptions are “Dynamic domain name, which returns the dynamic IP address of the node for the client, supports the failover function and is recommended for CIFS protocol sharing” and “Static domain name, which returns the domain name of the node's static external IP address to clients, is recommended for NFS protocol sharing.” The protocol recommendations in these descriptions are reversed. InfoEqualizer’s dynamic domain-name mechanism is intended for NFS sharing, allowing a client to obtain a dynamic external IP address and retain service availability when that address is taken over following a node fault. Static domain names resolve to static external IP addresses and are intended for CIFS sharing. This distinction reflects the different client-access and connection characteristics of NFS and CIFS in an OceanStor 9000 scale-out NAS deployment. Correctly selecting the access type ensures that DNS resolution, external-IP allocation, and failover behavior align with the selected NAS protocol. Huawei’s OceanStor 9000 product information describes the platform as a scale-out storage system providing NAS services and high availability; its documentation resources are available through the [Huawei OceanStor 9000 product page](#) and the [Huawei Enterprise Support documentation portal for OceanStor 9000](#).

QUESTION NO: 50

Simpana storage strategy is an important parameter is data retention rules to determine the need to back up the contents of the backup time and backup methods. Each copy has an independent data retention period, each copy can be set using_____. (Multiple choice.)

- A. Basic rules of retention
- B. Advanced retention rules
- C. Extended retention rules
- D. Special retention rules

ANSWER: A B

Explanation:

Basic rules of retention and **Advanced retention rules** are the applicable retention-setting methods for an individual Simpana (now Commvault) storage-policy copy. A copy's basic retention configuration establishes the normal aging threshold for its protected backup data. This baseline policy controls how long data remains available according to retention criteria such as time and backup-cycle requirements, allowing each copy to have a retention period appropriate to its operational purpose.

Advanced retention rules supplement that baseline configuration when the organization needs retention behavior beyond the standard copy-level period. These settings support more granular, policy-driven preservation of qualifying backup data, enabling longer-term retention requirements to be applied without changing the fundamental retention objective of every backup in the copy. This is particularly useful where a copy serves archive, compliance, or other long-duration recovery requirements while retaining independent retention control from other copies in the same storage policy.

Commvault documentation describes retention and data-aging management as storage-policy copy functions and documents the available retention configuration concepts in its official documentation portal. See the [Commvault Documentation portal](#) and Commvault's [Data Protection product documentation overview](#).

QUESTION NO: 51

Simpana index is divided into two sections, which used to record the information of each data object and storage location, the index of a large amount of information is managed by what?

- A. CommServe
- B. MediaAgent
- C. Library
- D. iDataAgent

ANSWER: B

Explanation:

MediaAgent is correct because, in the Simpana (now Commvault) architecture, the MediaAgent is responsible for managing the index cache associated with protected data. The index records the relationship between backed-up data objects and their locations on storage media, enabling the system to identify the required backup content and efficiently retrieve it during browse, search, and restore operations. The MediaAgent performs the data-movement function between clients and storage targets and maintains the relevant indexing information needed to locate data within backup media. This design keeps the operational, storage-oriented index processing close to the component that writes to and reads from the backup storage. The CommServe centrally coordinates the environment and maintains configuration and job-management information, while the MediaAgent handles the large-volume index-cache management required for locating protected data. Commvault's product documentation describes indexing as the mechanism that enables data discovery and recovery and documents the MediaAgent role within the data-protection architecture. See the [Commvault Documentation portal](#) and the [Commvault indexing documentation](#).

QUESTION NO: 52

SQL Server database administrator found that SQL Server is running out of space, you want to delete unnecessary databases to get more space. Which of the following databases will be deleted once the database engine is abnormal or the upper business interruption?

- A. pubs
- B. master
- C. msdb
- D. tempdb

ANSWER: B C D

Explanation:

The correct selections are **master**, **msdb**, and **tempdb**, because they are SQL Server system databases with essential instance or workload functions. The **master** database holds instance-wide configuration information and records information about all databases on the instance, including their files and locations. Its availability is fundamental to starting and operating the SQL Server instance, so removing or damaging it can make the database engine unavailable.

The **msdb** database supports operational services such as SQL Server Agent jobs, schedules, alerts, backup and restore history, Database Mail, and other management metadata. Removing it can interrupt scheduled processing, backups, maintenance, notifications, and other business-dependent operations. The **tempdb** database is required for temporary objects, row versioning workloads, sorts, hashes, and other internal operations. Although SQL Server recreates tempdb

during startup, it is required while the instance is operating; it cannot simply be dropped as a space-reclamation measure without disrupting database activity. Microsoft classifies all three as system databases and recommends protecting and managing them rather than deleting them. See [SQL Server system databases](#) and [tempdb database documentation](#).

QUESTION NO: 53

In the same operating system environment, what is the order to restore the oracle database through RMAN?

- A. controlfile, datafile, pfile
- B. pfile, datafile, controlfile
- C. datafile, pfile, controlfile
- D. pfile, controlfile, datafile

ANSWER: D

Explanation:

The correct restoration order is **pfile, controlfile, datafile**. Oracle must first have an initialization parameter file (PFILE), or an equivalent server parameter file, so that the instance can be started in the `NOMOUNT` state. At this stage Oracle reads instance configuration parameters such as memory settings, database name, and control-file locations; no control file is yet required to start the instance this way.

Once the instance is running in `NOMOUNT`, RMAN can restore the control file from its backup. The control file contains the database's physical structure and backup metadata required to identify data files and proceed with recovery. After restoring it, the database is mounted, which makes the restored control file available to Oracle. RMAN can then restore the data files to their required locations. In a complete recovery workflow, the restored data files are subsequently recovered by applying redo, and the database is opened when recovery has completed, commonly with `RESETLOGS` when required by the recovery scenario.

Oracle documents this sequence as starting the instance with a parameter file, restoring the control file, mounting the database, and then restoring database files. See the [Oracle RMAN recovery documentation](#) and the [Oracle Database Backup and Recovery User's Guide](#).

QUESTION NO: 54

The user backup client has sufficient resources and the production network resources are relatively tight. The amount of data scheduled to be backed up weekly is 1-2T. The backup window is 3 hours, using source-side de-duplication, deduplication ratio of 3, the bandwidth utilization value of 0.8. Based on this, what kind of networking can meet the backup business needs and investment least?

- A. GE LAN-free
- B. 10GE LAN-free
- C. FC LAN-free
- D. server-free

ANSWER: A

Explanation:

GE LAN-free meets the requirement with the lowest investment. Capacity planning should use the upper estimate of 2 TB. With source-side deduplication at a ratio of 3:1, the data transmitted during the backup is approximately $2 \text{ TB} \div 3 = 0.667 \text{ TB}$. Sending this amount within a three-hour backup window requires about 0.494 Gbit/s of effective throughput: $0.667 \text{ TB} \times 8 \div 3 \text{ hours}$. Factoring in the stated 0.8 bandwidth-utilization value, the required link rate is approximately 0.617 Gbit/s. A 1 GE connection therefore provides sufficient usable bandwidth while retaining reasonable operational headroom. A LAN-free

design places backup data traffic on a dedicated backup path rather than the constrained production LAN, which directly addresses the network constraint. Because client resources are already sufficient, source-side deduplication can reduce transmitted data before it enters that backup path. Huawei positions OceanProtect data-protection solutions around efficient backup and reduction technologies, including deduplication, as described on the [Huawei OceanProtect product page](#). Huawei also provides backup and recovery documentation through its [OceanProtect technical support portal](#).

QUESTION NO: 55

Enable SmartMigration on OceanStor V3 and start a migration task with SmartMigration. Before migration, suppose the LUN ID and the data volume ID of the source LUN are 1 and 2 respectively while the LUN ID and the data volume ID of the destination LUN are 3 and 4 respectively, what are they after migration?

- A. Source LUN: 1 and 3Destination LUN: 2 and 4
- B. Source LUN: 1 and 4Destination LUN: 3 and 2
- C. Source LUN: 3 and 1Destination LUN: 4 and 2
- D. Source LUN: 2 and 4Destination LUN: 1 and 3

ANSWER: C

Explanation:

Source LUN: 3 and 1Destination LUN: 4 and 2 is correct. SmartMigration performs an online migration while preserving the host-facing service identity of the original source LUN. At switchover, the source and destination LUN identities are exchanged so that hosts can continue accessing the service through the source LUN's original logical identity without requiring host-side mapping changes. In this scenario, the source LUN therefore takes the destination LUN ID, changing from 1 to 3, while the destination LUN takes the original source LUN ID, changing from 3 to 1. The corresponding underlying data volumes are also exchanged: the source-side LUN is associated with data volume ID 4, and the destination-side LUN is associated with data volume ID 2. This ID-exchange behavior is fundamental to transparent SmartMigration cutover and enables the migrated data to assume the production LUN's service role. Huawei documents SmartMigration as an online data migration capability in its OceanStor documentation; see the [Huawei OceanStor V3 support documentation](#) and the [OceanStor operation documentation portal](#).

QUESTION NO: 56

Which of the following statements about recovery point objective (RPO) and recovery time objective (RTO) are correct? (Multiple choice.)

- A. RPO describes the maximum acceptable amount of data loss measured in time.
- B. RTO describes the maximum acceptable time to restore a business service.
- C. When the required RPO is zero, synchronous remote replication can be considered.
- D. RPO is the time required to restart the application after a failure.

ANSWER: A B C

Explanation:

RPO defines the maximum acceptable amount of data loss measured in time. For example, if an application has an RPO of 15 minutes, the disaster-recovery design must ensure that no more than approximately 15 minutes of committed data can be lost after a failure. The synchronization interval of asynchronous remote replication is therefore an important factor in meeting the RPO.

RTO defines the maximum acceptable time required to restore a business service after an interruption. It includes not only storage recovery but also such activities as application start-up, host recovery, network switching, and service verification. Synchronous replication is commonly selected when the business requires an RPO of zero and the link latency and

architecture can support synchronous host-write acknowledgement. General contingency-planning terminology is described in [NIST SP 800-34 Rev. 1](#).

QUESTION NO: 57

Huawei OceanStor storage controller BBUs are both in a fault alarm. Which description is incorrect for the upper application?

- A. System does not work.
- B. If the device is powered off unexpectedly, some of the business data will be lost.
- C. Business performance is not affected.
- D. Upper business IO read and write will not have any impact.

ANSWER: B

Explanation:

If the device is powered off unexpectedly, some of the business data will be lost. is correct because the battery backup units protect controller cache contents during an unexpected power interruption. Write data can be acknowledged to a host after it reaches the storage controller cache but before it is permanently committed to backend disks or SSDs. Under normal conditions, a functioning BBU supplies sufficient temporary power for the controller to preserve or destage this cached write data, preventing acknowledged writes from being lost during a power event. When both controller BBUs have fault alarms, this protection is unavailable for cached data. Therefore, if power is unexpectedly removed before outstanding cached writes reach persistent media, the affected data may be lost, potentially causing application-level data inconsistency. This is why BBU faults require prompt maintenance and careful monitoring of cache-protection status. Huawei identifies controller battery/BBU components as part of the OceanStor storage hardware architecture and provides product and support resources for maintaining storage-system reliability: [Huawei OceanStor Dorado](#) and [Huawei Storage Support](#).

QUESTION NO: 58

A banking business running background storage system requirements are as follows:

- (1) Business Impact. The business impact generated by the operating system is about 100KB.
- (2) The service operation system is the whole GE network. The server and operating system are GE network. Follow-up to consider the business needs to upgrade to 10GE.
- (3) Only requires 72TB of total space, the latter expansion priority to expand the disk
- (4) Only 20U of rack space is reserved for the current storage devices.

The configuration of three nodes and switches is prioritized to meet the requirements of the machine room. Cabinet planning based on the above requirements.

What are the more appropriate networking solutions?

- A. Program One: 1 * OceanStor 9000 C Node36 * 2TB SATA configuration 2 * Front-end GE + 2 * Rear-end 10GE
- B. Scenario 2: 3 * Oceanstor 9000 C Node36 * 2TB SATA hard drive; 2 * Front-end GE + back-end 2 * 10GE
- C. program three: 3 * OceanStor 9000 P Node36 * 2TB SAS configuration 2 * Front-end GE + 2 * Rear-end 10GE
- D. Program Four: 3 * OceanStor 9000 M Node36 * 2TB SAS configuration 2 * Front-end GE + 2 * Rear-end 10GE

ANSWER: B

Explanation:

Scenario 2: 3 * Oceanstor 9000 C Node36 * 2TB SATA hard drive;2 * Front-end GE + back-end 2 * 10GE is the appropriate design because it directly satisfies the stated three-node deployment requirement while favoring capacity-oriented growth. Three 36-disk nodes populated with 2 TB SATA drives provide substantially more raw capacity than the immediate 72 TB requirement, leaving a practical capacity reserve and allowing later expansion by adding disks, which is explicitly the preferred expansion approach. SATA media is also appropriate where the requirement emphasizes usable capacity and incremental disk growth rather than a high-performance workload.

The two front-end GE connections match the existing GE server network. At the same time, the two 10GE back-end links provide the cluster interconnection bandwidth required for a scale-out storage architecture and establish an infrastructure path for the planned move toward 10GE. A three-node cluster also supplies the minimum distributed-node layout needed to build the OceanStor 9000 scale-out system, rather than relying on a single node. The capacity-node form factor and disk-focused expansion strategy help keep the initial deployment aligned with the limited rack-space allocation. Huawei describes OceanStor 9000 as a scale-out storage platform designed for capacity expansion and distributed operation; see [Huawei OceanStor 9000 support information](#) and [Huawei scale-out storage portfolio](#).

QUESTION NO: 59

True or False, during the disaster recovery drill, to you want to switch the service to the disaster recovery end. In order to prove the reliability of the disaster recovery system, you need to test bysimulation test and other means as much as possible.

- A. True
- B. False

ANSWER: B

Explanation:

False is correct because a disaster-recovery drill that is intended to switch production service to the disaster-recovery site must validate the real takeover procedure, including service cutover, host access, data availability, application behavior, and the subsequent recovery or switchback process. A simulation can be a useful preparation activity: it helps operators rehearse procedures, review dependencies, and reduce risk before a planned drill. However, simulation alone does not prove that the actual disaster-recovery environment can assume the production workload or that the end-to-end service path will operate correctly after takeover. Reliability evidence for a planned service switch therefore requires executing and recording the relevant operational checks under the approved drill plan, with appropriate maintenance controls and rollback preparations. This approach aligns with Huawei enterprise-storage operational guidance, where disaster-recovery capabilities must be configured, monitored, and verified through the applicable product procedures. See Huawei's [Enterprise Support documentation portal](#) for product-specific disaster-recovery guides and the [Huawei OceanStor storage overview](#) for the platform context in which these protection and recovery capabilities are used.

QUESTION NO: 60

Huawei N8000cluster NAS system which can meet the following requirements?

- A. A TV station needs 20 departments and 10-15 businesses in each department. Each business uses a separate file system
- B. a program distribution system, windows and linux client distribution using cifs and nfs protocol to write the program video, but the two are not written into the videoComplex, FTP client and then read the video file prepared broadcast
- C. A school multimedia video-on-demand business, you need to provide network disk address to the school's 5000 teachers and students, for everyone at any time for online courses.
- D. A research institute of high performance computing business, Xu provide 8GB / s stable read bandwidth

ANSWER: B

Explanation:

The requirement describing a program-distribution system using Windows and Linux clients is the suitable use case for a Huawei N8000 clustered NAS system. This workload needs file-level access from heterogeneous clients: Windows-based production users require CIFS/SMB access, while Linux-based production users require NFS access. The same storage platform must also make completed video files available to an FTP-based broadcast-preparation client. Clustered NAS is designed to provide centralized, scalable file services through these standard file-access protocols, allowing the application workflow to use the protocol appropriate to each client without requiring separate storage silos for the Windows, Linux, and FTP stages.

Video distribution is also a typical NAS workload because users work with files in a shared namespace and the system can scale file-serving resources as demand grows. The stated workflow separates content creation from subsequent FTP reading for broadcast preparation, which is well aligned with a multi-protocol NAS file-service design. Huawei identifies NAS as file storage for unstructured-data workloads and provides NAS products and solutions through its enterprise storage portfolio: [Huawei Enterprise Storage](#). NFS is a standardized network file-system protocol used by UNIX and Linux clients; its file-service model is defined in [RFC 8881](#).