

Google Cloud Certified - Professional Cloud Database Engineer

Google Professional-Cloud-Database-Engineer

Version Demo

Total Demo Questions: 19

Total Premium Questions: 192

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

PASSQUEEN.COM

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Designing highly scalable, available, and reliable cloud database solutions	65
Topic 2, Managing a cloud database solution	40
Topic 3, Designing secure cloud database solutions	29
Topic 4, Designing and implementing data migration solutions	29
Topic 5, Optimizing database performance	28
Topic 6, Mix Questions	1
Total	192

QUESTION NO: 1

Your organization has a critical business app that is running with a Cloud SQL for MySQL backend database. Your company wants to build the most fault-tolerant and highly available solution possible. You need to ensure that the application database can survive a zonal and regional failure with a primary region of us-central1 and the backup region of us-east1. What should you do?

- A.** Provision a Cloud SQL for MySQL instance in us-central1-a.
Create a multiple-zone instance in us-west1-b.
Create a read replica in us-east1-c.
- B.** Provision a high-availability Cloud SQL for MySQL instance in us-central1, with the primary and standby in separate zones.
Create a cross-region read replica in us-east1.
- C.** Provision a Cloud SQL for MySQL instance in us-central1-a.
Create a multiple-zone instance in us-east-b.
Create a read replica in us-east1-c.
- D.** Provision a Cloud SQL for MySQL instance in us-central1-a.
Create a multiple-zone instance in us-east1-b.
Create a read replica in us-central1-b.

ANSWER: B

Explanation:

Provision a high-availability Cloud SQL for MySQL instance in `us-central1`, with its primary and standby placed in separate zones, and create a cross-region read replica in `us-east1`. This combines the two required protection levels. Cloud SQL high availability uses a regional configuration: Cloud SQL maintains a primary instance and standby instance in different zones of the same region. If the primary zone becomes unavailable, Cloud SQL can fail over to the standby, providing resilience to a zonal failure while keeping the primary workload in `us-central1`.

A cross-region read replica in `us-east1` maintains an asynchronously replicated copy in the designated backup region. In the event that `us-central1` has a regional outage, the replica can be promoted to a standalone primary instance and the application can be redirected to it. This design therefore supplies local zonal failover plus a disaster-recovery path for regional failure. Recovery procedures should include tested replica-promotion and application connection-routing steps, because cross-region replicas are not automatically promoted during a regional outage. See the Cloud SQL [high-availability documentation](#) and [cross-region replica documentation](#).

QUESTION NO: 2

You released a popular mobile game and are using a 50 TB Cloud Spanner instance to store game data in a PITR-enabled production environment. When you analyzed the game statistics, you realized that some players are exploiting a loophole to gather more points to get on the leaderboard. Another DBA accidentally ran an emergency bugfix script that corrupted some of the data in the production environment. You need to determine the extent of the data corruption and restore the production environment. What should you do? (Choose two.)

- A.** If the corruption is significant, use backup and restore, and specify a recovery timestamp.
- B.** If the corruption is significant, perform a stale read and specify a recovery timestamp. Write the results back.
- C.** If the corruption is significant, use import and export.
- D.** If the corruption is insignificant, use backup and restore, and specify a recovery timestamp.
- E.** If the corruption is insignificant, perform a stale read and specify a recovery timestamp. Write the results back.

ANSWER: A E

Explanation:

For significant corruption, use the point-in-time recovery capability to restore the database to a chosen recovery timestamp. This is the appropriate recovery pattern when the affected data is widespread or when identifying and manually repairing every changed row would be impractical. Cloud Spanner restores the selected historical version into a new database, allowing the team to validate the recovered state and then safely direct the application to the restored environment as part of the recovery process. PITR is specifically intended to recover from accidental writes, deletions, and data corruption within the configured version-retention period.

For insignificant, well-scoped corruption, perform a stale read at a timestamp before the bad script ran, then write the known-good results back to the current database. A stale read provides a consistent historical view of the relevant data without requiring a full database recovery. This lets the DBA first determine the exact impact of the script and repair only the affected rows or records, minimizing disruption to a large production workload. Cloud Spanner supports timestamp-bound reads for this type of historical data inspection and targeted correction. See [Cloud Spanner point-in-time recovery documentation](#) and [Cloud Spanner read documentation](#).

QUESTION NO: 3

You are running a mission-critical application on a Cloud SQL for PostgreSQL database with a multi-zonal setup. The primary and read replica instances are in the same region but in different zones. You need to ensure that you split the application load between both instances. What should you do?

- A. Use Cloud Load Balancing for load balancing between the Cloud SQL primary and read replica instances.
- B. Use PgBouncer to set up database connection pooling between the Cloud SQL primary and read replica instances.
- C. Use HTTP(S) Load Balancing for database connection pooling between the Cloud SQL primary and read replica instances.
- D. Use the Cloud SQL Auth proxy for database connection pooling between the Cloud SQL primary and read replica instances.

ANSWER: B

Explanation:

Use PgBouncer to set up database connection pooling between the Cloud SQL primary and read replica instances. PgBouncer is a PostgreSQL-aware connection pooler that reduces the number and cost of direct database connections, which is particularly useful for applications with many short-lived connections. It can be configured with separate connection targets for the primary instance and for a read replica. The application can then send write transactions and any strongly consistent reads to the primary target, while directing eligible read-only operations to the replica target. This divides the database workload while preserving the required primary/replica roles.

Cloud SQL read replicas are intended to offload read traffic from the primary. Because replication is asynchronous, application design must account for replication lag: requests that must immediately observe a write should continue to use the primary. PgBouncer provides an appropriate pooling layer for PostgreSQL clients and helps avoid exhausting Cloud SQL connection resources as traffic is distributed. Google documents PgBouncer as a supported external connection-pooling solution for Cloud SQL for PostgreSQL, and documents read replicas as a way to scale read workloads. See [Cloud SQL PostgreSQL connection poolers](#) and [Cloud SQL PostgreSQL read replicas](#).

QUESTION NO: 4

Your organization is running a critical production database on a virtual machine (VM) on Compute Engine. The VM has an ext4-formatted persistent disk for data files. The database will soon run out of storage space. You need to implement a solution that avoids downtime. What should you do?

- A. In the Google Cloud Console, increase the size of the persistent disk, and use the `resize2fs` command to extend the disk.
- B. In the Google Cloud Console, increase the size of the persistent disk, and use the `fdisk` command to verify that the new space is ready to use
- C. In the Google Cloud Console, create a snapshot of the persistent disk, restore the snapshot to a new larger disk, unmount the old disk, mount the new disk, and restart the database service.

D. In the Google Cloud Console, create a new persistent disk attached to the VM, and configure the database service to move the files to the new disk.

ANSWER: A

Explanation:

In the Google Cloud Console, increase the size of the persistent disk, and use the `resize2fs` command to extend the disk. This is the appropriate online expansion workflow for an ext4 filesystem on a Compute Engine persistent disk. Compute Engine supports increasing the size of an attached persistent disk while the VM remains running. The guest operating system can then detect the larger block device capacity without detaching the disk or stopping the database workload. After the operating system sees the additional capacity, `resize2fs` expands the ext4 filesystem so that the database can allocate and use the newly available space.

This approach addresses both required layers: increasing the underlying persistent disk capacity in Google Cloud and extending the filesystem that presents usable storage to the database. Ext4 supports online filesystem growth, so this can be performed without unmounting the filesystem, helping the organization avoid downtime for its critical production database. If the filesystem is contained in a partition, the partition must also be enlarged before extending the filesystem; however, the essential no-downtime pattern remains online disk resizing followed by filesystem expansion. See the Google Cloud guidance for [resizing persistent disks](#) and the documented Linux procedure for [resizing partitions and filesystems](#).

QUESTION NO: 5

Your company uses Cloud Bigtable for a customer-facing application. The application must remain available if a Bigtable cluster or a region becomes unavailable. You also want client requests to be routed automatically to an available cluster. What should you do? (Choose two.)

- A. Create clusters for the Bigtable instance in at least two different regions.
- B. Create and use a multi-cluster routing app profile.
- C. Create all Bigtable clusters in separate zones within one region.
- D. Use a single-cluster routing app profile for all application requests.
- E. Configure a Cloud SQL cross-region read replica for the Bigtable instance.

ANSWER: A B

Explanation:

Create a replicated Bigtable instance with clusters in at least two different regions. Replication maintains copies of the data across clusters and provides the infrastructure required to continue serving traffic when an individual cluster or an entire region is unavailable.

Create and use a multi-cluster routing app profile. Multi-cluster routing automatically directs requests to an available cluster and is the recommended routing configuration for highly available Bigtable workloads. Applications must account for Bigtable replication being asynchronous when requests can be served in different clusters. See the [Cloud Bigtable replication overview](#) and [app profiles documentation](#).

QUESTION NO: 6

You are building an application that allows users to customize their website and mobile experiences. The application will capture user information and preferences. User profiles have a dynamic schema, and users can add or delete information from their profile. You need to ensure that user changes automatically trigger updates to your downstream BigQuery data warehouse. What should you do?

- A. Store your data in Bigtable, and use the user identifier as the key. Use one column family to store user profile data, and use another column family to store user preferences.

- B.** Use Cloud SQL, and create different tables for user profile data and user preferences from your recommendations model. Use SQL to join the user profile data and preferences
- C.** Use Firestore in Native mode to store each user profile as a document, and use the Firestore-to-BigQuery extension to stream user changes to BigQuery.
- D.** Use Firestore in Datastore mode, and store user profile data as a document. Update the user profile with preferences specific to that user and use the user identifier to query.

ANSWER: C

Explanation:

Use Firestore in Native mode to store each user profile as a document, and use the Firestore-to-BigQuery extension to stream user changes to BigQuery. Firestore's document-oriented data model is a strong fit for profiles whose fields vary by user and evolve over time. A profile document can contain a changing set of preference fields, nested maps, and other user-specific attributes without requiring a rigid, predefined relational schema. The application can efficiently read or update a profile by its user identifier, typically used as the document ID.

The official Firebase extension for streaming Cloud Firestore to BigQuery provides the required automated downstream integration. It listens for document writes in a configured Firestore collection and exports change events to BigQuery, enabling the warehouse to receive profile updates without the application having to implement and operate custom export logic. The exported data includes a change-history table and can support downstream transformations, analytics, and profile-based reporting in BigQuery. Firestore Native mode is the appropriate Firestore mode for Firebase application workloads and for this Firebase extension. See the [Stream Firestore to BigQuery extension documentation](#) and the [Cloud Firestore data model documentation](#).

QUESTION NO: 7

You are creating a new Cloud SQL for PostgreSQL instance that must use a customer-managed encryption key (CMEK). The instance will be located in us-central1. What should you do? (Choose two.)

- A.** Create or use a Cloud KMS key in us-central1, and specify the key when creating the Cloud SQL instance.
- B.** Grant the Cloud SQL service agent the Cloud KMS CryptoKey Encrypter/Decrypter role on the key.
- C.** Create the Cloud KMS key in a different region from the Cloud SQL instance.
- D.** Create the Cloud SQL instance with Google-managed encryption, and add CMEK after the instance is created.
- E.** Grant the application service account the Owner role on the Cloud KMS project.

ANSWER: A B

Explanation:

Create or use a Cloud KMS key in the same region as the Cloud SQL instance and specify that key when creating the instance. Cloud SQL CMEK keys must use a location compatible with the instance location. For a regional Cloud SQL instance in us-central1, use a key in us-central1.

Grant the Cloud SQL service agent the Cloud KMS CryptoKey Encrypter/Decrypter role on the selected key. This permission allows Cloud SQL to use the customer-managed key for its encryption operations. CMEK must be configured during instance creation, so planning the key and IAM permissions before provisioning is required. See the [Cloud SQL for PostgreSQL CMEK documentation](#).

QUESTION NO: 8

You are migrating your data center to Google Cloud. You plan to migrate your applications to Compute Engine and your Oracle databases to Bare Metal Solution for Oracle. You must ensure that the applications in different projects can communicate securely and efficiently with the Oracle databases. What should you do?

- A. Set up a Shared VPC, configure multiple service projects, and create firewall rules.
- B. Set up Serverless VPC Access.
- C. Set up Private Service Connect.
- D. Set up Traffic Director.

ANSWER: A

Explanation:

Set up a Shared VPC, configure multiple service projects, and create firewall rules. Shared VPC is designed for organizations that need centrally managed network connectivity while keeping workloads and administrative boundaries in separate projects. A host project owns the VPC network and its subnets, while the Compute Engine application projects can be attached as service projects. This lets their VM instances use the same centrally governed network and communicate privately with Oracle databases deployed through Bare Metal Solution connectivity. Central network administrators can define consistent routes, private connectivity, and firewall policy, while application teams retain control over resources in their respective service projects.

For Bare Metal Solution, Google documents configuring connectivity to a VPC network so that client workloads can access Oracle database services using private IP addressing. Combining this with Shared VPC avoids duplicating network infrastructure or relying on public paths between independently managed application projects and the database environment. Firewall rules can then explicitly permit only the required application-to-database flows, such as Oracle listener traffic, which supports least-privilege access and secure east-west communication. See the [Shared VPC documentation](#) and the [Bare Metal Solution Oracle networking documentation](#).

QUESTION NO: 9

Your company wants to move to Google Cloud. Your current data center is closing in six months. You are running a large, highly transactional Oracle application footprint on VMWare. You need to design a solution with minimal disruption to the current architecture and provide ease of migration to Google Cloud. What should you do?

- A. Migrate applications and Oracle databases to Google Cloud VMware Engine (VMware Engine).
- B. Migrate applications and Oracle databases to Compute Engine.
- C. Migrate applications to Cloud SQL.
- D. Migrate applications and Oracle databases to Google Kubernetes Engine (GKE).

ANSWER: A

Explanation:

Migrate applications and Oracle databases to Google Cloud VMware Engine (VMware Engine). This approach is designed for organizations that need to move VMware-based workloads to Google Cloud quickly while retaining their existing VMware operational model. VMware Engine provides a fully managed VMware Cloud Foundation environment, including familiar VMware vSphere, vSAN, NSX, and vCenter capabilities. The existing Oracle application servers and database virtual machines can therefore be migrated with substantially less architectural change than a redesign for native Compute Engine, containers, or a managed database platform would require.

Given the six-month data-center closure deadline, preserving the current VM-based architecture reduces migration risk and the amount of validation, refactoring, and operational retraining required before cutover. Google documents migration paths that enable moving on-premises VMware workloads to VMware Engine, including live migration capabilities through VMware HCX where appropriate. This supports a phased migration strategy with low disruption for transactional workloads. The organization should also validate Oracle licensing, support terms, capacity, network connectivity, performance, backup, and disaster-recovery requirements as part of detailed implementation planning. See the [Google Cloud VMware Engine overview](#) and [Google Cloud guidance for migrating VMWare workloads](#).

QUESTION NO: 10

You are planning to migrate a 10TB relational database from an on-premises environment to Cloud SQL for PostgreSQL. The database contains sensitive customer

Information. You want to follow Google-recommended practices to keep data secure during the migration. What should you do?

Choose 2 answers

- A. Establish a Private Service Connect connection between your on-premises environment and the Cloud SQL instance
- B. Set up identity Access Management (IAM) roles to restrict access with Cloud SQL with an internal IP address.
- C. Use an external IP address for the Cloud SQL Instance, and configure firewall rules.
- D. Configure Cloud SQL for automatic patching, and enable binary logging.
- E. Leverage Storage Transfer Service with client side encryption.

ANSWER: A B

Explanation:

Establish a Private Service Connect connection between your on-premises environment and the Cloud SQL instance and Set up identity Access Management (IAM) roles to restrict access with Cloud SQL with an internal IP address. provide complementary controls for protecting sensitive data throughout a large migration. Private Service Connect provides private connectivity to Cloud SQL, so database traffic can remain on private Google Cloud networking rather than requiring exposure through a public IP address. For an on-premises source, the private network path is extended through hybrid connectivity such as Cloud VPN or Cloud Interconnect. This reduces the network exposure of customer data while it is copied to the destination instance.

IAM should also follow least-privilege principles. Granting only the Cloud SQL permissions required by migration operators, service accounts, and administrators limits who can create connections, manage instances, or access related Cloud SQL resources. Combining IAM authorization with internal/private IP connectivity means that a caller must have both appropriate identity permissions and access to the intended private network path. Google recommends using private IP connectivity when possible and controlling administrative access through IAM. See the [Cloud SQL private services access documentation](#) and Google's [Cloud SQL IAM roles documentation](#).

QUESTION NO: 11

You are building a data warehouse on BigQuery. Sources of data include several MySQL databases located on-premises.

You need to transfer data from these databases into BigQuery for analytics. You want to use a managed solution that has low latency and is easy to set up. What should you do?

- A. Create extracts from your on-premises databases periodically, and push these extracts to Cloud Storage. Upload the changes into BigQuery, and merge them with existing tables.
- B. Use Cloud Data Fusion and scheduled workflows to extract data from MySQL. Transform this data into the appropriate schema, and load this data into your BigQuery database.
- C. Use Datastream to connect to your on-premises database and create a stream. Have Datastream write to Cloud Storage. Then use Dataflow to process the data into BigQuery.
- D. Use Database Migration Service to replicate data to a Cloud SQL for MySQL instance. Create federated tables in BigQuery on top of the replicated instances to transform and load the data into your BigQuery database.

ANSWER: C

Explanation:

Use Datastream to connect to your on-premises database and create a stream. Have Datastream write to Cloud Storage. Then use Dataflow to process the data into BigQuery. Datastream is Google Cloud's serverless change data capture service and supports MySQL sources, including sources hosted outside Google Cloud when private connectivity is configured. It performs an initial backfill and then continuously captures ongoing changes from the MySQL binary logs, providing the low-latency replication required for analytics rather than relying on periodic batch extracts.

Writing the change stream to Cloud Storage provides a durable, managed landing zone for the captured records. A managed Dataflow pipeline can then consume and apply the captured changes to BigQuery, including handling inserts, updates, and deletes to keep analytics tables current. This architecture separates capture from processing, scales without managing CDC infrastructure, and uses managed Google Cloud services throughout. Google documents Datastream's supported MySQL source configuration and connectivity requirements in the [MySQL source documentation](#). Google also provides guidance and templates for processing Datastream output with Dataflow and loading it into BigQuery in the [Datastream to BigQuery Dataflow template documentation](#).

QUESTION NO: 12

You are writing an application that will run on Cloud Run and require a database running in the Cloud SQL managed service. You want to secure this instance so that it only receives connections from applications running in your VPC environment in Google Cloud. What should you do?

- A.** 1. Create your instance with a specified external (public) IP address.2. Choose the VPC and create firewall rules to allow only connections from Cloud Run into your instance.3. Use Cloud SQL Auth proxy to connect to the instance.
- B.** 1. Create your instance with a specified external (public) IP address.2. Choose the VPC and create firewall rules to allow only connections from Cloud Run into your instance.3. Connect to the instance using a connection pool to best manage connections to the instance.
- C.** 1. Create your instance with a specified internal (private) IP address.2. Choose the VPC with private service connection configured.3. Configure the Serverless VPC Access connector in the same VPC network as your Cloud SQL instance.4. Use Cloud SQL Auth proxy to connect to the instance.
- D.** 1. Create your instance with a specified internal (private) IP address.2. Choose the VPC with private service connection configured.3. Configure the Serverless VPC Access connector in the same VPC network as your Cloud SQL instance.4. Connect to the instance using a connection pool to best manage connections to the instance.

ANSWER: D

Explanation:

Creating the Cloud SQL instance with an internal (private) IP address, selecting a VPC that has private service connection configured, configuring a Serverless VPC Access connector in that VPC, and using a connection pool is the appropriate design. A Cloud SQL private IP gives the instance an address reachable through the selected VPC rather than exposing the database through a public IP endpoint. Private connectivity for Cloud SQL depends on private services access being configured between the consumer VPC and Google-managed services.

Cloud Run is serverless and is not inherently attached to a customer VPC. A Serverless VPC Access connector supplies the required network path from the Cloud Run service to private addresses in that VPC, including the Cloud SQL private IP. Configure the Cloud Run service's VPC egress appropriately so traffic to private IP ranges is sent through this connector. Finally, Cloud Run can create many concurrent instances, each potentially opening database connections. A connection pool reuses a bounded set of database connections and helps prevent the Cloud SQL instance from reaching its connection limit while retaining the required private network path. See [Configure private IP for Cloud SQL](#) and [Connect Cloud Run to a VPC network](#).

QUESTION NO: 13

Your organization has an existing app that just went viral. The app uses a Cloud SQL for MySQL backend database that is experiencing slow disk performance while using hard disk drives (HDDs). You need to improve performance and reduce disk I/O wait times. What should you do?

- A.** Export the data from the existing instance, and import the data into a new instance with solid-state drives (SSDs).

- B. Edit the instance to change the storage type from HDD to SSD.
- C. Create a high availability (HA) failover instance with SSDs, and perform a failover to the new instance.
- D. Create a read replica of the instance with SSDs, and perform a failover to the new instance

ANSWER: A

Explanation:

Export the data from the existing instance, and import the data into a new instance with solid-state drives (SSDs). Cloud SQL storage type is selected when an instance is created and cannot be changed from HDD to SSD in place. Therefore, the appropriate migration approach is to create a replacement Cloud SQL for MySQL instance configured with SSD storage, export the source database, and import it into that new instance. SSD-backed persistent disk provides substantially better I/O performance and lower latency than HDD-backed storage, which directly addresses disk I/O wait as demand grows. Plan the migration carefully: create the target instance with sufficient CPU, memory, storage, network configuration, users, and database flags; take an export or other suitable logical migration copy; import the data; validate application behavior; and redirect application connections during a controlled cutover. For highly active workloads, the team should also account for writes that occur after the export begins and schedule an appropriate maintenance window or use a replication-based migration pattern where applicable. Google documents that the storage type is immutable after instance creation in the [Cloud SQL for MySQL instance settings documentation](#). The supported SQL export and import workflow is described in the [Cloud SQL import and export documentation](#).

QUESTION NO: 14

Your company uses Cloud Spanner for a mission-critical inventory management system that is globally available. You recently loaded stock keeping unit (SKU) and product catalog data from a company acquisition and observed hot-spots in the Cloud Spanner database. You want to follow Google-recommended schema design practices to avoid performance degradation. What should you do? (Choose two.)

- A. Use an auto-incrementing value as the primary key.
- B. Normalize the data model.
- C. Promote low-cardinality attributes in multi-attribute primary keys.
- D. Promote high-cardinality attributes in multi-attribute primary keys.
- E. Use bit-reverse sequential value as the primary key.

ANSWER: D E

Explanation:

Use a bit-reverse sequential value as the primary key and promote high-cardinality attributes in multi-attribute primary keys. Cloud Spanner stores rows in primary-key order and splits data ranges across servers. Consequently, a sequential identifier placed directly in the primary key concentrates new writes at the end of a key range, which can create a hotspot. Bit reversal disperses consecutively generated numeric values across the key space while retaining a sequential source value, allowing writes to be distributed among multiple splits. Google provides the `BIT_REVERSED_POSITIVE` identity option specifically for this pattern. For compound primary keys, placing a high-cardinality column first creates many distinct leading-key values and improves distribution of reads and writes. This is particularly relevant for catalog and SKU workloads, where concurrent updates or newly loaded records can otherwise target a limited set of adjacent key ranges. These schema choices distribute traffic more evenly while preserving efficient access patterns based on the primary key. See Google Cloud's [guidance for preventing hotspots with primary keys](#) and the documentation for [bit-reversed identity columns](#).

QUESTION NO: 15

Your organization uses Cloud Spanner for a financial application. An administrator accidentally deleted records from several tables. Point-in-time recovery is enabled, and the deletion time is known. You need to recover the data while minimizing risk to the current production database. What should you do? (Choose two.)

- A. Restore the database by using point-in-time recovery and specify a timestamp before the deletion.
- B. Validate the restored database before copying affected data or redirecting application traffic.
- C. Fail over the Cloud Spanner instance to recover the deleted records.
- D. Increase Cloud Spanner processing units to reverse the deletion.
- E. Delete the production database before starting the point-in-time recovery restore.

ANSWER: A B

Explanation:

Use Cloud Spanner point-in-time recovery to restore the database to a timestamp before the deletion. A point-in-time recovery restore creates a new database from the selected historical state rather than overwriting the existing production database. This preserves the current database for investigation and allows the recovered state to be validated safely.

Validate the restored database and then use an appropriate migration or application cutover procedure to recover the required data. For limited corruption, the team can compare the restored database with production and copy only the affected records. For widespread corruption, redirecting the application after validation can be appropriate. See [Cloud Spanner point-in-time recovery](#) and [Restore a database](#).

QUESTION NO: 16

Your company uses Bigtable for a customer-facing application. The application must remain available if a Bigtable cluster or an entire region becomes unavailable. You also need application requests to be routed automatically to a healthy cluster. What should you do? (Choose two.)

- A. Create clusters in at least two different regions in the same replicated Bigtable instance.
- B. Configure a multi-cluster routing app profile for the application.
- C. Create separate Bigtable instances in each region, and configure the application to use both instances.
- D. Configure a single-cluster routing app profile for all application requests.
- E. Use a read replica of the Bigtable cluster in a second region.

ANSWER: A B

Explanation:

Create a replicated Bigtable instance with clusters in at least two different regions. Replication maintains copies of Bigtable data across clusters, providing protection against the loss of a single cluster or a regional outage. The application can continue serving from another available cluster when a cluster becomes unavailable.

Configure a multi-cluster routing app profile for application traffic. Multi-cluster routing automatically directs requests to an available cluster in the instance and supports high availability without requiring the application to select a specific cluster. The application must account for Bigtable replication behavior because replication between clusters is asynchronous. See the [Bigtable replication overview](#) and [Bigtable app profiles documentation](#).

QUESTION NO: 17

You plan to use Database Migration Service to migrate data from a PostgreSQL on-premises instance to Cloud SQL. You need to identify the prerequisites for creating and automating the task. What should you do? (Choose two.)

- A. Drop or disable all users except database administration users.
- B. Disable all foreign key constraints on the source PostgreSQL database.

- C. Ensure that all PostgreSQL tables have a primary key.
- D. Shut down the database before the Data Migration Service task is started.
- E. Ensure that pglogical is installed on the source PostgreSQL database.

ANSWER: C E

Explanation:

Ensure that all PostgreSQL tables have a primary key. is required for reliable ongoing logical replication. Database Migration Service uses logical replication for a continuous PostgreSQL migration. A primary key provides the row identity needed to correctly apply updates and deletes at the Cloud SQL destination. Before creating the migration job, review the tables that must participate in ongoing replication and add an appropriate primary key where one is absent.

Ensure that pglogical is installed on the source PostgreSQL database. is also a source prerequisite for Database Migration Service PostgreSQL migrations that use pglogical-based logical replication. The source instance must be prepared to support logical decoding and replication, including installing and configuring the required pglogical extension and related replication settings. This preparation allows Database Migration Service to perform the initial data load and then continuously capture and apply subsequent source changes automatically. Google's PostgreSQL source-configuration documentation describes the required logical-replication and pglogical setup, while its migration guidance identifies the need for a usable replica identity, such as a primary key, for replicated tables. See [Configure a PostgreSQL source database](#) and [Database Migration Service for PostgreSQL overview](#).

QUESTION NO: 18

You want to migrate an on-premises 100 TB Microsoft SQL Server database to Google Cloud over a 1 Gbps network link. You have 48 hours allowed downtime to migrate this database. What should you do? (Choose two.)

- A. Use a change data capture (CDC) migration strategy.
- B. Move the physical database servers from on-premises to Google Cloud.
- C. Keep the network bandwidth at 1 Gbps, and then perform an offline data migration.
- D. Increase the network bandwidth to 2 Gbps, and then perform an offline data migration.
- E. Increase the network bandwidth to 10 Gbps, and then perform an offline data migration.

ANSWER: A E

Explanation:

Use a change data capture (CDC) migration strategy because it separates the long-running initial copy from the service interruption. The source SQL Server database can remain online while the bulk data is migrated and ongoing changes are captured and replicated. At cutover, writes are stopped briefly, remaining captured changes are applied, and clients are redirected to the destination. This approach is appropriate when the available downtime is substantially shorter than the time needed for an initial full transfer over the existing connection. Google Database Migration Service supports continuous migrations for SQL Server and uses an initial data load followed by continuous replication, enabling a low-downtime cutover.

Increase the network bandwidth to 10 Gbps, and then perform an offline data migration because a 100 TB transfer contains roughly 800 terabits. At an ideal sustained 10 Gbps, transferring that volume takes about 22.2 hours, leaving meaningful allowance within a 48-hour migration window for transfer overhead and migration operations. A 1 Gbps link would require about 9.3 days in ideal conditions, while 2 Gbps would still require about 4.6 days, so neither can complete the full offline copy within the permitted outage. Capacity planning must also account for protocol overhead, actual throughput, and validation time. See [Google Cloud SQL Server continuous migration documentation](#) and [Cloud Interconnect bandwidth documentation](#).

QUESTION NO: 19

You are the DBA of an online tutoring application that runs on a Cloud SQL for PostgreSQL database. You are testing the implementation of the cross-regional failover configuration. The database in region R1 fails over successfully to region R2, and the database becomes available for the application to process data. During testing, certain scenarios of the application work as expected in region R2, but a few scenarios fail with database errors. The application-related database queries, when executed in isolation from Cloud SQL for PostgreSQL in region R2, work as expected. The application performs completely as expected when the database fails back to region R1. You need to identify the cause of the database errors in region R2. What should you do?

- A. Determine whether the versions of Cloud SQL for PostgreSQL in regions R1 and R2 are different.
- B. Determine whether the database patches of Cloud SQL for PostgreSQL in regions R1 and R2 are different.
- C. Determine whether the failover of Cloud SQL for PostgreSQL from region R1 to region R2 is in progress or has completed successfully.
- D. Determine whether Cloud SQL for PostgreSQL in region R2 is a near-real-time copy of region R1 but not an exact copy.

ANSWER: D

Explanation:

Determine whether Cloud SQL for PostgreSQL in region R2 is a near-real-time copy of region R1 but not an exact copy. Cross-region read replicas provide asynchronous replication, so the replica can lag behind the primary. At the instant a failover or promotion occurs, transactions that were committed on the primary but had not yet been replayed on the replica might not be present in region R2. This can create application-level inconsistencies, such as a workflow referencing rows, state changes, or related records that are not yet available after promotion. Individual queries can still execute correctly when tested directly because their SQL syntax, permissions, and database objects are valid; the failure is caused by the application encountering an incomplete or inconsistent data state. The fact that the application works again after returning to region R1 further supports replication lag or unreplicated transactions as the relevant condition. For failover testing, Google recommends verifying that the replica has received and processed the primary's transactions before promoting it, using replication status and replay/replication lag information. See the [Cloud SQL cross-region replica documentation](#) and Google's guidance to [verify failover criteria](#).