

Administering Relational Databases on Microsoft Azure

Microsoft DP-300

Version Demo

Total Demo Questions: 45

Total Premium Questions: 457

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Plan and implement data platform resources	96
Topic 2, Implement a secure environment	90
Topic 3, Monitor and optimize operational resources	90
Topic 4, Optimize query performance	36
Topic 5, Perform automation of tasks	40
Topic 6, Plan and implement a High Availability and Disaster Recovery (HADR) environment	86
Topic 7, Mix Questions	2
Topic 8, Case Study 1	2
Topic 9, Case Study 2	3
Topic 10, Case Study 3	2
Topic 11, Case Study 4	3
Topic 12, Case Study 5	2
Topic 13, Case Study 6	2
Topic 14, Case Study 7	3
Total	457

QUESTION NO: 1 - (SIMULATION)

Task 10

You need to ensure that indexes are created automatically for all the databases on sql60152867. You may need to use SQL Server Management Studio and the Azure portal.

ANSWER: See the explanation for the answer

Explanation:

Configure automatic index creation at the **logical SQL server** level so that it applies to all databases.

1. In the Azure portal, open **SQL servers** and select **sql60152867**.
2. Under **Intelligent Performance**, select **Automatic tuning**.
3. Set **Create index to On**.
4. Select **Apply** (or **Save**).

Answer: On server sql60152867, enable Automatic tuning `CREATE_INDEX = ON`.

Server-level automatic tuning is inherited by the server's databases unless a database has an explicit override. If any database is configured with custom automatic-tuning settings, change it to **Inherit from server**, or enable index creation directly in that database:

```
ALTER DATABASE CURRENT  
SET AUTOMATIC_TUNING (CREATE_INDEX = ON);
```

QUESTION NO: 2

You have an Azure Data Factory pipeline that is triggered hourly.

The pipeline has had 100% success for the past seven days.

The pipeline execution fails, and two retries that occur 15 minutes apart also fail. The third failure returns the following error.

```
ErrorCode=UserErrorFileNotFound,  
'Type=Microsoft.DataTransfer.Common.Shared.HybridDeliveryException,Message=ADLS  
Gen2 operation failed for: Operation returned an invalid status code  
'NotFound'. Account: 'contosoproduksouth' FileSystem: wwi.Path:  
'BIKES/CARBON/year=2021/month=01/day=10/hour=06'. ErrorCode:  
'PathNotFound'.Message: 'The specified path does not exist.'. RequestId:  
'6d269b78-901f-001b-4924-e7a7bc000000'. Timestamp: 'Sun, 10 Jan 2021 07:45:05
```

What is a possible cause of the error?

- A. From 06:00 to 07:00 on January 10, 2021, there was no data in wwi/BIKES/CARBON.
- B. The parameter used to generate year=2021/month=01/day=10/hour=06 was incorrect.
- C. From 06:00 to 07:00 on January 10, 2021, the file format of data in wwi/BIKES/CARBON was incorrect.
- D. The pipeline was triggered too early.

ANSWER: B

Explanation:

The parameter used to generate `year=2021/month=01/day=10/hour=06` was incorrect. Azure Data Factory pipelines commonly use dynamic content to derive a partitioned source path from trigger metadata, pipeline parameters, or date functions. The evaluated values are then used to construct paths such as `wwi/BIKES/CARBON/year=2021/month=01/day=10/hour=06`. If a parameter expression uses an unexpected timestamp, applies an incorrect offset, or formats a date component differently from the source data's partition convention,

the activity resolves to a path that does not correspond to the expected hourly partition and can fail consistently on retry. Retrying does not change the already evaluated parameter values for the failed run, so repeated failures are consistent with an incorrect dynamically generated path. Azure Data Factory exposes trigger and pipeline system variables that can be used in expressions, including trigger time values; these should be validated against the naming convention used by the source folders. Review the activity input and resolved dynamic content for the failed run to confirm the generated year, month, day, and hour values. See [Azure Data Factory system variables](#) and [Azure Data Factory expression language and functions](#).

QUESTION NO: 3 - (DRAG DROP)

You have SQL Server on an Azure virtual machine that contains a database named DB1. DB1 is 30 TB and has a 1-GB daily rate of change.

You back up the database by using a Microsoft SQL Server Agent job that runs Transact-SQL commands. You perform a weekly full backup on Sunday, daily differential backups at 01:00, and transaction log backups every five minutes.

The database fails on Wednesday at 10:00.

Which three backups should you restore in sequence? To answer, move the appropriate backups from the list of backups to the answer area and arrange them in the correct order.

Actions

Monday, Tuesday, and then Wednesday differential backups

Wednesday, Tuesday, and then Monday log backups

full backup

Monday, Tuesday, and then Wednesday log backups

Wednesday, Tuesday, and then Monday differential backups

Wednesday log backups

Wednesday differential backup

Answer Area

>

<

↑

↓

ANSWER:

Explanation:

The restore sequence is **full backup**, followed by the **Wednesday differential backup**, followed by the **Wednesday log backups**. A full backup is always the starting point for this type of restore chain because it provides the complete baseline from which later changes can be recovered. In this case, the full backup was made on Sunday.

The differential backup from Wednesday at 01:00 is restored next. A differential backup contains every extent changed since the most recent full backup, rather than only the changes since the prior differential backup. Therefore, restoring the latest differential backup efficiently brings DB1 forward from the Sunday full-backup state to its Wednesday 01:00 state. There is no need to restore the Monday or Tuesday differential backups because the Wednesday differential already includes those changes.

Finally, restore the transaction log backups taken on Wednesday, in their original chronological sequence, beginning with the first log backup needed after the Wednesday differential and continuing through the last log backup available before the 10:00 failure. Since transaction log backups occur every five minutes, this provides a very small potential data-loss window. If the transaction log is accessible after the failure, a tail-log backup should also be taken and restored last to recover

transactions up to the point of failure. Microsoft documents the required full, differential, and log restore ordering in its [differential database restore guidance](#) and explains [tail-log backups](#) for point-of-failure recovery.

QUESTION NO: 4

You have an Azure Synapse Analytics dedicated SQL pool.

You run `PDW_SHOWSPACEUSED('dbo.FactInternetSales');` and get the results shown in the following table.

ROWS	RESERVED_SPACE	DATA_SPACE	INDEX_SPACE	UNUSED_SPACE	PDW_NODE_ID	DISTRIBUTION_ID
694	2776	616	48	2112	1	1
407	2704	576	48	2080	1	2
53	2376	512	16	1848	1	3
58	2376	512	16	1848	1	4
168	2632	528	32	2072	1	5
195	2696	536	32	2128	1	6
5995	3464	1424	32	2008	1	7
0	2232	496	0	1736	1	8
264	2576	544	40	1992	1	9
3008	3016	960	32	2024	1	10
---	---	---	---	---	---	---
1550	2832	752	48	2032	1	50
1238	2832	696	40	2096	1	51
192	2632	528	32	2072	1	52
1127	2768	680	48	2040	1	53
1244	3032	704	64	2264	1	54
409	2632	568	32	2032	1	55
0	2232	496	0	1736	1	56
1437	2832	728	40	2064	1	57
0	2232	496	0	1736	1	58
384	2632	560	32	2040	1	59
225	2768	544	40	2184	1	60

Which statement accurately describes the `dbo.FactInternetSales` table?

- A. The table contains less than 10,000 rows.
- B. All distributions contain data.
- C. The table uses round-robin distribution
- D. The table is skewed.

ANSWER: D

Explanation:

The table is skewed. In an Azure Synapse Analytics dedicated SQL pool, a distributed table is divided across 60 distributions. Efficient massively parallel processing depends on rows being distributed relatively evenly, because a query step can complete only as quickly as the distribution with the greatest workload. `PDW_SHOWSPACEUSED` reports row counts and space information by distribution, making it a useful diagnostic tool for identifying an uneven data spread.

The displayed results show a material difference between the numbers of rows held in the distributions. This exceeds the commonly used 10% variance guideline for a balanced table and indicates distribution skew. Skew can cause some compute resources to process substantially more data than others, increasing query duration and reducing parallelism. The appropriate next investigation is the table's distribution strategy and, for a hash-distributed table, whether its distribution column has sufficiently high cardinality and an even value frequency. Microsoft recommends choosing distribution columns that support even row distribution and minimize data movement for common joins. See [Guidance for table distribution in dedicated SQL pools](#) and [PDW_SHOWSPACEUSED documentation](#).

QUESTION NO: 5

You have a SQL Server on Azure Virtual Machines instance named SQLVM1.

You need to configure automated backups for the databases on SQLVM1. The backups must support point-in-time restore and be retained for 35 days.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Enable automated backup for SQLVM1.
- B. Configure the automated backup retention period to 35 days.
- C. Enable geo-replication for each database.
- D. Create a SQL Server Agent job that runs DBCC CHECKDB.
- E. Configure a database-scoped credential.

ANSWER: A B

Explanation:

Enable automated backup for SQLVM1 and configure the retention period to 35 days. Automated backup for SQL Server on Azure Virtual Machines configures scheduled full, differential, and transaction log backups to Azure Storage. Transaction log backups provide the backup chain required for point-in-time recovery.

The automated backup retention setting controls how long backup files are retained in the configured storage account. A value of 35 days meets the stated retention requirement. SQL Server Agent jobs can be used for custom backup scheduling, but they do not provide the managed automated-backup configuration requested in this scenario.

For more information, see [Automated backup for SQL Server on Azure Virtual Machines](#).

QUESTION NO: 6

You have an Azure Synapse Analytics dedicated SQL pool that contains a large fact table named FactSales and several small dimension tables.

FactSales is frequently joined to the dimension tables by using a common product key. The dimension tables are small enough to fit on every compute node and change infrequently.

You need to minimize data movement during queries.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Hash distribute FactSales on the product key.
- B. Replicate the dimension tables.
- C. Round-robin distribute FactSales.
- D. Hash distribute each dimension table on a different key.
- E. Create clustered indexes on the dimension tables.

ANSWER: A B

Explanation:

Distribute FactSales by using hash distribution on the product key. Hash distribution places rows that have the same distribution-key value on the same distribution. When the fact table is joined to another hash-distributed table on the same key, the query can perform the join locally without redistributing the large fact table.

Replicate the small dimension tables. A replicated table is copied to every compute node. Because a local copy of each dimension is available on every node, joins between the distributed fact table and the dimensions do not require data movement for the dimension rows. Replication is appropriate only for tables small enough to be stored on every compute node.

For more information, see [Distribution options for tables in dedicated SQL pool](#).

QUESTION NO: 7

You need to recommend a solution to meet the security requirements and the business requirements for DB3. What should you recommend as the first step of the solution?

- A. Run the `sys.sp_cdc_enable_db` stored procedure.
- B. Run the `ALTER TABLE` statement and specify the `ENABLE CHANGE_TRACKING` clause.
- C. Run the `ALTER DATABASE` statement and specify `SET CHANGE_TRACKING = ON`.
- D. Run the `sp_addarticle` stored procedure.

ANSWER: C

Explanation:

Run the `ALTER DATABASE` statement and specify `SET CHANGE_TRACKING = ON`. Change tracking must first be enabled at the database level before it can be enabled for individual tables. This database-level setting establishes the change-tracking infrastructure and configuration for DB3, including the retention period and automatic cleanup behavior if those settings are specified. After change tracking is enabled for the database, the required tables can be configured individually by using `ALTER TABLE ... ENABLE CHANGE_TRACKING`. This sequencing is important because table-level change tracking depends on the database-level feature already being active. Change tracking is designed to efficiently identify rows that have changed since a specified synchronization version, which supports applications or processes that need to synchronize only incremental data changes rather than repeatedly processing entire tables. Microsoft documents database-level enablement as the prerequisite for enabling change tracking on tables. See [Enable and disable change tracking](#) and [ALTER DATABASE SET options](#).

QUESTION NO: 8

You have 12 Azure SQL managed instances that host databases for separate development teams.

The managed instances use low amounts of compute and storage resources. The instances are located in the same Azure region.

You need to reduce costs while retaining separate managed instance boundaries for each team.

What should you use?

- A. an Azure SQL Database elastic pool
- B. an Azure SQL Managed Instance instance pool
- C. an Azure SQL Database serverless database
- D. a failover group

ANSWER: B

Explanation:

An **Azure SQL Managed Instance instance pool** is the appropriate solution. Instance pools provide a cost-efficient way to deploy multiple managed instances that have predictable and relatively low resource usage. The instances share allocated virtual cores and storage within the pool while retaining separate managed instance boundaries.

Each team can continue to use its own managed instance, including separate databases, logins, and instance-level configuration, without requiring capacity to be provisioned independently for every instance. This can reduce costs when multiple smaller instances have unused capacity.

For more information, see [Azure SQL Managed Instance pools](#).

QUESTION NO: 9

You have a Microsoft SQL Server 2019 instance in an on-premises datacenter. The instance contains a 4-TB database named DB1.

You plan to migrate DB1 to an Azure SQL Database managed instance.

What should you use to minimize downtime and data loss during the migration?

- A. database mirroring
- B. distributed availability groups
- C. Always On Availability Group
- D. Azure Database Migration Service

ANSWER: D

Explanation:

Azure Database Migration Service is the appropriate choice because it supports online migrations from SQL Server to Azure SQL Managed Instance. An online migration first restores a full database backup to the managed instance, then continuously restores subsequent transaction-log backups while the source database remains online and available to applications. This ongoing log restore process keeps the target closely synchronized with DB1 throughout the bulk of the migration, which is particularly valuable for a 4-TB database because copying and restoring the initial full backup can take substantial time.

When the organization is ready to cut over, it stops application activity briefly, takes and restores the final transaction-log backup, and points applications to the managed instance. This approach reduces the outage to the final cutover window and minimizes potential data loss by applying transaction-log changes generated after the initial backup. Azure Database Migration Service also provides migration monitoring and validation support, helping administrators track progress and complete the cutover in a controlled manner. Microsoft documents online migration to Azure SQL Managed Instance with Azure DMS at [Migrate SQL Server to Azure SQL Managed Instance online](#) and describes Azure DMS capabilities at [Azure Database Migration Service overview](#).

QUESTION NO: 10

You have two instances of SQL Server on Azure Virtual Machines that are configured as replicas in an Always On availability group.

You need to offload reporting queries from the primary replica to the secondary replica.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Configure the secondary replica to allow read-only connections.
- B. Configure read-only routing for the availability group.
- C. Configure synchronous commit for every replica.

D. Create a failover cluster instance.

E. Enable transactional replication.

ANSWER: A B

Explanation:

You must configure the secondary replica to allow read-only connections and configure read-only routing. A readable secondary replica can accept connections intended for reporting workloads. Read-only routing defines the destination replica for connections that specify read-only application intent.

Applications use `ApplicationIntent=ReadOnly` in their connection strings when connecting through the availability group listener. This allows the listener to direct eligible read-only connections according to the configured routing list while normal read-write activity continues to use the primary replica.

For more information, see [Configure read-only routing for an availability group](#).

QUESTION NO: 11

You have an Azure SQL managed instance named SQLMI1.

You need to grant a Microsoft Entra group named DBAdmins permissions to manage users in a database named DB1 on SQLMI1.

Which three actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Configure a Microsoft Entra administrator for SQLMI1.
- B. Create a login for DBAdmins FROM EXTERNAL PROVIDER.
- C. Create a user in DB1 for the DBAdmins login.
- D. Create a SQL authentication login for DBAdmins.
- E. Assign DBAdmins the Owner role for the SQLMI1 resource.

ANSWER: A B C

Explanation:

First, configure a Microsoft Entra administrator for SQLMI1. The Microsoft Entra administrator is required to establish Microsoft Entra authentication and to create Microsoft Entra server principals on the managed instance.

Create a login for DBAdmins by using `FROM EXTERNAL PROVIDER`. Azure SQL Managed Instance supports Microsoft Entra server principals, which can then be mapped to users in individual databases. Create a user in DB1 from that login, and grant the required database permissions or database-role membership to the user.

This approach scopes the group's permissions to DB1 rather than granting broad instance-level administrative permissions. For more information, see [Configure Microsoft Entra authentication for Azure SQL Managed Instance](#) and [CREATE LOGIN \(Transact-SQL\)](#).

QUESTION NO: 12

You are designing a streaming data solution that will ingest variable volumes of data.

You need to ensure that you can change the partition count after creation.

Which service should you use to ingest the data?

- A. Azure Event Hubs Standard
- B. Azure Stream Analytics
- C. Azure Data Factory
- D. Azure Event Hubs Dedicated

ANSWER: D

Explanation:

Azure Event Hubs Dedicated is the appropriate ingestion service because it supports increasing an event hub's partition count after the event hub is created. Partitions provide the parallelism used for event ingestion and consumption, so the ability to add partitions is valuable when a streaming workload has variable or growing throughput requirements. A dedicated Event Hubs cluster provides isolated capacity for an organization and supports this post-creation partition expansion capability, allowing the event stream to be scaled without having to recreate the event hub solely to use a larger partition count.

Partition planning remains important because partition count affects how event data is distributed and how consumers process it. Increasing the count adds capacity for future parallel processing, but it does not reduce the count or redistribute existing events already written to prior partitions. Therefore, using Azure Event Hubs Dedicated gives the solution a supported path to accommodate increased ingestion and processing parallelism as requirements evolve. Microsoft documents the partition behavior and the ability to increase partition counts for eligible Event Hubs offerings in [Event Hubs features and terminology](#). Additional capacity and scaling considerations for dedicated clusters are available in [Azure Event Hubs Dedicated overview](#).

QUESTION NO: 13

You have an Azure subscription. The subscription contains an instance of SQL Server on Azure Virtual Machines named SQL1 and an Azure Automation account named account1. You need to configure account1 to restart the SQL Server Agent service if the service stops. Which setting should you configure?

- A. Start/Stop VM
- B. Update management
- C. State configuration (DSC)
- D. Change tracking

ANSWER: C

Explanation:

State configuration (DSC) is the appropriate Azure Automation setting because it maintains a machine's intended configuration over time. A DSC configuration can define the SQL Server Agent Windows service as required to be present and in the `Running` state. When the Azure Automation State Configuration node checks in and performs its consistency evaluation, DSC detects configuration drift, including a required service that is stopped, and applies the configuration needed to return that service to its desired running state. This makes DSC suitable for enforcing the operational state of SQL Server Agent on SQL1 without requiring an administrator to manually monitor and restart it. The target VM must be registered as an Azure Automation State Configuration node and have the appropriate DSC configuration assigned. Microsoft describes the DSC Service resource as managing Windows service state, including ensuring that a service is running, and documents Azure Automation State Configuration as a service for consistently deploying and monitoring desired configurations on machines. See [DSC Service resource documentation](#) and [Azure Automation State Configuration overview](#).

QUESTION NO: 14

You have an Azure SQL database that contains a table named factSales. FactSales contains the columns shown in the following table.

Name	Data type
SalesID	Int
Product	Int
Total Number	Numeric(8,4)
Tax Number	Numeric(8,4)
SalesRep	Varchar(30)

FactSales has 6 billion rows and is loaded nightly by using a batch process.

Which type of compression provides the greatest space reduction for the database?

- A. page compression
- B. row compression
- C. columnstore compression
- D. columnstore archival compression

ANSWER: D

Explanation:

columnstore archival compression provides the greatest space reduction. A fact table with billions of rows is a typical analytical workload, and data loaded in nightly batches is well suited to columnstore storage. Columnstore indexes organize data by column into compressed rowgroups, which allows values from the same column to be encoded and compressed highly effectively, particularly for large fact tables that commonly contain repetitive values and numeric measures.

The `COLUMNSTORE_ARCHIVE` compression setting applies additional compression beyond standard `COLUMNSTORE` compression. Microsoft documents it as providing a higher compression ratio and therefore requiring less storage. This mode is intended where minimizing storage footprint is more important than query performance, because accessing archival-compressed rowgroups can require additional CPU resources. Since the question specifically asks for the greatest database space reduction, that tradeoff does not alter the choice. The scale of factSales makes the additional reduction especially meaningful.

For implementation details and compression behavior, see [Columnstore indexes overview](#) and [Columnstore index data compression](#).

QUESTION NO: 15

Note: This question is part of a series of questions that present the same scenario. Each question in the series contains a unique solution that might meet the stated goals. Some question sets might have more than one correct solution, while others might not have a correct solution.

After you answer a question in this section, you will NOT be able to return to it. As a result, these questions will not appear in the review screen.

You have an Azure Data Lake Storage account that contains a staging zone.

You need to design a daily process to ingest incremental data from the staging zone, transform the data by executing an R script, and then insert the transformed data into a data warehouse in Azure Synapse Analytics.

Solution: You use an Azure Data Factory schedule trigger to execute a pipeline that copies the data to a staging table in the data warehouse, and then uses a stored procedure to execute the R script.

Does this meet the goal?

- A. Yes

B. No

ANSWER: B

Explanation:

No is correct because the required transformation must execute an R script, and this processing should be performed by a supported compute environment orchestrated from Azure Data Factory rather than by a stored procedure in an Azure Synapse Analytics dedicated SQL pool. Azure Data Factory can schedule the daily incremental ingestion through a schedule trigger and can coordinate copying data from Azure Data Lake Storage. For custom transformation logic such as R code, Data Factory supports Custom activities, which run user-developed processing code on an associated compute service, such as Azure Batch. The R processing environment can read the staged input, execute the transformation, and write output that is then loaded into the Synapse Analytics warehouse.

This separates orchestration, transformation compute, and warehouse loading appropriately: Data Factory controls the daily workflow and dependencies, R runs where its runtime and required packages are available, and Synapse Analytics receives the transformed result. Microsoft documents Custom activity as the mechanism for running custom code in Data Factory pipelines; see [Custom activity in Azure Data Factory](#). For scheduling and recurring pipeline execution, see [Create a schedule trigger in Azure Data Factory](#).

QUESTION NO: 16

You need to identify the cause of the performance issues on SalesSQLDb1.

Which two dynamic management views should you use? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. sys.dm_pdw_nodes_tran_locks
- B. sys.dm_exec_compute_node_errors
- C. sys.dm_exec_requests
- D. sys.dm_cdc_errors
- E. sys.dm_pdw_nodes_os_wait_stats
- F. sys.dm_tran_locks

ANSWER: A E

Explanation:

`sys.dm_pdw_nodes_tran_locks` and `sys.dm_pdw_nodes_os_wait_stats` are the appropriate diagnostic views for investigating performance problems in an Azure Synapse Analytics dedicated SQL pool.

`sys.dm_pdw_nodes_tran_locks` exposes lock requests across the compute nodes, including the requested lock mode, lock status, session information, and the resource involved. This lets an administrator identify concurrency pressure and determine whether sessions are waiting on locks held by other activity.

`sys.dm_pdw_nodes_os_wait_stats` provides cumulative wait statistics for each distribution and compute node. Wait types and their durations help identify where query execution time is being spent, such as waits related to CPU scheduling, I/O, data movement, or locking. Reviewing these waits alongside active lock information provides evidence for the underlying source of a slowdown rather than relying only on query duration.

Microsoft documents the Synapse node-level lock DMV at [sys.dm_pdw_nodes_tran_locks](#) and the node-level wait-statistics DMV at [sys.dm_pdw_nodes_os_wait_stats](#).

QUESTION NO: 17

You have an Azure SQL database.

You identify a long running query.

You need to identify which operation in the query is causing the performance issue.

What should you use to display the query execution plan in Microsoft SQL Server Management Studio (SSMS)?

- A. Live Query Statistics
- B. an estimated execution plan
- C. an actual execution plan
- D. Client Statistics

ANSWER: C

Explanation:

Use **an actual execution plan**. In SSMS, selecting *Include Actual Execution Plan* causes SQL Server to execute the query and then return the plan that was used, together with runtime execution information. This is especially useful for investigating a long-running query because the graphical plan shows the operators involved and includes actual row counts, elapsed and CPU-related runtime details, warnings, and other operator properties. Comparing actual versus estimated row counts can reveal cardinality-estimation discrepancies that lead SQL Server to select an inefficient operation, join strategy, memory grant, or access method. Operator-level metrics and warnings help focus investigation on the part of the plan consuming resources or producing unexpectedly many rows. The actual plan is generated only after execution completes, so it represents the execution behavior of the submitted query rather than a prediction made before execution. In SSMS, open the query window, choose *Query > Include Actual Execution Plan*, execute the query, and inspect the *Execution Plan* result tab. Azure SQL Database supports displaying actual execution plans in SSMS. See [Display an Actual Execution Plan](#) and [Execution Plans](#).

QUESTION NO: 18

You need to identify the event_file target for monitoring DB3 after the migration to Azure SQL Database. The solution must meet the management requirements.

What should you use as the event_file target?

- A. an Azure SQL database
- B. an Azure Blob Storage container
- C. a SQL Server filegroup
- D. an Azure Files share

ANSWER: B

Explanation:

an Azure Blob Storage container is the appropriate event_file target for an Extended Events session in Azure SQL Database. Azure SQL Database does not provide access to a local server file system for persisting Extended Events output. Instead, the event_file target writes event data to a file in Azure Blob Storage, allowing the monitoring data to remain durable and accessible after the session runs or the database is migrated. The target configuration uses a Blob Storage URL, typically for a container that holds the generated .xel event files, and authentication is configured through a database-scoped credential or managed identity as appropriate. Administrators can subsequently query or review the event files by using supported Extended Events functions and tools. This design provides storage that is independent of the Azure SQL Database compute instance and is suited to operational monitoring, retention, and centralized access requirements. Microsoft documents Azure Blob Storage as the storage destination required when using the event_file target with Extended Events in Azure SQL Database. See [Extended Events differences for Azure SQL Database](#) and [Targets for Extended Events](#).

QUESTION NO: 19

You have an Azure SQL managed instance named SQL1 and two Azure web apps named App1 and App2.

You need to limit the number of IOPs that App2 queries generate on SQL1.

Which two actions should you perform on SQL1? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Enable query optimizer fixes.
- B. Enable Resource Governor.
- C. Enable parameter sniffing.
- D. Create a workload group.
- E. Configure In-memory OLTP.
- F. Run the Database Engine Tuning Advisor.
- G. Reduce the Max Degree of Parallelism value.

ANSWER: B D

Explanation:

Enable Resource Governor and create a workload group. Resource Governor provides the SQL Server Database Engine mechanism for governing resources consumed by sessions and requests. After its configuration is applied, it can direct connections into defined workload groups using a classifier function. In this scenario, the classifier can identify connections from App2, such as through its application name in the connection string, and assign those sessions to the intended workload group.

The workload group is where the policy that governs App2 is defined. Its configuration supports an I/O limit through `MAX_IOPS_PER_VOLUME`, allowing the managed instance to constrain the IOPS generated by requests assigned to that group. The workload group can be associated with the applicable resource pool and then activated as part of the Resource Governor configuration. This scopes the control to App2's classified workload rather than imposing an indiscriminate limit on all application activity, so App1 can continue to use the normal/default workload path.

Microsoft documents Resource Governor architecture, including classifier functions and workload groups, at [Resource Governor](#). The workload-group syntax and the IOPS configuration setting are documented at [CREATE WORKLOAD GROUP](#).

QUESTION NO: 20 - (HOTSPOT)

HOTSPOT

You are performing exploratory analysis of bus fare data in an Azure Data Lake Storage Gen2 account by using an Azure Synapse Analytics serverless SQL pool.

You execute the Transact-SQL query shown in the following exhibit.

```

SELECT
    payment_type,
    SUM(fare_amount) AS fare_total
FROM OPENROWSET(
    BULK 'csv/busfare/tripdata_2020*.csv',
    DATA_SOURCE = 'BusData',
    FORMAT = 'CSV', PARSER_VERSION = '2.0',
    FIRSTROW = 2
)
WITH (
    payment_type INT 10,
    fare_amount FLOAT 11
) AS nyc
GROUP BY payment_type
ORDER BY payment_type;

```

Use the drop-down menus to select the answer choice that completes each statement based on the information presented in the graphic.

Hot Area:

Answer Area

The query results include only [answer choice]
in the csv/busfare folder.

	▼
CSV files in the tripdata_2020 subfolder	
files that have files names beginning with "tripdata_2020"	
CSV files that have file names containing "tripdata_202"	
CSV files that have file named beginning with "tripdata_2020"	

The query assumes that the first row in a CSV file is
[answer choice] row.

	▼
a header	
a data	
an empty	

ANSWER:

Answer Area

The query results include only [answer choice]
in the csv/busfare folder.

	▼
CSV files in the tripdata_2020 subfolder	
files that have files names beginning with "tripdata_2020"	
CSV files that have file names containing "tripdata_202"	
CSV files that have file named beginning with "tripdata_2020"	

The query assumes that the first row in a CSV file is
[answer choice] row.

	▼
a header	
a data	
an empty	

Explanation:

The query reads CSV data through the serverless SQL pool by using `OPENROWSET`. In the BULK file path, the value includes the `tripdata_2020` prefix followed by a wildcard. That pattern selects CSV files in the specified `csv/busfare` location whose file names begin with `tripdata_2020`. This lets one query cover multiple matching fare-data files, such as files differentiated by a remaining date or period portion of their names.

The query also specifies `FIRSTROW = 2`. For delimited text files, this tells the reader that data loading begins with the second row. As a result, the first row is expected to contain the column names rather than a bus-fare record. The first row is therefore assumed to be a header row, while all rows beginning with row two are interpreted according to the columns defined by the query. This is a common pattern for exploratory queries over CSV files because it prevents header text from being returned as data.

Microsoft documents that wildcard patterns can be used in the file path supplied to `OPENROWSET`, and that `FIRSTROW` controls the first row read from a delimited file. See [How to use OPENROWSET in serverless SQL pool](#) and [OPENROWSET \(BULK\) Transact-SQL reference](#).

QUESTION NO: 21

Note: This question is part of a series of questions that present the same scenario. Each question in the series contains a unique solution that might meet the stated goals. Some question sets might have more than one correct solution, while others might not have a correct solution.

After you answer a question in this section, you will NOT be able to return to it. As a result, these questions will not appear in the review screen.

You have SQL Server 2019 on an Azure virtual machine.

You are troubleshooting performance issues for a query in a SQL Server instance.

To gather more information, you query `sys.dm_exec_requests` and discover that the wait type is `PAGELATCH_UP` and the `wait_resource` is `2:3:905856`.

You need to improve system performance.

Solution: You reduce the use of table variables and temporary tables.

Does this meet the goal?

A. Yes

B. No

ANSWER: A

Explanation:

Yes. The wait resource format is database ID:file ID:page ID. Database ID 2 is `tempdb`, and page 905856 is a page allocation metadata page because it falls on a PFS-page interval. A `PAGELATCH_UP` wait on this type of `tempdb` page indicates in-memory latch contention while concurrent sessions allocate or deallocate pages. It is not an I/O latency issue.

Temporary tables and table variables are created in `tempdb`. Workloads that create, populate, modify, or drop these objects frequently increase allocation activity and can contend for PFS, GAM, and SGAM metadata pages. Reducing unnecessary use of temporary tables and table variables therefore reduces the allocation/deallocation pressure that produces this wait pattern and can improve throughput and query response time. This is an appropriate workload-level mitigation, particularly when the objects are used repeatedly in highly concurrent code paths. Microsoft identifies reducing allocation contention in `tempdb` as a key response to these waits; operational measures such as multiple equally sized `tempdb` data files may also be considered after workload review. See [Recommendations to reduce allocation contention in SQL Server](#) and [tempdb database documentation](#).

QUESTION NO: 22

You have an Azure SQL database named DB1.

You have an Azure App Service web app that uses a system-assigned managed identity.

You need to enable the web app to authenticate to DB1 by using its managed identity. The solution must minimize credential management.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Configure a Microsoft Entra administrator for the logical server.
- B. Create a contained database user for the managed identity in DB1.
- C. Create a SQL authentication login for the web app.
- D. Add the managed identity to the db_owner role in master.
- E. Enable Transparent Data Encryption (TDE) for DB1.

ANSWER: A B

Explanation:

You must configure a Microsoft Entra administrator for the logical server and create a contained database user for the managed identity. The Microsoft Entra administrator is authorized to create Microsoft Entra principals in Azure SQL Database. A database user created from the external provider maps the web app's managed identity to a principal in DB1.

After the user is created, grant only the database permissions required by the application. A SQL authentication login and a stored password are unnecessary because the managed identity obtains Microsoft Entra tokens automatically.

For more information, see [Microsoft Entra authentication for Azure SQL Database](#) and [Managed identities for App Service](#).

QUESTION NO: 23

You have a SQL pool in Azure Synapse that contains a table named dbo.Customers. The table contains a column named Email.

You need to prevent nonadministrative users from seeing the full email addresses in the Email column. The users must see values in a format of aXXX@XXXX.com instead.

What should you do?

- A. From the Azure portal, set a mask on the Email column.
- B. From the Azure portal, set a sensitivity classification of Confidential for the Email column.
- C. From Microsoft SQL Server Management Studio, set an email mask on the Email column.
- D. From Microsoft SQL Server Management Studio, grant the SELECT permission to the users for all the columns in the dbo.Customers table except Email.

ANSWER: C

Explanation:

From Microsoft SQL Server Management Studio, set an email mask on the Email column. Dynamic data masking is supported for Azure Synapse dedicated SQL pools and can be configured by executing Transact-SQL from SQL Server Management Studio. The built-in `email()` masking function is designed specifically for email-address values: it exposes the first character of the address and returns the rest in the standardized masked form `aXXX@XXXX.com`. For example, a user querying a value such as `alex@contoso.com` receives a masked result beginning with `aXXX@XXXX.com`, rather than the underlying address.

An administrator can apply this behavior with an `ALTER TABLE ... ALTER COLUMN ... ADD MASKED WITH (FUNCTION = 'email()')` statement against `dbo.Customers.Email`. Dynamic data masking is applied at query time, so the original value remains stored in the table while users without permission to view unmasked data receive the masked representation. Administrative users and principals granted the appropriate unmasking permission can still access the original data when required. Microsoft documents dynamic data masking for Synapse SQL pools and the built-in email masking function in its SQL security documentation: [Dynamic data masking in Azure Synapse Analytics](#) and [Dynamic Data Masking](#).

QUESTION NO: 24

You have an Azure SQL database. The database contains a table that uses a columnstore index and is accessed infrequently.

You enable columnstore archival compression.

What are two possible results of the configuration? Each correct answer presents a complete solution.

NOTE: Each correct selection is worth one point.

- A. Queries that use the index will consume more disk I/O.
- B. Queries that use the index will retrieve fewer data pages.
- C. The index will consume more disk space.
- D. The index will consume more memory.
- E. Queries that use the index will consume more CPU resources.

ANSWER: B E

Explanation:

Columnstore archival compression is intended for columnstore data that is accessed infrequently and for which minimizing storage consumption is more important than achieving the fastest query performance. It applies an additional compression algorithm to columnstore segments, producing a smaller on-disk representation than standard columnstore compression. Therefore, queries that use the index will retrieve fewer data pages because the compressed index occupies less storage and less data must be read from storage for applicable operations. This can be valuable for historical or cold data, where reduced storage and I/O volume are desirable.

The additional compression also requires extra processing when SQL Server or Azure SQL Database stores and reads the compressed data. As a result, queries that use the index will consume more CPU resources to decompress archival-compressed columnstore data before it can be processed. This storage-versus-CPU tradeoff is precisely why archival compression is recommended for infrequently accessed data rather than performance-sensitive workloads. Microsoft documents that archival compression further reduces columnstore data size but increases the CPU resources needed to compress and decompress it. See [Columnstore indexes: data compression](#) and [Columnstore indexes overview](#).

QUESTION NO: 25

You receive numerous alerts from Azure Monitor for an Azure SQL Database instance.

You need to reduce the number of alerts. You must only receive alerts if there is a significant change in usage patterns for an extended period.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Set Threshold Sensitivity to High
- B. Set the Alert logic threshold to Dynamic

- C. Set the Alert logic threshold to Static
- D. Set Threshold Sensitivity to Low
- E. Set Force Plan to On

ANSWER: B D

Explanation:

Set the Alert logic threshold to Dynamic because dynamic thresholds use Azure Monitor machine learning to evaluate a metric against its own historical behavior rather than against one fixed numeric limit. This approach establishes an expected range for the metric and adapts as it identifies normal recurring patterns, including hourly, daily, and weekly seasonality. It is therefore appropriate when alerting should be based on a meaningful departure from the database's established usage pattern.

Set Threshold Sensitivity to Low to require a larger deviation from the learned normal range before the alert is triggered. Low sensitivity places the dynamic thresholds farther from the expected metric values, reducing notifications caused by small fluctuations or short-lived variation. Used with a dynamic threshold, this setting focuses alerting on substantial, sustained anomalies and helps reduce alert noise while retaining detection of significant usage changes. Azure Monitor also supports configuring evaluation frequency and failing periods when creating the alert rule, allowing the rule to require repeated violations before notification.

Microsoft documents dynamic thresholds and sensitivity behavior at [Dynamic thresholds in Azure Monitor](#). For metric alert evaluation and failing-period configuration, see [Create or edit a metric alert rule](#).

QUESTION NO: 26 - (HOTSPOT)

You need to recommend which service and target endpoint to use when migrating the databases from SVR1 to Instance1. The solution must meet the availability requirements.

What should you recommend? To answer, select the appropriate options in the answer area.

NOTE Each correct selection is worth one point.

Answer Area

Migration service:

Target endpoint:

ANSWER:

Answer:

Migration service: Managed Instance link
 Log Replay Service (LRS)
Managed Instance link
 SQL Data Sync

Target endpoint: A VNet-local endpoint
 A private endpoint
 A public endpoint
A VNet-local endpoint

Answer:

Explanation:

Basic Concept: This question tests high availability and disaster recovery design for Azure SQL, SQL Server on Azure VMs, and regional failure scenarios. **Why the selected answer is Correct:** The selected values fit this scenario because they apply the required Azure SQL configuration at the correct scope and sequence: You need to recommend which service and target endpoint to use when migrating the databases from SVR1 to Instance1. **Why the alternate choices are Wrong:** The alternate choices do not satisfy the stated availability design because quorum and listener resources must match Azure failover-cluster behavior and the required SLA.

Explanation:

Use **Managed Instance link** with a **VNet-local endpoint**. Managed Instance link is designed for moving SQL Server databases to Azure SQL Managed Instance while preserving a high-availability-oriented migration path. It establishes a link between the SQL Server source and the managed instance, continuously replicating the selected databases. This permits the target to remain current during the migration period and supports a planned cutover when the organization is ready. That approach is especially suitable where database availability must be maintained instead of relying on a single offline export, restore, or one-time copy operation.

The SQL Server source must be able to communicate with the Azure SQL Managed Instance that is participating in the link. A VNet-local endpoint supplies an endpoint address reachable through the managed instance virtual network. This keeps the replication communication on the private Azure network path and provides the connectivity required by the link architecture. Before creating the link, the required network routing, name resolution, and firewall or network-security rules must permit the source SQL Server and managed instance to reach each other on the documented ports.

After synchronization has caught up, the migration can be completed by failing over the linked databases to the managed instance. Microsoft documents the feature, its prerequisites, networking requirements, and the failover process in [Managed Instance link overview](#) and [Prepare for Managed Instance link](#).

QUESTION NO: 27 - (SIMULATION)

Task 9

You need to ensure that when non-administrative users query the SalesLT.Customer table in db1, email addresses are obscured. For example, an email address of alice@contoso.com must appear as aXXX@XXXX.com.

You may need to use SQL Server Management Studio and the Azure portal.

ANSWER: See the explanation for the answer

Explanation:

Configure Dynamic Data Masking for SalesLT.Customer.EmailAddress in db1 using the built-in email() mask. This displays an address such as alice@contoso.com as aXXX@XXXX.com to users who do not have UNMASK permission.

In SSMS, connect to db1 and run:

```
ALTER TABLE [SalesLT].[Customer]
```

```
ALTER COLUMN [EmailAddress]  
ADD MASKED WITH (FUNCTION = 'email()');
```

Alternatively, in the Azure portal, open `db1`, select **Security > Dynamic Data Masking**, and add a rule with:

Schema: SalesLT

Table: Customer

Column: EmailAddress

Masking format: Email/email()

Ensure normal users have not been granted `UNMASK`; otherwise, they will see the original email values.

QUESTION NO: 28

You create five Azure SQL Database instances on the same logical server.

In each database, you create a user for an Azure Active Directory (Azure AD) user named User1.

User1 attempts to connect to the logical server by using Azure Data Studio and receives a login error.

You need to ensure that when User1 connects to the logical server by using Azure Data Studio, User1 can see all the databases.

What should you do?

- A. Create User1 in the master database.
- B. Assign User1 the `db_datareader` role for the master database.
- C. Assign User1 the `db_datareader` role for the databases that User1 creates.
- D. Grant `SELECT` on `sys.databases` to public in the master database.

ANSWER: A

Explanation:

Create User1 in the master database. Azure SQL Database uses the logical server's `master` database as the initial connection context when a client connects without specifying a target database. Although User1 already exists as a contained database user in each individual database, that does not provide the user with an identity in `master` for the initial server-level connection and database discovery process. Creating the corresponding Microsoft Entra ID user in `master` enables User1 to authenticate through the logical server and connect successfully from Azure Data Studio. Once connected, Azure SQL Database can display the databases to which User1 has access; in this case, the five databases where the user was created. The user in `master` is a database user mapped to the same Microsoft Entra ID principal, rather than a traditional SQL Server login. An administrator connected to `master` can create it by using `CREATE USER [User1] FROM EXTERNAL PROVIDER`. Microsoft documents that Microsoft Entra users who need to connect without a specific database name require a user account in the logical server's `master` database. See [Microsoft Learn: Manage logins and users in Azure SQL Database](#) and [Microsoft Learn: Configure Microsoft Entra authentication](#).

QUESTION NO: 29

You have an Azure Stream Analytics job.

You need to ensure that the job has enough streaming units provisioned.

You configure monitoring of the SU % Utilization metric.

Which two additional metrics should you monitor? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Late Input Events

- B. Out of order Events
- C. Backlogged Input Events
- D. Watermark Delay
- E. Function Events

ANSWER: C D

Explanation:

Backlogged Input Events and **Watermark Delay** provide the key workload-impact signals to monitor alongside SU % Utilization. Backlogged Input Events measures input events that the Stream Analytics job has received but has not yet processed. A nonzero value, especially one that persists or rises over time, indicates that processing capacity is not keeping pace with the incoming event rate. This is a direct indicator that the job may require additional streaming units.

Watermark Delay measures the gap between the time of the latest input event processed and the current wall-clock time. An increasing delay can show that the job is falling behind while processing streaming data. Monitoring this metric together with backlog and utilization helps distinguish a sustained capacity issue from a short-lived workload spike. When utilization is high and either backlog or watermark delay trends upward, scaling the job by increasing its streaming units is appropriate.

Microsoft recommends monitoring these job metrics to identify whether an Azure Stream Analytics job is keeping up with input data and to support scaling decisions. See [Monitor Azure Stream Analytics jobs](#) and [Scale Azure Stream Analytics jobs](#).

QUESTION NO: 30

You have SQL Server on an Azure virtual machine that contains a database named DB1. DB1 contains a table named CustomerPII.

You need to record whenever users query the CustomerPII table.

Which two options should you enable? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. server audit specification
- B. SQL Server audit
- C. database audit specification
- D. a server principal

ANSWER: B C

Explanation:

To record queries against the CustomerPII table on SQL Server running in an Azure virtual machine, enable a **SQL Server audit** and a **database audit specification**. A SQL Server audit is the audit object that defines where audit events are written, such as a file, the Windows Application event log, or the Windows Security event log. It must be enabled before SQL Server can persist audit records.

A database audit specification is then created in DB1 and enabled to define the database-level events to capture. For this requirement, it can contain an object-level audit action for SELECT on dbo.CustomerPII, which records access attempts when users query that table. This pairing separates the audit event destination from the database-specific actions that generate events, allowing the organization to capture and review access to sensitive customer data.

Microsoft documents that SQL Server Audit uses an Audit object together with server or database audit specifications, and that object-level database audit actions can audit SELECT activity on specific securables. See [SQL Server Audit \(Database Engine\)](#) and [SQL Server audit action groups and actions](#).

QUESTION NO: 31

You have two Azure virtual machines named Server 1 and Server2 that run Windows Server 2022 and are joined to an Active Directory Domain Services (AD DS) domain named contoso.com.

Both virtual machines have a default instance of Microsoft SQL Server 2019 installed. Server1 is configured as a master server, and Server2 is configured as a target server.

On Server1, you create a proxy account named contoso\sqlproxy.

You need to ensure that the SQL Server Agent job steps can be downloaded from Server1 and run on Server2.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. On Server2, grant the contoso\sqlproxy account the Impersonate a client after authentication user right.
- B. On Server2, grant the contoso\sqlproxy account the Access this computer from the network user right.
- C. On Server2, create a proxy account.
- D. On Server1, set the AllowDownloadedJobsToMatchProxyName registry entry to 1.
- E. On Server2, set the AllowDownloadedJobsToMatchProxyName registry entry to 1.

ANSWER: C E

Explanation:

Creating a proxy account on Server2 is required because a target server runs the downloaded SQL Server Agent job locally. The target must therefore have a local SQL Server Agent proxy with the same proxy name as the proxy referenced by the job step. That proxy supplies the credential and subsystem authorization used when the job executes on the target instance.

Setting the `AllowDownloadedJobsToMatchProxyName` registry entry to 1 on Server2 enables SQL Server Agent to associate a proxy referenced in a downloaded multiserver job with the local target-server proxy that has the matching name. This setting supports the intended master-server and target-server job-distribution behavior while ensuring that execution on the target uses its locally configured proxy definition. After configuring the proxy and setting, restart the SQL Server Agent service on the target so that the Agent reads the updated registry configuration.

Microsoft documents proxy accounts as SQL Server Agent objects that provide job-step security contexts, and documents multiserver administration as the mechanism through which jobs are downloaded from a master server to target servers. See [Create a SQL Server Agent proxy](#) and [Automated administration across an enterprise](#).

QUESTION NO: 32

You have SQL Server on an Azure virtual machine that contains a database named DB1.

You have an application that queries DB1 to generate a sales report.

You need to see the parameter values from the last time the query was executed.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Enable Last_Query_Plan_Stats in the master database
- B. Enable Lightweight_Query_Profiling in DB1
- C. Enable Last_Query_Plan_Stats in DB1
- D. Enable Lightweight_Query_Profiling in the master database

E. Enable PARAMETER_SNIFFING in DB1

ANSWER: B C

Explanation:

Enable `LAST_QUERY_PLAN_STATS` in DB1 to make SQL Server retain the last known actual execution-plan statistics for queries executed in that database. This database-scoped configuration enables the `sys.dm_exec_query_plan_stats` dynamic management function to return a plan comparable to an actual execution plan for the most recent execution of a cached query. The returned plan XML can include runtime information, including the parameter values associated with that execution, allowing the sales-report query to be investigated after it has run.

Enable `LIGHTWEIGHT_QUERY_PROFILING` in DB1 so SQL Server collects runtime query-profiling information with substantially lower overhead than traditional execution profiling. Because the application queries DB1, the relevant database-scoped settings must be enabled in DB1. Together, lightweight profiling and last-query-plan-statistics collection provide the runtime plan data needed to inspect the latest execution and its parameter details. Microsoft documents `LAST_QUERY_PLAN_STATS` as the setting that exposes last actual-plan statistics through `sys.dm_exec_query_plan_stats`, and documents lightweight profiling as the infrastructure for collecting runtime profiling data. See [sys.dm_exec_query_plan_stats](#) and [Query profiling infrastructure](#).

QUESTION NO: 33

Note: This question is part of a series of questions that present the same scenario. Each question in the series contains a unique solution that might meet the stated goals. Some question sets might have more than one correct solution, while others might not have a correct solution.

After you answer a question in this section, you will NOT be able to return to it. As a result, these questions will not appear in the review screen.

You have an Azure Synapse Analytics dedicated SQL pool that contains a table named Table1.

You have files that are ingested and loaded into an Azure Data Lake Storage Gen2 container named container1.

You plan to insert data from the files into Table1 and transform the data. Each row of data in the files will produce one row in the serving layer of Table1.

You need to ensure that when the source data files are loaded to container1, the DateTime is stored as an additional column in Table1.

Solution: In an Azure Synapse Analytics pipeline, you use a Get Metadata activity that retrieves the DateTime of the files.

Does this meet the goal?

A. Yes

B. No

ANSWER: B

Explanation:

No is correct because retrieving a file DateTime through a Get Metadata activity does not itself add that value to rows written to Table1. Get Metadata returns information about a dataset, such as file metadata, to the pipeline's activity output. The pipeline must subsequently use that output in an operation that writes data to the dedicated SQL pool.

To meet the stated requirement, the retrieved DateTime would need to be mapped into the destination as a value for an additional column for every loaded row. For example, a Copy activity can use its additional-columns capability to add a source-independent value to each copied row, with the value supplied dynamically from pipeline metadata. Alternatively, a data flow or SQL transformation could explicitly project the DateTime into the destination schema. Since the proposed solution stops at obtaining metadata and does not configure a load or transformation that persists it in Table1, it does not ensure that the DateTime is stored in the serving-layer rows. See [Get Metadata activity in Azure Data Factory and Synapse Analytics](#) and [Add additional columns during copy](#).

QUESTION NO: 34

You have an Azure SQL database named DB1.

A user named User 1 has an Azure AD account.

You need to provide User1 with the ability to add and remove columns from the tables in DB1. The solution must use the principle of least privilege.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point

- A. Assign the database user the db_ddladmin role.
- B. Assign the database user the db.owner role.
- C. Create a contained database user.
- D. Create a login and an associated database user.

ANSWER: A C

Explanation:

Use a contained database user for the Microsoft Entra ID (formerly Azure AD) identity, and assign that user to the db_ddladmin fixed database role. In Azure SQL Database, users are created at the database level, and a Microsoft Entra user can be created as a contained user with `CREATE USER [user@domain] FROM EXTERNAL PROVIDER`. This maps the authenticated Entra identity to a principal in DB1 without granting unnecessary server-level permissions.

The db_ddladmin role is the least-privileged built-in database role appropriate for changing table definitions. Its members can run Data Definition Language statements in the database, including `ALTER TABLE ... ADD` to add columns and `ALTER TABLE ... DROP COLUMN` to remove columns. Therefore, it provides the required schema-modification capability while keeping privileges scoped to DDL administration rather than broad database ownership.

Microsoft documents Microsoft Entra contained database users for Azure SQL Database at [Microsoft Entra service principals with Azure SQL](#) and documents the permissions of fixed database roles, including db_ddladmin, at [Database-level roles](#).

QUESTION NO: 35

You have an on-premises Microsoft SQL Server 2022 instance named SQL1. You have an Azure subscription that contains an Azure SQL managed instance named SQLM11. You need to configure a bidirectional disaster recovery solution between SQL1 and SQLM11. What should you include in the solution?

- A. Log Replay Service (LRS)
- B. the Managed Instance link
- C. Azure Database Migration Service
- D. a failover group

ANSWER: B

Explanation:

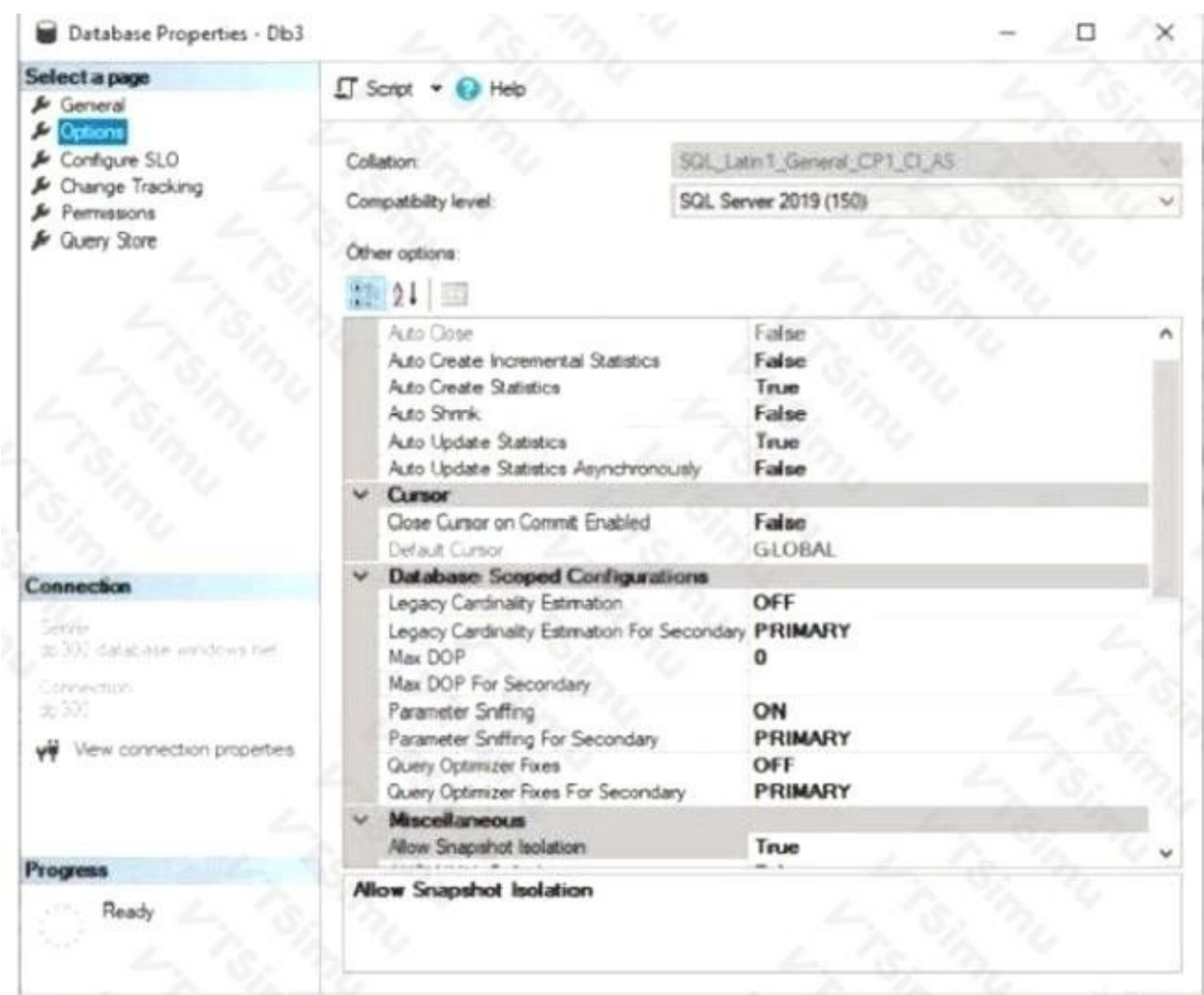
The Managed Instance link is the appropriate solution because it is designed to establish a link between SQL Server 2022 and Azure SQL Managed Instance for near-real-time data replication and disaster-recovery scenarios. The feature uses distributed availability groups to replicate databases from the SQL Server instance to the managed instance. It also supports bidirectional failover: after an initial failover to Azure SQL Managed Instance, the linked environment can be used to fail back to SQL Server when the on-premises environment is available again. This makes it suitable when both the on-premises SQL Server 2022 instance and Azure SQL Managed Instance must participate in a continuing disaster-recovery topology rather

than a one-time migration. The linked databases remain individually manageable, and the solution is intended to provide an Azure-based standby environment while preserving the ability to return workloads to SQL Server. Before configuring the link, ensure that the SQL Server edition, version, networking, and required availability group configuration meet the feature prerequisites. Microsoft documents the architecture, requirements, and failover operations in [Managed Instance link overview](#) and provides deployment guidance in [Prepare SQL Server for Managed Instance link](#).

QUESTION NO: 36

You have an Azure SQL database named DB3.

You need to provide a user named DevUser with the ability to view the properties of DB3 from Microsoft SQL Server Management Studio (SSMS) as shown in the exhibit. (Click the Exhibit tab.)



Which Transact-SQL command should you run?

- A. GRANT SHOWPLAN TO DevUser
- B. GRANT VIEW DEFINITION TO DevUser
- C. GRANT VIEW DATABASE STATE TO DevUser
- D. GRANT SELECT TO DevUser

ANSWER: C

Explanation:

GRANT VIEW DATABASE STATE TO DevUser is the appropriate command because the State information displayed by the SSMS Database Properties dialog is based on database state and metadata exposed through database-scoped system

views and dynamic management views. `VIEW DATABASE STATE` grants a database principal permission to view the current state of the database and the database-scoped DMVs that SSMS can use to populate state-related properties. The permission is granted within DB3, so DevUser can inspect the relevant status information for that database without being granted broader server-level visibility or permissions to alter database configuration.

This follows the SQL Server and Azure SQL permission model: database state permissions control access to information about internal database activity, state, and operational details. In Azure SQL Database, granting the database-scoped permission is particularly appropriate because access is managed at the individual database level. Run the statement while connected to DB3 and ensure that DevUser is a database user in DB3. Microsoft documents `VIEW DATABASE STATE` as a database permission and documents the permissions required for dynamic management views in the SQL Database Engine. See [Database Engine permissions](#) and [System dynamic management views](#).

QUESTION NO: 37

You have an Azure subscription that contains the resources shown in the following table.

Name	Type
App1	Azure web app
db1	Azure SQL database in the serverless tier

App1 experiences transient connection errors and timeouts when it attempts to access db1 after extended periods of inactivity. You need to modify db1 to resolve the issues experienced by App1 as soon as possible, without considering immediate costs. What should you do?

- A. Increase the number Of vCores allocated to db1.
- B. Disable auto-pause delay for db1.
- C. Decrease the auto-pause delay for db1.
- D. Enable automatic tuning for db1.

ANSWER: B

Explanation:

Disable auto-pause delay for db1. is correct because the described behavior is characteristic of an Azure SQL Database configured in the serverless compute tier with auto-pause enabled. After the database has been inactive for its configured delay period, Azure pauses it to reduce compute charges. When App1 next attempts to connect, the database must resume before it can process requests. This resume operation can introduce a delay and can lead to application connection timeouts or transient connectivity failures when the application does not retry for long enough.

Disabling auto-pause keeps the serverless database online even during idle periods, so it is immediately available when App1 resumes activity. This directly addresses the requirement to eliminate the inactivity-related connection delays as quickly as possible. The tradeoff is that compute billing continues while the database is idle, but the scenario explicitly states that immediate cost is not a consideration. Azure SQL Database serverless still provides compute scaling based on workload; disabling auto-pause specifically prevents the transition to the paused state that causes the observed startup latency. For details, see [Azure SQL Database serverless tier overview](#) and [serverless auto-pause behavior](#).

QUESTION NO: 38

You have an Azure virtual machine named VM1 on a virtual network named VNet1. Outbound traffic from VM1 to the internet is blocked.

You have an Azure SQL database named SqlDb1 on a logical server named SqlSrv1.

You need to implement connectivity between VM1 and SqlDb1 to meet the following requirements:

Ensure that all traffic to the public endpoint of SqlSrv1 is blocked.

Minimize the possibility of VM1 exfiltrating data stored in SqlDb1.

What should you create on VNet1?

- A. a VPN gateway
- B. a service endpoint
- C. a private endpoint
- D. an ExpressRoute gateway

ANSWER: C

Explanation:

A private endpoint is the appropriate resource to create on VNet1. It assigns a private IP address from a VNet subnet to the Azure SQL logical server, allowing VM1 to reach SqlDb1 through Azure Private Link rather than through the server's public endpoint. The connection remains on the Microsoft network backbone and can be governed by network security controls applicable to the virtual network.

To fully meet the requirement that traffic to the public endpoint be blocked, public network access on SqlSrv1 should be disabled after private-endpoint connectivity and DNS resolution are configured. Azure Private DNS integration enables the SQL server's normal fully qualified domain name to resolve to the private endpoint's IP address from VNet1. This avoids requiring VM1 to use public routing and reduces the data-exfiltration path by limiting database access to the approved private network connection. Azure SQL Database private endpoints are specifically designed to provide private, VNet-based access to logical SQL servers. See [Azure SQL Database private endpoint overview](#) and [Azure Private Endpoint overview](#).

QUESTION NO: 39

You are designing a star schema for a dataset that contains records of online orders. Each record includes an order date, an order due date, and an order ship date.

You need to ensure that the design provides the fastest query times of the records when querying for arbitrary date ranges and aggregating by fiscal calendar attributes.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Create a date dimension table that has a DateTime key.
- B. Create a date dimension table that has an integer key in the format of YYYYMMDD.
- C. Use built-in SQL functions to extract date attributes.
- D. Use integer columns for the date fields.
- E. Use DateTime columns for the date fields.

ANSWER: B D

Explanation:

Create a date dimension table that has an integer key in the format of YYYYMMDD, and use integer columns for the date fields. A dedicated date dimension centralizes calendar and fiscal attributes, such as fiscal year, fiscal period, quarter, month, and day, so queries can filter or aggregate by those precomputed attributes rather than deriving them repeatedly from order dates. The integer YYYYMMDD surrogate-style key is compact, naturally supports equality joins between the fact table and the date dimension, and is appropriate for range predicates when the values are stored in chronological YYYYMMDD order. The order date, due date, and ship date columns in the orders fact table should each store the corresponding integer date key. Each is then a role-playing relationship to the same date dimension, enabling efficient

analysis for any of the business date roles while retaining one consistent fiscal calendar definition. This follows star-schema guidance to use dimension tables for descriptive attributes and fact tables for foreign keys and measures. Microsoft also recommends modeling date-related attributes in a date table to support consistent time-based analysis. See [Microsoft guidance on star schema design](#) and [Microsoft guidance on date tables](#).

QUESTION NO: 40

You have an Azure SQL database named DB1.

A user named User 1 has an Azure AD account.

You need to provide User1 with the ability to add and remove columns from the tables in DB1. The solution must use the principle of least privilege.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point

- A. Assign the database user the db_ddladmin role.
- B. Assign the database user the db.owner role.
- C. Create a contained database user.
- D. Create a login and an associated database user.

ANSWER: A C

Explanation:

Assign the database user the db_ddladmin role and create a contained database user. In Azure SQL Database, a Microsoft Entra ID user is created in the database as a contained database user, typically by using `CREATE USER [user@domain] FROM EXTERNAL PROVIDER`. This maps the Microsoft Entra identity to a database principal without requiring a traditional server-level SQL login. After the user exists in DB1, membership in the fixed database role db_ddladmin grants authority to execute data definition language commands in that database. Commands such as `ALTER TABLE ... ADD` and `ALTER TABLE ... DROP COLUMN` are therefore available, allowing the user to add and remove table columns as required. This is consistent with least privilege because the permission scope is limited to DDL activity within the database rather than granting broad database ownership or administrative control. Microsoft documents creating Microsoft Entra contained database users in [Microsoft Entra authentication for Azure SQL Database](#) and documents the permissions of fixed database roles, including db_ddladmin, in [Database-level roles](#).

QUESTION NO: 41

You have a new Azure SQL database. The database contains a column that stores confidential information.

You need to track each time values from the column are returned in a query. The tracking information must be stored for 365 days from the date the query was executed.

Which three actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Turn on auditing and write audit logs to an Azure Storage account.
- B. Add extended properties to the column.
- C. Turn on auditing and write audit logs to an Event Hub
- D. Apply sensitivity labels named Highly Confidential to the column.
- E. Turn on Azure Defender for SQL

F. Configure the auditing retention period to 365 days.

ANSWER: A D F

Explanation:

To meet the requirement, enable auditing and send audit logs to an Azure Storage account. Azure SQL Database Auditing records database events and query activity in an audit log, and a storage account is an appropriate destination for retaining those records. Configure the audit retention period as 365 days so that the auditing service retains the stored audit files for the required duration.

Apply sensitivity labels named Highly Confidential to the column. Sensitivity classification adds metadata to the column that identifies its data as sensitive. When Azure SQL auditing processes a query that returns classified data, the audit record can include *data_sensitivity_information*, enabling investigators to identify that sensitive column data was returned. This combination provides both the activity record and the sensitive-data context needed for monitoring access.

These settings work together: classification identifies the protected data, auditing captures access activity and sensitivity details, and the storage retention configuration preserves the audit trail for the required 365 days. For implementation details, see [Auditing for Azure SQL Database](#) and [Data discovery and classification for Azure SQL Database](#).

QUESTION NO: 42 - (DRAG DROP)

DRAG DROP

You are building an Azure virtual machine.

You allocate two 1-TiB, P30 premium storage disks to the virtual machine. Each disk provides 5,000 IOPS.

You plan to migrate an on-premises instance of Microsoft SQL Server to the virtual machine. The instance has a database that contains a 1.2-TiB data file. The database requires 10,000 IOPS.

You need to configure storage for the virtual machine to support the database.

Which three objects should you create in sequence? To answer, move the appropriate objects from the list of objects to the answer area and arrange them in the correct order.

Select and Place:

Actions

- a virtual disk that uses the stripe layout
- a virtual disk that uses the mirror layout
- a volume
- a virtual disk that uses the simple layout
- a storage pool

Answer Area

ANSWER:

Actions

a virtual disk that uses the stripe layout

a virtual disk that uses the mirror layout

a volume

a virtual disk that uses the simple layout

a storage pool

Answer Area

a storage pool

a virtual disk that uses the stripe layout

a volume



Explanation:

Create a storage pool first, using the two attached 1-TiB P30 premium SSD managed disks. Storage Spaces uses a pool as the collection of physical disks from which it provisions virtual disks. With both P30 disks in the pool, the available raw capacity is 2 TiB and the potential disk I/O capacity is the aggregate of the disks.

Next, create a virtual disk that uses the stripe layout. Striping distributes data across both physical disks, which allows the SQL Server data file to use the combined capacity and IOPS capability. Each P30 disk supplies 5,000 IOPS, so a striped virtual disk across the two disks can support the required 10,000 IOPS, assuming the selected VM SKU also provides sufficient uncached disk IOPS and throughput. The striped capacity is also sufficient for the 1.2-TiB database data file.

Finally, create a volume on the striped virtual disk. The volume is the usable, formatted file-system location that receives a drive letter or mount point and can be used for the SQL Server data file. This ordering follows the Storage Spaces provisioning model: physical disks are grouped into a pool, the pool is used to create the virtual disk, and the virtual disk is formatted and exposed through a volume. Microsoft documents this Storage Spaces workflow in [Storage Spaces overview](#) and describes Azure VM disk performance considerations in [Managed disks performance](#).

QUESTION NO: 43

You plan to build a structured streaming solution in Azure Databricks. The solution will count new events in five-minute intervals and report only events that arrive during the interval.

The output will be sent to a Delta Lake table.

Which output mode should you use?

- A. complete
- B. append
- C. update

ANSWER: B

Explanation:

The correct output mode is **append**. Append mode writes only rows that are new since the preceding streaming trigger, which aligns with the requirement to report only the results for newly completed five-minute event intervals rather than rewriting previously reported results. For a streaming aggregation such as a count grouped into time windows, the query should define an event-time watermark. The watermark establishes when a window can be treated as complete, allowing Structured Streaming to emit that window's final count once and append it to the Delta Lake target. This produces an incremental, durable table of interval-level counts suitable for downstream reporting and avoids repeatedly sending the entire aggregation state. Azure Databricks supports writing Structured Streaming output to Delta tables, and Delta provides transactional writes for each micro-batch. Microsoft's Structured Streaming documentation describes append mode as

emitting only new rows, while its Delta Lake streaming guidance covers using Delta tables as streaming sinks. See [Structured Streaming output modes](#) and [Azure Databricks Delta Lake streaming documentation](#).

QUESTION NO: 44

You have an Azure SQL Database logical server named Server1.

You need to configure transparent data encryption (TDE) for databases on Server1 by using a customer-managed key stored in Azure Key Vault.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Assign a managed identity to Server1.
- B. Grant the Server1 managed identity permission to use the key in Azure Key Vault.
- C. Create a shared access signature for the Key Vault.
- D. Create a database-scoped credential in every database.
- E. Enable Always Encrypted for all database columns.

ANSWER: A B

Explanation:

You must assign a managed identity to Server1 and grant that identity permissions to use the key in Azure Key Vault. The logical server uses its managed identity to access the customer-managed TDE protector. The identity requires permissions that allow it to retrieve and use the key, including key permissions such as `get`, `wrapKey`, and `unwrapKey`, depending on the Key Vault authorization model.

After access is configured, the key can be assigned as the server key protector. A shared access signature and a database-scoped credential do not provide Azure SQL Database with the authorization required to use a Key Vault key for TDE.

For more information, see [Transparent Data Encryption with customer-managed keys](#).

QUESTION NO: 45

You have 40 Azure SQL databases, each for a different customer. All the databases reside on the same Azure SQL Database server.

You need to ensure that each customer can only connect to and access their respective database.

Which two actions should you perform? Each correct answer presents part of the solution.

NOTE: Each correct selection is worth one point.

- A. Implement row-level security (RLS).
- B. Create users in each database.
- C. Configure the database firewall.
- D. Configure the server firewall.
- E. Create logins in the master database.
- F. Implement Always Encrypted.

ANSWER: B C

Explanation:

Creating users in each database establishes database-scoped identities and permissions for each customer. In Azure SQL Database, users can be created as contained database users, meaning that their authentication and authorization are managed within the individual database rather than granting broad access across the logical server. Granting each customer's user only the required permissions in that customer's database ensures that access is limited to the intended database.

Configuring the database firewall adds a network-level restriction that applies specifically to an individual database. Database-level IP firewall rules enable connections from approved customer IP address ranges to that database without opening equivalent access to every database hosted by the same logical server. Used with database-specific users and permissions, this provides both network access control and database authorization isolation for each customer.

Azure SQL Database documents the use of contained database users for database-level access management and describes database-level firewall rules as rules that apply to a specific database. See [Manage logins and users in Azure SQL Database](#) and [Configure Azure SQL Database firewall rules](#).