

DUMPSBOSS.

GIAC Cloud Forensics Responder (GCFR)

GIAC GCFR

Version Demo

Total Demo Questions: 8

Total Premium Questions: 88

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

QUESTION NO: 1

Which Azure binary large object type is usually used to store virtual hard drive files?

- A. Append
- B. Page
- C. Block

ANSWER: B

Explanation:

Page is the Azure Blob Storage type designed for virtual hard disk files. Page blobs are optimized for random read/write operations and organize data as fixed-size 512-byte pages. This makes them appropriate for the access patterns of virtual machine disks, where the platform may need to update or retrieve data at arbitrary offsets rather than only sequentially. Azure uses page blobs for VHD files in scenarios involving unmanaged disks; a page blob can be up to 8 TiB in size and supports efficient range-based updates. In Azure Storage forensic work, recognizing that an unmanaged VM disk is represented as a VHD in a page blob helps responders identify the relevant storage artifacts, acquire the disk correctly, and interpret the evidence source in its proper cloud context. Microsoft documents page blobs as the blob type used to store virtual hard disk files and notes their optimization for random read and write operations. See [Azure Page Blobs overview](#) and [Introduction to Azure Blob Storage](#).

QUESTION NO: 2

What is the recommended storage type when creating an initial snapshot of a VM in Azure for forensic analysis?

- A. Standard SSD
- B. Ultra Disk
- C. Premium SSD
- D. Standard HDD

ANSWER: D

Explanation:

Standard HDD is the recommended choice for the initial forensic snapshot because the immediate objective is to preserve a point-in-time copy of the VM disk as evidence, rather than to provide high-performance storage for an actively running workload. Standard HDD snapshots provide economical, durable storage that is appropriate for retaining the original acquired state while the investigation is scoped and evidence-handling decisions are made. Using a lower-cost snapshot tier at this stage helps an investigator preserve the source disk promptly without incurring unnecessary performance-tier charges during what may become a lengthy retention period. Azure snapshots are point-in-time copies of managed disks, and Azure supports Standard HDD as a snapshot storage SKU. The snapshot can later be used to create a separate managed disk for controlled examination, helping keep analysis activity isolated from the original preserved acquisition. Microsoft's disk documentation describes the available snapshot storage SKUs and their use with managed disks: [Azure managed disk snapshots](#). For the characteristics and intended use of Azure disk storage tiers, including Standard HDD, see [Azure managed disk types](#).

QUESTION NO: 3

Which performance feature of an Amazon EC2 instance is configured to add additional resources based on set trigger points?

- A. Burstable
- B. Optimized

- C. Managed
- D. Accelerated
- E. Amazon EC2 Auto Scaling

ANSWER: E

Explanation:

Amazon EC2 Auto Scaling is the AWS capability that automatically adjusts compute capacity in response to defined conditions. An Auto Scaling group maintains a collection of EC2 instances and can use scaling policies tied to Amazon CloudWatch metrics or scheduled criteria. For example, a target-tracking policy can launch additional instances when average CPU utilization exceeds the configured target, then remove capacity when demand falls. This allows an application to maintain availability and performance while aligning the number of running instances with workload demand.

The question's phrase "add additional resources based on set trigger points" describes dynamic scaling: capacity is increased after a metric alarm or policy condition is met. Auto Scaling can also define minimum, desired, and maximum instance counts, providing controlled and repeatable resource expansion. AWS documents that EC2 Auto Scaling monitors applications and automatically adjusts capacity to maintain steady, predictable performance at the lowest practical cost. See the [Amazon EC2 Auto Scaling User Guide](#) and AWS's [documentation on scaling based on demand](#).

QUESTION NO: 4

A Google Workspace responder needs to determine whether a suspicious message sent from an external domain was delivered to a particular user, rejected, or routed to Spam. Which Google Workspace evidence source is most appropriate?

- A. Google Workspace Admin audit log
- B. Google Drive audit log
- C. Email Log Search
- D. Google Cloud Audit Logs

ANSWER: C

Explanation:

Email Log Search is correct because it provides message-tracking information that can be used to determine delivery status and routing details for email involving Google Workspace users. The responder can search using attributes such as sender, recipient, subject, and time range to identify whether the message was delivered, rejected, quarantined, or otherwise processed. Admin audit logs record administrative activity, Drive audit logs record file-related events, and Google Cloud Audit Logs record activity involving Google Cloud resources rather than Gmail message delivery.

QUESTION NO: 5

The Azure PowerShell output below is an example of which of the following?

```

{
  "Name": "Contributor",
  "Id": "b24988vf-6180-42a0-ab55-20f7382dd24c",
  "IsCustom": false,
  "Description": "Manage everything except access to resources.",
  "Actions": [
    "*"
  ],
  "NotActions": [
    "Microsoft.Authorization/*/Delete",
    "Microsoft.Authorization/*/Write",
    "Microsoft.Authorization/elevateAccess/Action",
    "Microsoft.Blueprint/blueprintAssignments/write",
    "Microsoft.Blueprint/blueprintAssignments/delete"
  ],
  "DataActions": [],
  "NotDataActions": [],
  "AssignableScopes": [
    "/"
  ]
}

```

- A. Role assignment
- B. Managed identity
- C. Role definition
- D. Service principal

ANSWER: B

Explanation:

Managed identity is correct. In Microsoft Entra ID, an Azure managed identity is represented by a service-principal object that Azure creates and manages on behalf of an Azure resource. Azure PowerShell output for these identities commonly exposes directory-object properties such as an object ID, application/client ID, display name, and a service-principal type that identifies the object as a managed identity. This representation is expected because managed identities authenticate to Microsoft Entra ID using a service principal, but their credential lifecycle is handled by Azure rather than by an administrator or application developer.

Managed identities enable a workload running in Azure to request tokens for supported services without storing passwords, client secrets, or certificates in application code or configuration. This makes their presence especially relevant to cloud incident response: the displayed identity can be correlated with Azure resource configuration, Microsoft Entra sign-in and audit activity, RBAC permissions, and resource access logs to establish what workload identity was available and what actions it could perform. Microsoft documents that managed identities provide an automatically managed identity in Microsoft Entra ID and that the platform manages its credentials. See [Managed identities for Azure resources overview](#) and [Get-AzADServicePrincipal documentation](#).

QUESTION NO: 6

An engineer has set up log forwarding for a new data source and wants to use that data to run reports and create dashboards in Kibana. What needs to be created in order to properly handle these logs?

- A. Row
- B. Parser
- C. ingest script
- D. Beat

ANSWER: B

Explanation:

Parser is correct because forwarded logs must be parsed into consistently named, typed, and searchable fields before they can support useful Kibana reporting and visualizations. A parser extracts meaningful elements from each event—such as timestamps, source and destination addresses, usernames, event actions, status values, and messages—and maps them into a structured schema. This turns raw log text into fields that Elasticsearch can index and Kibana can aggregate, filter, and display. Accurate parsing also ensures timestamps are recognized correctly, which is essential for time-based dashboard views and report queries. In Elastic environments, this parsing function can be implemented through mechanisms such as ingest-pipeline processors, including Grok or Dissect, but the required logical artifact is the parser that defines how the new source's events are interpreted. Once parsing produces usable fields, the data can be selected in a Kibana data view and used in Discover, Lens, dashboards, and reports. Elastic documents how ingest pipelines transform documents before indexing and how processors extract structured data from log content: [Elasticsearch ingest pipelines](#) and [Grok processor reference](#).

QUESTION NO: 7

A company using PaaS to host and develop their software application is experiencing a DOS attack. What challenge will a DFIR analyst experience when investigating this attack?

- A. Restricted access to their application logs
- B. Resource scaling will affect access to logs
- C. Network logs are unavailable for review
- D. Network monitoring disabled by the company

ANSWER: C

Explanation:

Network logs are unavailable for review is correct because a platform-as-a-service environment places the underlying network infrastructure under the cloud provider's control. During a denial-of-service investigation, the most useful evidence may include packet-level data, network-flow records, firewall events, load-balancer telemetry, routing information, and mitigation-service events. In a PaaS model, the customer and its DFIR team commonly do not have direct access to that provider-managed infrastructure or to its native network telemetry. The provider may retain relevant evidence, but it is generally obtained only through provider support, a formal incident process, or legal disclosure procedures, and its availability and retention are outside the customer's direct control.

This visibility gap can limit the analyst's ability to establish the traffic source, volume, protocol characteristics, attack path, and precise timeline at the network layer. The investigation must therefore rely on evidence exposed by the PaaS service, such as application telemetry and service-level diagnostic logs, while requesting any provider-held network evidence promptly. Microsoft describes that customers use platform diagnostic settings to route available resource logs and metrics, while the service provider operates the underlying cloud infrastructure. See [Azure Monitor diagnostic settings](#) and [AWS Shared Responsibility Model](#).

QUESTION NO: 8

Which is a limitation when adding GPUs to Google cloud VMs?

- A. They can only be added at VM creation
- B. Preemptible VMs do not support GPU addition
- C. Google limits the GPUs assigned to a single VM
- D. They are only available in specific zones

ANSWER: D

Explanation:

They are only available in specific zones is correct because GPU capacity and accelerator availability in Compute Engine are location-dependent. Google does not make every GPU model available in every zone, so an investigator or cloud administrator planning a GPU-enabled VM must first select a region and zone that offers the required accelerator type. Availability can also vary over time because GPU resources are finite within a zone; a supported zone might temporarily be unable to satisfy a request if capacity is exhausted. This makes zone selection an operational constraint when provisioning GPU-backed virtual machines, particularly where an existing workload, evidence-processing environment, or other dependent resources are tied to a particular location. Google's accelerator documentation directs users to identify the regions and zones in which each GPU model is available before creating the VM. In practice, the Google Cloud console, gcloud CLI, or Compute Engine API can be used to inspect accelerator types for a chosen zone and then request an appropriate machine configuration. See the [Compute Engine GPU regions and zones documentation](#) and Google's [GPU platform documentation](#) for supported accelerator locations and provisioning guidance.