

IBM Cloud Pak for Data v4.7 Architect

IBM C1000-173

Version Demo

Total Demo Questions: 7

Total Premium Questions: 78

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Planning	5
Topic 2, Architecture	8
Topic 3, Installation	11
Topic 4, Security	19
Topic 5, Data Fabric	28
Topic 6, Data Science	7
Total	78

QUESTION NO: 1

If a Cloud Pak for Data cluster is air-gapped, it is unable to reach the internet, and a private registry must be configured for installations and upgrades. Beyond container images, what other components need to be addressed?

- A. Utilities, like pip, still require a connection to the internet by default.
- B. License validation requires periodic updates via internet proxy.
- C. Platform monitoring messages are sent via an internet messaging hub.
- D. Data movement between nodes must be blocked via firewall rules.

ANSWER: A

Explanation:

Utilities, like pip, still require a connection to the internet by default. An air-gapped Cloud Pak for Data deployment therefore needs planning for external software dependencies in addition to mirroring the container images that the platform services use. Package-management tooling commonly attempts to contact its configured public package index when it resolves or installs dependencies. This is especially relevant when Python packages are installed for workloads, custom environments, notebooks, or solution-specific extensions. To keep these operations functional without outbound network access, an organization should make approved package artifacts available from an internal package repository or mirror and configure the relevant tooling to use that internal source. Pip supports offline installation patterns through local package directories and index controls, including use of a local wheelhouse with `--no-index` and `--find-links`. Cloud Pak for Data air-gapped installation guidance likewise treats disconnected installation as a process that requires locally available installation artifacts and registries rather than direct access to public endpoints. This ensures that deployment and later operational activities remain repeatable while preserving the network isolation requirement. See the [IBM Cloud Pak for Data installation documentation](#) and the [pip install documentation](#).

QUESTION NO: 2

How many service instances can be provisioned for Watson Discovery at one time?

- A. 15
- B. 20
- C. 5
- D. 10

ANSWER: D

Explanation:

10 is correct because IBM Cloud Pak for Data limits Watson Discovery provisioning to a maximum of 10 service instances per deployment. This is a deployment-level limit, meaning that the count applies to all Watson Discovery instances created within the same Cloud Pak for Data deployment rather than to a particular project, user, or namespace. Administrators can provision instances up to this limit when separate Discovery configurations or workloads require isolation. Once the tenth instance has been created, Cloud Pak for Data prevents further provisioning through the user interface: the option to create a new instance is no longer displayed. This behavior helps maintain supported service capacity and makes the maximum enforceable without requiring users to calculate availability manually. The limit is specific to the Watson Discovery service in the documented Cloud Pak for Data environment and should be considered during architecture planning, especially where multiple teams expect dedicated Discovery instances. IBM documents service provisioning and service-instance management in its Cloud Pak for Data documentation: [Watson Discovery service documentation](#) and [Provisioning service instances](#).

QUESTION NO: 3

Which two features are valid only when deploying Cloud Pak for Data on-premises?

- A. The number of OpenShift nodes is configured by the administrator.
- B. Compute resources are automatically scaled.
- C. Services are automatically updated.
- D. Persistent storage is always used.
- E. Network security is managed by IBM.

ANSWER: A D

Explanation:

“The number of OpenShift nodes is configured by the administrator.” is correct because an on-premises Cloud Pak for Data installation runs on an OpenShift cluster owned and operated by the customer. Before installation, the organization plans the cluster capacity and provides the required worker nodes, including sufficient CPU, memory, and infrastructure capacity for the Cloud Pak for Data services it intends to run. This direct responsibility for defining and maintaining the OpenShift node count is characteristic of a customer-managed, on-premises deployment.

“Persistent storage is always used.” is also correct because Cloud Pak for Data on-premises workloads require persistent volumes supplied by a supported storage solution in the customer’s OpenShift environment. Persistent storage supports retained application data, service metadata, and other stateful workload requirements across pod restarts or rescheduling. The storage classes, persistent-volume provisioning, capacity, performance characteristics, and availability are planned and administered as part of the on-premises platform design. IBM’s installation planning documentation emphasizes that storage must be prepared and available before deploying the platform and its services. See the [IBM Cloud Pak for Data 4.7 documentation](#) and IBM’s [storage planning considerations](#).

QUESTION NO: 4

Which two Cloud Pak for Data services implement data masking to support secure data sharing?

- A. Db2 Data Gate
- B. IBM Knowledge Catalog
- C. DataStage
- D. Data Privacy
- E. SPSS

ANSWER: B D

Explanation:

IBM Knowledge Catalog and Data Privacy provide complementary masking capabilities for secure data sharing in Cloud Pak for Data. IBM Knowledge Catalog supports governed access through data protection rules. These rules use governed metadata, including classifications for sensitive information, to define who can access data and what protection is applied. When enforced with the platform’s governed data-access capabilities, protection rules can mask sensitive values for users who are permitted to use a data asset but should not see its confidential contents. This lets organizations share useful data while consistently applying centrally managed privacy policies.

Data Privacy is designed to de-identify data through privacy and masking flows. It provides techniques for transforming sensitive values while retaining data utility for development, testing, analytics, and collaboration. Masking can be configured to meet the protection needs of particular data elements, allowing teams to prepare a safer copy of data for an approved downstream purpose. Together, IBM Knowledge Catalog supplies policy-based governance and protection enforcement, while Data Privacy supplies transformation-based masking for data copies and workflows. IBM documents data protection rules in its [Cloud Pak for Data governance documentation](#) and masking capabilities in the [Data Privacy documentation](#).

QUESTION NO: 5

An organization imports a Knowledge Accelerator and wants to add organization-specific business terms and policies. The governance team must be able to distinguish its additions from the IBM-provided accelerator content during ongoing maintenance.

How should the organization customize the Knowledge Accelerator?

- A. Customize the content in a separate namespace.
- B. Export the accelerator content to a Git repository.
- C. Create a separate Cloud Pak for Data project for the content.
- D. Install a separate IBM Knowledge Catalog service instance.

ANSWER: A

Explanation:

Customize the content in a separate namespace. A separate namespace isolates organization-specific governance artifacts from the base Knowledge Accelerator content while allowing the organization to relate its local terms, policies, and categories to the supplied vocabulary. This approach supports clearer ownership and maintenance of custom content.

Inline changes are appropriate when modifications are intended to become part of the installed accelerator content, but they do not provide the same separation between IBM-provided and organization-specific artifacts.

QUESTION NO: 6

A retailer has customer records from e-commerce, loyalty, and in-store systems. The records use different identifiers and formats, but the business needs a trusted mastered view of each customer and the relationships among the source records.

Which Cloud Pak for Data service is most appropriate?

- A. IBM Data Virtualization
- B. DataStage
- C. IBM Match 360
- D. IBM Knowledge Catalog

ANSWER: C

Explanation:

IBM Match 360 is designed for master data management and entity resolution. It can ingest records from disparate sources, identify records that represent the same real-world entity, and create mastered entities that users can explore and analyze.

IBM Data Virtualization provides logical access to distributed data but does not itself resolve entities into mastered records. DataStage transforms and integrates data, while IBM Knowledge Catalog governs and catalogs data assets. These capabilities can support a master-data solution, but Match 360 addresses the entity-resolution requirement directly.

QUESTION NO: 7

An architect is selecting persistent storage for a production Cloud Pak for Data deployment. The planned services have different persistent-volume requirements, and several workloads will process high volumes of data.

Which two considerations should guide the storage selection? Select two.

- A. The storage supports the services that will be installed.

- B. The storage provides sufficient I/O performance.
- C. The storage network supports a minimum transmission speed of 4 Gbps.
- D. The storage provides a fixed minimum throughput of 16 MB/s per disk.
- E. NFS storage supports all Cloud Pak for Data services.

ANSWER: A B

Explanation:

The selected storage must **support the services that will be installed**, including their required access modes and persistent-volume characteristics. It must also **provide sufficient I/O performance** for the expected workload, including the required latency, throughput, and IOPS.

Capacity planning is important, but a fixed network-speed or per-disk throughput value is not a general Cloud Pak for Data storage requirement. NFS is not automatically suitable for every service.