

DUMPSBOSS.

Nokia Bell Labs Distributed Cloud Networks Exam

Nokia BL0-220

Version Demo

Total Demo Questions: 8

Total Premium Questions: 83

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

QUESTION NO: 1

A cloud-native service must recover automatically when a workload fails at an edge site. Which two capabilities are required to implement an effective auto-healing workflow? (Select 2)

- A. Fault detection
- B. Automated corrective action
- C. Manual remediation for every event
- D. A fixed capacity plan

ANSWER: A B

Explanation:

Fault detection is required to identify that the workload is no longer operating as intended. Automated corrective action is then required to restart, replace, or otherwise recover the failed workload according to the desired state.

A warning message alone does not ensure recovery, and a fixed capacity plan does not detect or remediate a fault. Manual remediation can restore a service, but it does not provide auto-healing.

QUESTION NO: 2

Which of the following are properties related to a public cloud? (Select 2)

- A. For users, there is no cloud infrastructure operation
- B. An In-house solution
- C. Easy scaling
- D. O&A requirement
- E. Very complex management

ANSWER: A C

Explanation:

For users, there is no cloud infrastructure operation is a defining practical characteristic of the public-cloud consumption model. A public-cloud provider owns and operates the underlying data-center facilities, servers, storage, networking, and cloud platform. The customer consumes services through a provider interface and remains responsible for its workloads and configurations, but does not perform day-to-day operation of the provider's physical cloud infrastructure. This provider-operated model lets organizations obtain compute, storage, and platform capabilities without building and running the underlying environment themselves.

Easy scaling is also a core public-cloud property. Public-cloud services expose a pool of shared, on-demand resources that can be provisioned and released rapidly as workload demand changes. This elasticity allows users to increase or reduce capacity without purchasing, installing, and operating new hardware for every change in demand. NIST identifies rapid elasticity and on-demand self-service as essential cloud-computing characteristics, while the Cloud Security Alliance describes public cloud as infrastructure made available for open use and operated by a cloud provider. These characteristics support the operational convenience and scalable consumption expected from public-cloud services. See the [NIST cloud-computing definition](#) and the [Cloud Security Alliance Security Guidance](#).

QUESTION NO: 3

What event will trigger the auto-healing operation?

- A. Threshold

- B. Warning message
- C. New status
- D. Fault detection

ANSWER: D

Explanation:

Fault detection is correct because auto-healing is initiated when cloud-management monitoring identifies an actual fault affecting a resource, service, or workload. Detection provides the operational signal that a failure condition exists and that recovery action is required. The auto-healing workflow can then apply the appropriate remediation for the detected condition, such as restarting an affected workload, restoring service through an available replacement, or otherwise returning the resource to its intended operational state. This closed-loop sequence—monitoring, fault detection, and automated corrective action—is a fundamental operational capability for distributed cloud environments. The trigger is therefore the detection of the fault itself, rather than a general informational state change. Nokia's cloud and network automation approach emphasizes assurance-driven operations in which telemetry and fault information support automated lifecycle and service-recovery processes. See Nokia's [cloud-native network automation overview](#) and the [Nokia Cloud Platform overview](#) for Nokia information on cloud-native automation and operations.

QUESTION NO: 4

In public networks, what are the parts of a distributed cloud used for the 5G Core control plane. (Select 2)

- A. Metro Edge Cloud
- B. Core/Central Cloud
- C. Far Edge Cloud
- D. On-premise Edge Cloud

ANSWER: A B

Explanation:

The correct distributed-cloud locations for a public-network 5G Core control plane are **Metro Edge Cloud** and **Core/Central Cloud**. The control plane contains functions that establish, authorize, control, and manage service sessions and mobility. Its cloud placement must therefore balance proximity to the radio-access network with resilient, centralized operation. A Metro Edge Cloud provides a regional location close to access and aggregation networks. This placement supports low-latency control interactions and regional scaling where proximity is beneficial for subscribers and services. The Core/Central Cloud provides the highly available, secure, and scalable cloud environment used for centralized control-plane capabilities and network-wide coordination. Together, these locations form a practical distributed deployment model: regional metro sites extend cloud resources toward the network edge, while central sites provide robust shared infrastructure for core functions. Nokia describes its cloud-native 5G Core as supporting flexible deployment across distributed cloud infrastructure, and its edge-cloud portfolio describes metro-edge infrastructure as extending cloud capabilities toward the access network. See [Nokia 5G Core](#) and [Nokia Edge Cloud](#).

QUESTION NO: 5

What is the approximate maximum distance of an Edge Cloud from the end-user?

- A. 10000KM
- B. 300KM
- C. 5000KM
- D. 10Q0KM

ANSWER: B

Explanation:

300KM is the approximate maximum distance associated with an Edge Cloud in the distributed-cloud model used in this context. Edge Cloud places compute, storage, and networking capabilities substantially nearer to users and data sources than centralized cloud regions. This geographic proximity is essential for applications whose experience or operation depends on low and predictable latency, such as interactive video, industrial automation, connected transport, and real-time analytics.

Rather than identifying a single fixed site type, “edge” describes a placement on the cloud continuum that is close enough to the access network and end user to avoid the delay introduced by carrying traffic to a distant centralized data center. The approximate 300-kilometre boundary represents the outer end of that edge placement range in the course’s model; deployments may be much closer depending on service requirements and available network locations. Nokia describes edge cloud as bringing cloud capabilities closer to where data is generated and consumed, while ETSI’s MEC work similarly focuses on hosting applications close to users and network access points. See [Nokia Edge Cloud](#) and [ETSI Multi-access Edge Computing](#).

QUESTION NO: 6

Select the best option below to complete the following sentence.

Your apps need cloud resources which are _____ and then ____

- A. commissioned, instantiated
- B. allocated, provisioned
- C. started, commissioned
- D. distributed, terminated

ANSWER: B

Explanation:

The correct completion is “**allocated, provisioned.**” Cloud-resource planning first allocates resources: the platform identifies and reserves the required compute, storage, networking, and related capacity for the application according to its requested policy, placement, and available capacity. Allocation therefore establishes that the resources have been assigned to the workload and are available for its use.

Provisioning follows allocation. At this stage, the assigned capacity is configured into usable cloud resources—for example, virtual machines or containers are created, storage and network connectivity are configured, required images and software settings are applied, and access is made available to the application. The application can then consume those resources as an operational environment. This ordered distinction is important in distributed cloud networks, where resource availability and placement must be determined before the infrastructure can be configured for a particular workload.

This usage aligns with the resource-planning lifecycle taught in the Nokia Bell Labs Distributed Cloud Networks material and with cloud-management practice, in which provisioning prepares assigned infrastructure for consumption. Nokia’s cloud networking portfolio provides context for the distributed-cloud environment in which these lifecycle functions operate: [Nokia Cloud and Network Services](#). A general technical description of provisioning operations is also available in the [OpenStack documentation](#).

QUESTION NO: 7

Hyperscale computing relies on scalable server architecture.

- A. True
- B. False

ANSWER: A

Explanation:

“True” is correct. Hyperscale computing is designed to provide compute, storage, and networking capacity at extremely large scale while maintaining efficient operations and the ability to expand incrementally. It therefore depends on a scalable server architecture: standardized server resources can be deployed in large numbers, interconnected through high-capacity networks, and managed as pooled infrastructure. This architecture allows operators to add capacity as demand increases rather than redesigning the entire environment, supporting the elastic growth expected of cloud-scale services. Hyperscale data centers also use extensive automation and software-defined management to provision, monitor, repair, and optimize large fleets of servers consistently. The resulting combination of scale-out hardware, resource pooling, and automation supports high workload volumes, resilient service delivery, and more efficient utilization of infrastructure. These characteristics align with Nokia’s distributed-cloud focus, in which cloud infrastructure must scale while supporting workloads across centralized and distributed locations. For additional background, see the [Nokia cloud solutions overview](#) and the National Institute of Standards and Technology’s cloud-computing definition, including rapid elasticity and resource pooling, at [NIST SP 800-145](#).

QUESTION NO: 8

Which of the following are characteristics of Cloud Native services. (Select 2)

- A. Low Scalability
- B. Very light weight application
- C. Fixed capacity
- D. Very fast deployment

ANSWER: B D

Explanation:

Cloud-native services are designed as lightweight applications, commonly decomposed into small, independently deployable components such as microservices. Packaging these components in containers provides a consistent runtime environment and reduces the overhead associated with deploying a complete, monolithic application stack. This lightweight approach supports efficient use of cloud infrastructure and enables services to be operated flexibly across cloud environments.

Very fast deployment is also a core cloud-native characteristic. Cloud-native architectures use automation throughout the delivery and operations lifecycle, including CI/CD pipelines, declarative configuration, container orchestration, and automated rollout or rollback processes. These capabilities allow teams to release changes rapidly and reliably, while scaling service instances according to demand. The Cloud Native Computing Foundation describes cloud-native technologies as enabling organizations to build and run scalable applications in dynamic environments, using containers, service meshes, microservices, immutable infrastructure, and declarative APIs. Nokia similarly presents cloud-native network software as containerized and automated for faster innovation and deployment. See the [CNCF Cloud Native Definition](#) and Nokia’s [Cloud-native network software](#) overview.