

Using HPE AI and Machine Learning

HP HPE2-N69

Version Demo

Total Demo Questions: 6

Total Premium Questions: 68

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

QUESTION NO: 1

What role do HPE ProLiant DL325 servers play in HPE Machine Learning Development System?

- A. They run validation and checkpoint workloads.
- B. They run training workloads that do not require GPUs.
- C. They host management software such as the conductor and HPCM.
- D. They run non-distributed training workloads.

ANSWER: C

Explanation:

They host management software such as the conductor and HPCM. In an HPE Machine Learning Development System deployment, the HPE ProLiant DL325 platform serves as the management-server foundation for the environment. This management role provides the CPU, memory, storage, and networking resources needed to operate the cluster-management and orchestration components, including the conductor and HPCM. Centralizing these services on a dedicated, general-purpose ProLiant server enables the GPU-equipped compute resources to remain focused on accelerated AI and machine-learning work while the management plane handles provisioning, monitoring, scheduling, configuration, and operational control of the system. The DL325 is well suited to this purpose because its single-socket AMD EPYC architecture delivers substantial core density and I/O capability in a compact rack form factor, making it an efficient platform for infrastructure services. HPE positions ProLiant DL325 servers as versatile compute platforms for demanding enterprise workloads, while HPE's AI portfolio describes the need for an integrated infrastructure and management layer to support the AI lifecycle. See the [HPE ProLiant DL325 server overview](#) and [HPE Artificial Intelligence solutions](#).

QUESTION NO: 2

Refer to the exhibit.



You are demonstrating HPE Machine Learning Development Environment, and you show details about an experiment, as shown in the exhibits. The customer asks about what "validation loss" means. What should you respond?

- A. Validation refers to testing how well the current model performs on new data; the lower the loss, the better the performance.
- B. Validation refers to an assessment of how efficient the model code is; the lower the loss the lower the demand on GPU memory resources.
- C. Validation loss refers to the loss detected during the backward pass of training, while training loss refers to loss during the forward pass.
- D. Validation loss is metadata that indicates how many updates were lost between the conductor and agents.

ANSWER: A

Explanation:

Validation refers to evaluating the current model with a validation dataset: data held out from the model's parameter-learning process. Validation loss is the value of the selected loss function when the model produces predictions for that held-out data and those predictions are compared with the known target values. It is therefore a practical indicator of how well the model is likely to generalize to data it has not learned from directly. For a given experiment using the same validation set and loss definition, a lower validation-loss value generally indicates predictions that are closer to the expected results and thus better performance on new data. Monitoring this metric over training is especially useful because it helps teams assess generalization while they compare trials, checkpoints, hyperparameters, and model architectures. HPE Machine Learning Development Environment tracks metrics reported by training workloads, including validation metrics, in experiment and trial views; the specific numerical meaning remains defined by the training code's loss function. See the [Determined Keras training API documentation](#) for validation-data metric reporting and the [HPE AI solutions overview](#) for HPE's AI platform context.

QUESTION NO: 3

You are meeting with a customer how has several DL models deployed. Out wants to expand the projects.

The ML/DL team is growing from 5 members to 7 members. To support the growing team, the customer has assigned 2 dedicated IT staff. The customer is trying to put together an on-prem GPU cluster with at least 14 CPUs.

What should you determine about this customer?

- A. The customer is not ready for an HPE Machine Learning Development solution, but you could recommend open-source Determined AI.
- B. The customer is not ready for an HPE Machine Learning Development solution. Out you could recommend an educational HPE Pointnext ASPS workshop.
- C. The customer is a key target for HPE Machine Learning Development Environment, but not HPE Machine Learning Development System.
- D. The customer is a key target for an HPE Machine Learning Development solution, and you should continue the discussion.

ANSWER: D

Explanation:

The customer is a key target for an HPE Machine Learning Development solution, and you should continue the discussion. The customer has already deployed several deep-learning models, which shows that AI/ML work is moving beyond an initial proof of concept. Their plan to expand projects, increase the ML/DL team from five to seven people, and assign dedicated IT personnel indicates a need for a more structured, scalable on-premises machine-learning platform. A GPU cluster with at least 14 CPUs also demonstrates that the customer is planning the compute infrastructure needed for shared, accelerated model development and training.

HPE Machine Learning Development solutions are intended to help teams operationalize and scale AI initiatives by providing an environment for managing users, workloads, experiments, and GPU-based training resources. The presence of both a growing data-science team and dedicated IT support is especially important because successful on-premises ML platforms require collaboration between ML practitioners and infrastructure administrators. These customer characteristics justify continuing discovery around workload requirements, GPU configuration, storage, networking, user access, and operational processes. HPE describes its ML development offerings and their role in accelerating and scaling AI initiatives at [HPE Machine Learning Development System](#) and [HPE Artificial Intelligence solutions](#).

QUESTION NO: 4

An ML engineer trains a neural network and observes that the training loss fluctuates sharply and does not converge. The engineer is using gradient descent.

What change should the engineer consider first?

- A. Reduce the learning rate.
- B. Increase the number of output classes.
- C. Remove the validation dataset.
- D. Increase the prediction threshold.

ANSWER: A

Explanation:

The engineer should **reduce the learning rate**. The learning rate controls the size of the parameter updates made by the optimizer. If it is too large, updates can repeatedly overshoot a minimum of the loss function, causing unstable loss values or preventing convergence.

A smaller learning rate makes updates more gradual and can improve training stability, although it may increase training time. Other factors can affect convergence, but an excessively large learning rate is a common cause of rapidly oscillating or divergent loss. See Google's [hyperparameters guide](#) for more information about learning-rate selection.

QUESTION NO: 5

A trial is running on a GPU slot within a resource pool on HPE Machine Learning Development Environment. That GPU fails. What happens next?

- A. The trial fails, and the ML engineer must restart it manually by re-running the experiment.
- B. The conductor reschedules the trial on another available GPU in the pool, and the trial restarts from the state of the latest training workload.
- C. The conductor reschedules the trial on another available GPU in the pool, and the trial restarts from the latest checkpoint.
- D. The trial fails, and the ML engineer must manually restart it from the latest checkpoint using the WebUI.

ANSWER: C

Explanation:

The conductor reschedules the trial on another available GPU in the pool, and the trial restarts from the latest checkpoint. HPE Machine Learning Development Environment, which is powered by Determined AI technology, is designed to provide fault tolerance for distributed training workloads. The platform's conductor monitors trial execution and the availability of resources in the configured resource pool. When a GPU or its hosting resource becomes unavailable, the conductor can requeue the affected trial and schedule it onto suitable remaining capacity in that pool.

Training progress is preserved through checkpoints written by the trial. After the trial is rescheduled, the training process restores the most recent checkpoint, allowing it to continue from the latest saved model and training state rather than beginning again from scratch. This behavior makes resource-pool execution resilient to transient hardware failures and avoids requiring an ML engineer to intervene merely to restart a recoverable workload. Checkpointing is therefore an important part of a fault-tolerant experiment configuration. See the Determined documentation on [fault tolerance](#) and [checkpoints](#).

QUESTION NO: 6

An ML engineer wants to evaluate randomly selected combinations of learning rate, batch size, and dropout rate from defined hyperparameter ranges in HPE Machine Learning Development Environment.

Which searcher should the engineer use?

- A. random
- B. single
- C. checkpoint
- D. distributed

ANSWER: A

Explanation:

The engineer should use the **random** searcher. Random search samples hyperparameter configurations from the specified search spaces, creating trials that evaluate different combinations of the candidate values. It is useful when the engineer wants broader exploration without evaluating every possible combination as a grid search would.

The experiment configuration defines both the hyperparameter ranges and the searcher behavior. A `single` searcher runs one fixed configuration, while ASHA-based searchers also use intermediate results to allocate training resources. HPE MLDE documents supported searchers in the [experiment configuration reference](#).