

Testpassport**Q&A**



H i g h e r Q u a l i t y

B e t t e r S e r v i c e !

We offer free update service for one year
[Http://www.testpassport.com](http://www.testpassport.com)

Exam : NCP-ADS

**Title : NVIDIA-Certified-
Professional Accelerated
Data Science**

Version : DEMO

1.A data scientist is training a deep learning model on an NVIDIA GPU but is encountering out-of-memory (OOM) errors.

To optimize GPU memory usage while maintaining efficient training performance, which of the following strategies should they prioritize?

- A. Increasing batch size without adjusting the optimizer settings
- B. Using mixed precision training with automatic loss scaling
- C. Storing all training data in GPU memory at once
- D. Using single-precision (FP32) calculations for better accuracy

Answer: B

2.You are training a large-scale random forest model on a dataset with millions of rows and hundreds of features. The training time is significantly high when using traditional CPU-based machine learning frameworks.

Which NVIDIA technology should you use to accelerate training while maintaining compatibility with common ML frameworks like scikit-learn?

- A. NVIDIA DeepStream to preprocess tabular data and optimize random forest model execution.
- B. NVIDIA RAPIDS cuML to accelerate random forest training using GPU-optimized implementations.
- C. NVIDIA Triton Inference Server to distribute random forest model training across multiple GPUs.
- D. NVIDIA TensorRT to accelerate random forest model training by optimizing tree-based algorithms.

Answer: B

3.A data scientist is working on training a deep learning model in a cloud-based environment. The dataset is large, and model convergence is taking too long on a standard CPU instance.

To optimize performance through GPU acceleration, which of the following strategies should the data scientist implement?

- A. Use a cloud instance with multiple GPUs and enable mixed-precision training.
- B. Store all training data in RAM and load it directly to the CPU for processing.
- C. Disable CUDA and use only OpenMP to parallelize computations across CPU cores.
- D. Increase the number of CPU cores and distribute training across multiple CPU threads.

Answer: A

4.You are training a machine learning model using RAPIDS cuML and need to ensure that all numeric features are standardized for better model performance.

Which of the following is the best approach for scaling data using RAPIDS?

- A. `df_scaled = df / df.max()`
- B. `df_scaled = (df - df.min()) / (df.max() - df.min())`
- C. `scaler = cuml.preprocessing.StandardScaler()`
- D. `df_scaled = scaler.fit_transform(df)`
- E. `df_scaled = df.apply(lambda x: x / np.linalg.norm(x))`

Answer: C

5.When deciding whether to use GPU acceleration or a traditional CPU approach for a machine learning task, which of the following factors should be considered to determine if the data qualifies as "big data" and whether GPU acceleration is beneficial? (Select two)

- A. CPU-based machine learning methods are always more effective for small datasets, regardless of the algorithm used.
- B. The size of the dataset in terms of rows and columns is irrelevant when determining if it qualifies as big data.
- C. GPU acceleration is beneficial when the dataset can be divided into independent chunks that can be processed in parallel.
- D. The complexity of the algorithm being used plays a crucial role in deciding whether to use GPU acceleration, with more complex algorithms benefiting from parallel computation.
- E. The dataset must be over 100GB in size to qualify as big data and warrant GPU acceleration.

Answer: C, D