

NCP-AII Fragen&Antworten & NCP-AII Prüfungsvorbereitung

Question: 3

You notice that one of the fans in your GPU server is running at a significantly higher RPM than the others, even under minimal load. 'ipmitool sensor' output shows a normal temperature for that GPU. What could be the potential causes?

- A. The fan's PWM control signal is malfunctioning, causing it to run at full speed.
- B. The fan bearing is wearing out, causing increased friction and requiring higher RPM to maintain airflow.
- C. The fan is attempting to compensate for restricted airflow due to dust buildup.
- D. The server's BMC (Baseboard Management Controller) has a faulty temperature sensor reading, causing it to overcompensate.
- E. A network connectivity issue is causing higher CPU utilization, leading to increased system-wide heat.

Answer: A,B,C

Explanation:

A malfunctioning PWM control signal, worn fan bearings, or restricted airflow can all cause a fan to run at higher RPMs. While a faulty BMC sensor could be a cause, the question states that 'ipmitool sensor' shows a normal temperature. Network connectivity issues are less likely to cause an isolated fan to run high, if the GPU temperature is normal.

Question: 4

After upgrading the network card drivers on your AI inference server, you experience intermittent network connectivity issues, including packet loss and high latency. You've verified that the physical connections are secure. Which of the following steps would be most effective in troubleshooting this issue?

- A. Roll back the network card drivers to the previous version.
- B. Check the system logs for error messages related to the network card or driver.
- C. Run network diagnostic tools like 'ping', 'traceroute', and 'iperf3' to assess the network performance.
- D. Reinstall the operating system.
- E. Update the server's BIOS.

Answer: A,B,C

Explanation:

Rolling back drivers is a quick way to revert to a known working state. Checking system logs will provide valuable information about driver errors or network issues. Network diagnostic tools will quantify the network performance and help isolate the problem. Reinstalling the OS is drastic and should be a last resort. Updating the BIOS is unlikely to resolve driver-related network issues unless specifically recommended for the network card.

Visit us at: <https://p2pexam.com/ncp-aai>

2026 Die neuesten DeutschPrüfung NCP-AII PDF- Versionen Prüfungsfragen und NCP-AII Fragen und Antworten sind kostenlos verfügbar: https://drive.google.com/open?id=1FU5RLNdCO_Q22whH-w31ubgXJymB_Kys

Die NVIDIA NCP-AII Zertifizierungsprüfung wird jetzt immer populärer. Es gibt viele verschiedene IT-Zertifizierungsprüfungen. Welche Prüfung haben Sie abgelegt? Lassen Wir hier NVIDIA NCP-AII Zertifizierungsprüfung als Beispiel erklären. Wenn Sie an der NCP-AII Prüfung teilnehmen, NVIDIA NCP-AII Dumps von DeutschPrüfung Ihnen helfen, sehr leicht die Prüfung zu bestehen.

Die Zertifikat der NVIDIA NCP-AII ist international anerkannt. Sie zu erwerben bedeutet, dass Sie den Schlüssel zur höheren Stelle besitzen. Die NVIDIA NCP-AII Prüfungsunterlagen von DeutschPrüfung werden von erfahrenen IT-Profis herstellt und immer wieder aktualisiert. Jetzt können Sie mit günstigem Preis die verlässliche NVIDIA NCP-AII Prüfungsunterlagen genießen. Nachdem Sie die Zertifizierung erworbt haben, können Sie leicht eine höhere Arbeitsposition oder Gehalten bekommen.

>> NCP-AII Fragen&Antworten <<

NCP-AII Studienmaterialien: NVIDIA AI Infrastructure - NCP-AII Torrent Prüfung & NCP-AII wirkliche Prüfung

Die Zertifikat der NVIDIA NCP-AII ist international anerkannt. Sie zu erwerben bedeutet, dass Sie den Schlüssel zur höheren Stelle besitzen. Die NVIDIA NCP-AII Prüfungsunterlagen von DeutschPrüfung werden von erfahrenen IT-Profis herstellt und immer wieder aktualisiert. Jetzt können Sie mit günstigem Preis die verlässliche NVIDIA NCP-AII Prüfungsunterlagen genießen. Nachdem Sie die Zertifizierung erworbt haben, können Sie leicht eine höhere Arbeitsposition oder Gehalten bekommen.

NVIDIA AI Infrastructure NCP-AII Prüfungsfragen mit Lösungen (Q122-Q127):

122. Frage

You are building a cloud-native application that uses both CPU and GPU resources. You want to optimize resource utilization and cost by scheduling CPU-intensive tasks on nodes without GPUs and GPU-intensive tasks on nodes with GPUs. How would you achieve this node selection and workload placement in Kubernetes?

- A. Use node affinity rules to schedule CPU-intensive tasks on nodes with GPUs and GPU-intensive tasks on nodes without GPUs.
- **B. Use node affinity rules to schedule CPU-intensive tasks on nodes without GPUs and GPU-intensive tasks on nodes with GPUs.**
- C. Use taints and tolerations to dedicate nodes without GPUs to CPU-intensive tasks and nodes with GPUs to GPU-intensive tasks.
- D. Use resource quotas to limit the CPU resources available on nodes with GPUs and the GPU resources available on nodes without GPUs.
- E. Use labels to identify the CPU and GPU-intensive nodes.

Antwort: B

Begründung:

Node affinity rules allow you to specify constraints on which nodes a pod can be scheduled on. By using node affinity rules, you can ensure that CPU-intensive tasks are scheduled on nodes without GPUs and GPU-intensive tasks are scheduled on nodes with GPUs. This optimizes resource utilization and cost. Taints and tolerations can be used, but affinity is more flexible. Resource quotas limit resource usage but do not control placement.

123. Frage

Consider an AI server equipped with two NVIDIA A100 GPUs interconnected with NVLink. You want to maximize the memory bandwidth available to a CUDA application. You observe that the application's performance doesn't scale linearly with the number of GPUs. Which of the following coding techniques or configurations could potentially improve inter-GPU memory access performance?

- **A. Manually manage data transfers between GPUs using 'cudaMemcpyPeer' to exploit NVLink bandwidth. Choose the GPU with more free memory for allocations.**
- B. Employ Unified Memory (I-M) with prefetching to automatically migrate data between GPUs as needed.
- C. Ensure all memory allocations are performed on GPU 0 to minimize data transfer.
- D. Disable NVLink to force the application to use PCIe, which might provide more consistent performance.
- E. Use CUDA-aware MPI for inter-GPU communication to leverage NVLink.

Antwort: A

Begründung:

ScudaMemcpyPeer allows explicit, optimized data transfers between GPUs using NVLink. Unified Memory with prefetching can simplify development, but might not always provide the best performance. CUDA-aware MPI is typically used for inter-node communication, not intra-node GPU-GPU. Allocating all memory on one GPU defeats the purpose of multi-GPU acceleration. PCIe will be slower than NVLink. Manually managing memory transfers, while complex, gives the programmer the most control over leveraging NVLink bandwidth.

124. Frage

You are implementing a distributed deep learning training setup using multiple servers connected via NVLink switches. You want to ensure optimal utilization of the NVLink interconnect. Which of the following strategies would be MOST effective in achieving this goal?

- A. Use a standard TCP/IP socket connection for inter-GPU communication across servers.
- **B. Increase the batch size to reduce the frequency of inter-GPU communication.**
- **C. Configure NCCL to use GPUDirect RDMA for inter-GPU communication across servers.**
- D. Disable peer-to-peer GPU memory access within each server to avoid contention.
- E. Implement a data compression algorithm that can be processed by the CPU before sending data over NVLink.

Antwort: B,C

Begründung:

Configuring NCCL to use GPUDirect RDMA is key as it bypasses the CPU and reduces latency. Increasing batch size helps amortize the cost of communication. TCP/IP will not utilize NVLink; CPU-based compression adds overhead, and disabling peer-to-peer access hurts performance. GPUDirect RDMA allows direct memory access and it is the most critical piece in NVLink switch based system setup. Large batch size would also help to saturate the link bandwidth and less frequently communicate between GPUs.

125. Frage

You are designing an AI infrastructure cluster for training large language models (LLMs). The dataset consists of 10TB of image data and 5TB of text data. You estimate that intermediate training data (checkpoints, temporary files) will require an additional 20TB of storage. You want to use a parallel file system for optimal performance. Considering a replication factor of 2 for data redundancy and a 20% overhead for file system metadata, what is the minimum raw storage capacity you should provision?

- A. 92.4 TB
- B. 70 TB
- C. 42 TB
- D. 84 TB
- E. 100.8 TB

Antwort: A

Begründung:

Total data size: 10TB + 5TB + 20TB = 35TB. With a replication factor of 2, the storage required is $35\text{TB} \times 2 = 70\text{TB}$. Adding 20% overhead for metadata, we get $70\text{TB} \times 1.2 = 84\text{TB}$. Therefore, the minimum raw storage capacity is $84 + 8.4 = 92.4\text{TB}$. Overhead needs to be calculated from after replication is implemented, so replication + 20% overhead.

126. Frage

You are configuring a server with multiple GPUs for CUDA-aware MPI. Which environment variable is critical for ensuring proper GPU affinity, so that each MPI process uses the correct GPU?

- A. LD_LIBRARY_PATH
- B. CUDA_DEVICE_ORDER
- C. CUDA_VISIBLE_DEVICES
- D. CUDA_LAUNCH_BLOCKING
- E. MPI_GPU_SUPPORT

Antwort: C

Begründung:

'CUDA_VISIBLE_DEVICES' is essential for GPU affinity. It allows you to specify which GPUs are visible to a particular process. Without it, all processes might try to use the same GPU, leading to performance bottlenecks. It controls the order in which GPUs are enumerated. It specifies the path to shared libraries. It is hypothetical. It forces synchronous CUDA calls.

127. Frage

.....

Wenn Sie ein Ziel haben, sollen Sie Ihr Ziel ganz mutig erzielen. Jeder IT-Fachmann wird mit den jetzigen einfachen Lebensverhältnissen zufrieden sein. Der Druck in allen Branchen und Gewerben ist sehr groß. In der IT-Branche ist es auch so. Wenn Sie ein Ziel haben, sollen Sie mutig Ihren Traum erfüllen. Auch in der NVIDIA NCP-AII Zertifizierungsprüfung herrscht große Konkurrenz. Durch die NVIDIA NCP-AII Prüfung wird Ihre Berufskarriere sicher ganz anders. Eine glänzende Zukunft wartet schon auf Sie. Unser DeutschPrüfung bietet Ihnen die genauesten und richtigsten NVIDIA NCP-AII Schulungsunterlagen und Ihnen helfen, die Zertifizierungsprüfung zu bestehen und Ihr Ziel zu erreichen.

NCP-AII Prüfungsvorbereitung: <https://www.deutschpruefung.com/NCP-AII-deutsch-pruefungsfragen.html>

Die Anzahl der Angestellten in NCP-AII beträgt mehr als 6,000+, NCP-AII Hilfsmittel Prüfung bietet Sie das Sicherheitsgefühl,

Außerdem sind jetzt einige Teile dieser DeutschPrüfung NCP-AII Prüfungsfragen kostenlos erhältlich:
https://drive.google.com/open?id=1FU5RLNdCO_Q22whH-w31ubgXJyMB_Kys