

Huawei

H13-313_V1.0

Huawei Certified ICT Associate-AI Solution

For More Information – Visit link below:

<https://www.examsempire.com/>

Product Version

1. Up to Date products, reliable and verified.
2. Questions and Answers in PDF Format.



<https://examsempire.com/>

Visit us at: <https://www.examsempire.com/h13-313-v1-0>

Latest Version: 6.0

Question: 1

Throughout the evolution of AI hardware, Moore's Law played a significant role. Why has the slowing of Moore's Law in recent years forced the industry to innovate in specific AI hardware architectures?

- A. It necessitates general-purpose CPUs for all AI tasks.
- B. It requires domain-specific accelerators like NPUs.
- C. It mandates the reduction of all AI model accuracy.
- D. It prevents the development of any future AI models.

Answer: B

Question: 2

When deploying a Large Language Model (LLM) into a production environment, an enterprise must choose between "Full Fine-Tuning" and "Parameter-Efficient Fine-Tuning" (PEFT). Which characteristic of PEFT makes it the preferred approach for many businesses with limited hardware resources?

- A. It updates only a small subset of weights.
- B. It requires training all model parameters.
- C. It deletes the pre-trained knowledge base.
- D. It relies solely on randomized guesses.

Answer: A

Question: 3

How does the deployment of a "Local LLM" (running on-premise) differ from a "Cloud LLM" in terms of enterprise data strategy?

- A. Local LLMs are always faster than Cloud LLMs.
- B. Local LLMs do not require any electricity to run.
- C. Local LLMs offer higher data privacy and sovereignty.
- D. Cloud LLMs are physically impossible to use.

Answer: C

Question: 4

A company is evaluating the "Model-as-a-Service" (MaaS) business model for their internal AI operations. What is the primary operational benefit of MaaS compared to the traditional "On-Premise Custom Build" model?

- A. It eliminates the need for any data usage.
- B. It prevents the use of any internet connection.
- C. It guarantees a 0% error rate for all tasks.
- D. It offloads infrastructure maintenance to providers.

Answer: D

Question: 5

What is the function of a "Model Repository" (e.g., ModelArts model management) in the deployment lifecycle?

- A. To version, store, and manage deployed models.
- B. To keep the server room clean.
- C. To prevent the model from being run.
- D. To manually rewrite the model code.

Answer: A

Question: 6

Consider the historical evolution of AI algorithms. Why was the emergence of the Backpropagation algorithm in the 1980s a critical turning point for multi-layer neural networks?

- A. It replaced the need for GPU-based hardware acceleration.
- B. It completely eliminated the risk of vanishing gradients.
- C. It allowed AI to run on low-power mobile microprocessors.
- D. It provided a method to efficiently train deep network weights.

Answer: D

Question: 7

Which of the following is considered a major organizational challenge when implementing AI in traditional enterprises?

- A. The availability of too many GPUs.
- B. The inability to use programming languages.
- C. Data silos and poor data quality.
- D. The lack of electricity in offices.

Answer: C

Question: 8

During the "Data Labeling" stage, a project manager decides to use "Active Learning." What is the main benefit of this approach?

- A. It forces humans to label every single data point.
- B. It identifies the most informative samples for humans to label.
- C. It completely eliminates the need for human input.
- D. It prevents the model from ever being trained.

Answer: B

Question: 9

Why is "Memory Bandwidth" often a bottleneck in high-performance AI computing scenarios?

- A. Data cannot reach the processor fast enough for math.
- B. Because processors are faster at reading than memory.
- C. Because all memory is stored in the cloud off-site.
- D. Because memory has no impact on AI performance.

Answer: A

Question: 10

In a smart city project, what is the primary advantage of deploying AI inference at the "edge" (e.g., on the camera) rather than sending all video to the cloud?

- A. It increases the cost of network bandwidth.
- B. It forces the cloud to work faster.
- C. It reduces latency and conserves bandwidth.
- D. It makes the model less accurate.

Answer: C

Thank You for Trying Our Product

Special 16 USD Discount Coupon: NSZUBG3X

Email: support@examsempire.com

**Check our Customer Testimonials and ratings
available on every product page.**

Visit our website.

<https://examsempire.com/>