



IAPP

Exam Questions AIGP

Artificial Intelligence Governance Professional

NEW QUESTION 1

- (Topic 1)

CASE STUDY

Please use the following answer the next question:

ABC Corp, is a leading insurance provider offering a range of coverage options to individuals. ABC has decided to utilize artificial intelligence to streamline and improve its customer acquisition and underwriting process, including the accuracy and efficiency of pricing policies.

ABC has engaged a cloud provider to utilize and fine-tune its pre-trained, general purpose large language model (“LLM”). In particular, ABC intends to use its historical customer data—including applications, policies, and claims—and proprietary pricing and risk strategies to provide an initial qualification assessment of potential customers, which would then be routed a human underwriter for final review.

ABC and the cloud provider have completed training and testing the LLM, performed a readiness assessment, and made the decision to deploy the LLM into production. ABC has designated an internal compliance team to monitor the model during the first month, specifically to evaluate the accuracy, fairness, and reliability of its output. After the first month in production, ABC realizes that the LLM declines a higher percentage of women's loan applications due primarily to women historically receiving lower salaries than men.

What is the best strategy to mitigate the bias uncovered in the loan applications?

- A. Retrain the model with data that reflects demographic parity.
- B. Procure a third-party statistical bias assessment tool.
- C. Document all instances of bias in the data set.
- D. Delete all gender-based data in the data set.

Answer: A

Explanation:

Retraining the model with data that reflects demographic parity is the best strategy to mitigate the bias uncovered in the loan applications. This approach addresses the root cause of the bias by ensuring that the training data is representative and balanced, leading to more equitable decision-making by the AI model.

Reference: The AIGP Body of Knowledge stresses the importance of using high-quality, unbiased training data to develop fair and reliable AI systems. Retraining the model with balanced data helps correct biases that arise from historical inequalities, ensuring that the AI system makes decisions based on equitable criteria.

NEW QUESTION 2

- (Topic 1)

An EU bank intends to launch a multi-modal AI platform for customer engagement and automated decision-making assist with the opening of bank accounts. The platform has been subject to thorough risk assessments and testing, where it proves to be effective in not discriminating against any individual on the basis of a protected class.

What additional obligations must the bank fulfill prior to deployment?

- A. The bank must obtain explicit consent from users under the privacy Directive.
- B. The bank must disclose how the AI system works under the EII Digital Services Act.
- C. The bank must subject the AI system an adequacy decision and publish its appropriate safeguards.
- D. The bank must disclose the use of the AI system and implement suitable measures for users to contest automated decision-making.

Answer: D

Explanation:

Under the EU regulations, particularly the GDPR, banks using AI for decision-making must inform users about the use of AI and provide mechanisms for users to contest decisions. This is part of ensuring transparency and accountability in automated processing. Explicit consent under the privacy directive (A) and disclosing under the Digital Services Act (B) are not specifically required in this context. An adequacy decision is related to data transfers outside the EU (C).

NEW QUESTION 3

- (Topic 1)

CASE STUDY

Please use the following answer the next question:

ABC Corp, is a leading insurance provider offering a range of coverage options to individuals. ABC has decided to utilize artificial intelligence to streamline and improve its customer acquisition and underwriting process, including the accuracy and efficiency of pricing policies.

ABC has engaged a cloud provider to utilize and fine-tune its pre-trained, general purpose large language model (“LLM”). In particular, ABC intends to use its historical customer data—including applications, policies, and claims—and proprietary pricing and risk strategies to provide an initial qualification assessment of potential customers, which would then be routed a human underwriter for final review.

ABC and the cloud provider have completed training and testing the LLM, performed a readiness assessment, and made the decision to deploy the LLM into production. ABC has designated an internal compliance team to monitor the model during the first month, specifically to evaluate the accuracy, fairness, and reliability of its output. After the first month in production, ABC realizes that the LLM declines a higher percentage of women's loan applications due primarily to women historically receiving lower salaries than men.

Which of the following is the most important reason to train the underwriters on the model prior to deployment?

- A. To provide a reminder of a right appeal.
- B. To solicit on-going feedback on model performance.
- C. To apply their own judgment to the initial assessment.
- D. To ensure they provide transparency applicants on the model.

Answer: C

Explanation:

Training underwriters on the model prior to deployment is crucial so they can apply their own judgment to the initial assessment. While AI models can streamline the process, human judgment is still essential to catch nuances that the model might miss or to account for any biases or errors in the model's decision-making process.

Reference: The AIGP Body of Knowledge emphasizes the importance of human oversight in AI systems, particularly in high-stakes areas such as underwriting and loan approvals. Human underwriters can provide a critical review and ensure that the model's assessments are accurate and fair, integrating their expertise and understanding of complex cases.

NEW QUESTION 4

- (Topic 1)

Which of the following best defines an "AI model"?

- A. A system that applies defined rules to execute tasks.
- B. A system of controls that is used to govern an AI algorithm.
- C. A corpus of data which an AI algorithm analyzes to make predictions.
- D. A program that has been trained on a set of data to find patterns within the data.

Answer: D

Explanation:

An AI model is best defined as a program that has been trained on a set of data to find patterns within that data. This definition captures the essence of machine learning, where the model learns from the data to make predictions or decisions. Reference: AIGP BODY OF KNOWLEDGE, which provides a detailed explanation of AI models and their training processes.

NEW QUESTION 5

- (Topic 1)

According to the Singapore Model AI Governance Framework, all of the following are recommended measures to promote the responsible use of AI EXCEPT?

- A. Determining the level of human involvement in algorithmic decision-making.
- B. Adapting the existing governance structure algorithmic decision-making.
- C. Employing human-over-the-loop protocols for high-risk systems.
- D. Establishing communications and collaboration among stakeholders.

Answer: C

Explanation:

The Singapore Model AI Governance Framework recommends several measures to promote the responsible use of AI, such as determining the level of human involvement in decision-making, adapting governance structures, and establishing communications and collaboration among stakeholders. However, employing human-over-the-loop protocols is not specifically mentioned in this framework. The focus is more on integrating human oversight appropriately within the decision-making process rather than exclusively employing such protocols. Reference: AIGP Body of Knowledge, section on AI governance frameworks.

NEW QUESTION 6

- (Topic 1)

All of the following are common optimization techniques in deep learning to determine weights that represent the strength of the connection between artificial neurons EXCEPT?

- A. Gradient descent, which initially sets weights arbitrary values, and then at each step changes them.
- B. Momentum, which improves the convergence speed and stability of neural network training.
- C. Autoregression, which analyzes and makes predictions about time-series data.
- D. Backpropagation, which starts from the last layer working backwards.

Answer: C

Explanation:

Autoregression is not a common optimization technique in deep learning to determine weights for artificial neurons. Common techniques include gradient descent, momentum, and backpropagation. Autoregression is more commonly associated with time-series analysis and forecasting rather than neural network optimization. Reference: AIGP BODY OF KNOWLEDGE, which discusses common optimization techniques used in deep learning.

NEW QUESTION 7

- (Topic 1)

According to the EU AI Act, providers of what kind of machine learning systems will be required to register with an EU oversight agency before placing their systems in the EU market?

- A. AI systems that are harmful based on a legal risk-utility calculation.
- B. AI systems that are "strong" general intelligence.
- C. AI systems trained on sensitive personal data.
- D. AI systems that are high-risk.

Answer: D

Explanation:

According to the EU AI Act, providers of high-risk AI systems are required to register with an EU oversight agency before these systems can be placed on the market. This requirement is part of the Act's framework to ensure that high-risk AI systems comply with stringent safety, transparency, and accountability standards. High-risk systems are those that pose significant risks to health, safety, or fundamental rights. Registration with oversight agencies helps facilitate ongoing monitoring and enforcement of compliance with the Act's provisions. Systems categorized under other criteria, such as those trained on sensitive personal data or exhibiting "strong" general intelligence, also fall under scrutiny but are primarily covered under different regulatory requirements or classifications.

NEW QUESTION 8

- (Topic 1)

If it is possible to provide a rationale for a specific output of an AI system, that system can best be described as?

- A. Accountable.
- B. Transparent.
- C. Explainable.
- D. Reliable.

Answer: C

Explanation:

If it is possible to provide a rationale for a specific output of an AI system, that system can best be described as explainable. Explainability in AI refers to the ability to interpret and understand the decision-making process of the AI system. This involves being able to articulate the factors and logic that led to a particular output or decision. Explainability is critical for building trust, enabling users to understand and validate the AI system's actions, and ensuring compliance with ethical and regulatory standards. It also facilitates debugging and improving the system by providing insights into its behavior.

NEW QUESTION 9

- (Topic 1)

What is the 1956 Dartmouth summer research project on AI best known as?

- A. A meeting focused on the impacts of the launch of the first mass-produced computer.
- B. A research project on the impacts of technology on society.
- C. A research project to create a test for machine intelligence.
- D. A meeting focused on the founding of the AI field.

Answer: D

Explanation:

The 1956 Dartmouth summer research project on AI is best known as a meeting focused on the founding of the AI field. This conference is historically significant because it marked the formal beginning of artificial intelligence as an academic discipline. The term "artificial intelligence" was coined during this event, and it laid the foundation for future research and development in AI.

Reference: The AIGP Body of Knowledge highlights the importance of the Dartmouth Conference as a pivotal moment in the history of AI, which established AI as a distinct field of study and research.

NEW QUESTION 10

- (Topic 1)

CASE STUDY

Please use the following answer the next question:

XYZ Corp., a premier payroll services company that employs thousands of people globally, is embarking on a new hiring campaign and wants to implement policies and procedures to identify and retain the best talent. The new talent will help the company's product team expand its payroll offerings to companies in the healthcare and transportation sectors, including in Asia.

It has become time consuming and expensive for HR to review all resumes, and they are concerned that human reviewers might be susceptible to bias.

Address these concerns, the company is considering using a third-party AI tool to screen resumes and assist with hiring. They have been talking to several vendors about possibly obtaining a third-party AI-enabled hiring solution, as long as it would achieve its goals and comply with all applicable laws.

The organization has a large procurement team that is responsible for the contracting of technology solutions. One of the procurement team's goals is to reduce costs, and it often prefers lower-cost solutions. Others within the company are responsible for integrating and deploying technology solutions into the organization's operations in a responsible, cost-effective manner.

The organization is aware of the risks presented by AI hiring tools and wants to mitigate them. It also questions how best to organize and train its existing personnel to use the AI hiring tool responsibly. Their concerns are heightened by the fact that relevant laws vary across jurisdictions and continue to change.

The frameworks that would be most appropriate for XYZ's governance needs would be the NIST AI Risk Management Framework and?

- A. NIST Information Security Risk (NIST SP 800-39).
- B. NIST Cyber Security Risk Management Framework (CSF 2.0).
- C. IEEE Ethical System Design Risk Management Framework (IEEE 7000-21).
- D. Human Rights, Democracy, and Rule of Law Impact Assessment (HUDERIA).

Answer: C

Explanation:

The IEEE Ethical System Design Risk Management Framework (IEEE 7000-21) would be most appropriate for XYZ Corp's governance needs in addition to the NIST AI Risk Management Framework. The IEEE framework specifically addresses ethical concerns during system design, which is crucial for ensuring the responsible use of AI in hiring. It complements the NIST framework by focusing on ethical risk management, aligning well with XYZ Corp's goals of deploying AI responsibly and mitigating associated risks.

NEW QUESTION 10

- (Topic 1)

Which of the following disclosures is NOT required for an EU organization that developed and deployed a high-risk AI system?

- A. The human oversight measures employed.
- B. How an individual may contest a decision.
- C. The location(s) where data is stored.
- D. The fact that an AI system is being used.

Answer: C

Explanation:

Under the EU AI Act, organizations that develop and deploy high-risk AI systems are required to provide several key disclosures to ensure transparency and accountability. These include the human oversight measures employed, how individuals can contest decisions made by the AI system, and informing individuals that an AI system is being used. However, there is no specific requirement to disclose the exact locations where data is stored. The focus of the Act is on the transparency of the AI system's operation and its impact on individuals, rather than on the technical details of data storage locations.

NEW QUESTION 11

- (Topic 1)

CASE STUDY

Please use the following answer the next question:

XYZ Corp., a premier payroll services company that employs thousands of people globally, is embarking on a new hiring campaign and wants to implement policies and procedures to identify and retain the best talent. The new talent will help the company's product team expand its payroll offerings to companies in the

healthcare and transportation sectors, including in Asia.

It has become time consuming and expensive for HR to review all resumes, and they are concerned that human reviewers might be susceptible to bias.

Address these concerns, the company is considering using a third-party AI tool to screen resumes and assist with hiring. They have been talking to several vendors about possibly obtaining a third-party AI-enabled hiring solution, as long as it would achieve its goals and comply with all applicable laws.

The organization has a large procurement team that is responsible for the contracting of technology solutions. One of the procurement team's goals is to reduce costs, and it often prefers lower-cost solutions. Others within the company are responsible for integrating and deploying technology solutions into the organization's operations in a responsible, cost-effective manner.

The organization is aware of the risks presented by AI hiring tools and wants to mitigate

them. It also questions how best to organize and train its existing personnel to use the AI hiring tool responsibly. Their concerns are heightened by the fact that relevant laws vary across jurisdictions and continue to change.

If XYZ does not deploy and use the AI hiring tool responsibly in the United States, its liability would likely increase under all of the following laws EXCEPT?

- A. Anti-discrimination laws.
- B. Product liability laws.
- C. Accessibility laws.
- D. Privacy laws.

Answer: B

Explanation:

In the United States, the use of AI hiring tools must comply with anti-discrimination laws, accessibility laws, and privacy laws to avoid increasing liability. Anti-discrimination laws (A) ensure that hiring practices do not unlawfully discriminate against protected classes. Accessibility laws (C) require that hiring tools are accessible to all applicants, including those with disabilities. Privacy laws (D) govern the handling of personal data during the hiring process. Product liability laws (B), however, typically apply to the safety and reliability of physical products and would not generally increase liability specifically related to the responsible use of AI hiring tools in the employment context.

NEW QUESTION 15

- (Topic 1)

What is the key feature of Graphical Processing Units (GPUs) that makes them well-suited to running AI applications?

- A. GPUs run many tasks concurrently, resulting in faster processing.
- B. GPUs can access memory quickly, resulting in lower latency than CPUs.
- C. GPUs can run every task on a computer, making them more robust than CPUs.
- D. The number of transistors on GPUs doubles every two years, making the chips smaller and lighter.

Answer: A

Explanation:

GPUs (Graphical Processing Units) are well-suited to running AI applications due to their ability to run many tasks concurrently, which significantly enhances processing speed. This parallel processing capability makes GPUs ideal for handling the large-scale computations required in AI and deep learning tasks.

Reference: AIGP BODY OF KNOWLEDGE, which explains the importance of compute infrastructure in AI applications.

NEW QUESTION 19

- (Topic 1)

Which of the following is NOT a common type of machine learning?

- A. Deep learning.
- B. Cognitive learning.
- C. Unsupervised learning.
- D. Reinforcement learning.

Answer: B

Explanation:

The common types of machine learning include supervised learning, unsupervised learning, reinforcement learning, and deep learning. Cognitive learning is not a type of machine learning; rather, it is a term often associated with the broader field of cognitive science and psychology. Reference: AIGP BODY OF KNOWLEDGE and standard AI/ML literature.

NEW QUESTION 20

- (Topic 1)

Which of the following most encourages accountability over AI systems?

- A. Determining the business objective and success criteria for the AI project.
- B. Performing due diligence on third-party AI training and testing data.
- C. Defining the roles and responsibilities of AI stakeholders.
- D. Understanding AI legal and regulatory requirements.

Answer: C

Explanation:

Defining the roles and responsibilities of AI stakeholders is crucial for encouraging accountability over AI systems. Clear delineation of who is responsible for different aspects of the AI lifecycle ensures that there is a person or team accountable for monitoring, maintaining, and addressing issues that arise. This accountability framework helps in ensuring that ethical standards and regulatory requirements are met, and it facilitates transparency and traceability in AI operations. By assigning specific roles, organizations can better manage and mitigate risks associated with AI deployment and use.

NEW QUESTION 25

- (Topic 1)

A Canadian company is developing an AI solution to evaluate candidates in the course of job interviews.

Before offering the AI solution in the EU market, the company must take all of the following steps EXCEPT?

- A. Register the AI solution in a public EU database.
- B. Establish a risk and quality management system.
- C. Engage a third-party auditor to perform a bias audit.
- D. Draw up technical documentation and instructions for use.

Answer: A

Explanation:

Before offering an AI solution in the EU market, a Canadian company must take several steps to comply with the EU AI Act. These steps include establishing a risk and quality management system (B), engaging a third-party auditor to perform a bias audit (C), and drawing up technical documentation and instructions for use (D). However, there is no requirement to register the AI solution in a public EU database (A). This registration step is not specified as part of the compliance requirements under the EU AI Act for such solutions.

NEW QUESTION 27

- (Topic 1)

Which of the following is a subcategory of AI and machine learning that uses labeled datasets to train algorithms?

- A. Segmentation.
- B. Generative AI.
- C. Expert systems.
- D. Supervised learning.

Answer: D

Explanation:

Supervised learning is a subcategory of AI and machine learning where labeled datasets are used to train algorithms. This process involves feeding the algorithm a dataset where the input-output pairs are known, allowing the algorithm to learn and make predictions or decisions based on new, unseen data. Reference: AIGP BODY OF KNOWLEDGE, which describes supervised learning as a model trained on labeled data (e.g., text recognition, detecting spam in emails).

NEW QUESTION 28

- (Topic 2)

All of the following are elements of establishing a global AI governance infrastructure EXCEPT?

- A. Providing training to foster a culture that promotes ethical behavior.
- B. Creating policies and procedures to manage third-party risk.
- C. Understanding differences in norms across countries.
- D. Publicly disclosing ethical principles.

Answer: D

Explanation:

Establishing a global AI governance infrastructure involves several key elements, including providing training to foster a culture that promotes ethical behavior, creating policies and procedures to manage third-party risk, and understanding differences in norms across countries. While publicly disclosing ethical principles can enhance transparency and trust, it is not a core element necessary for the establishment of a governance infrastructure. The focus is more on internal processes and structures rather than public disclosure. Reference: AIGP Body of Knowledge on AI Governance and Infrastructure.

NEW QUESTION 29

- (Topic 2)

What is the best reason for a company adopt a policy that prohibits the use of generative AI?

- A. Avoid using technology that cannot be monetized.
- B. Avoid needing to identify and hire qualified resources.
- C. Avoid the time necessary to train employees on acceptable use.
- D. Avoid accidental disclosure to its confidential and proprietary information.

Answer: D

Explanation:

The primary concern for a company adopting a policy prohibiting the use of generative AI is the risk of accidental disclosure of confidential and proprietary information. Generative AI tools can inadvertently leak sensitive data during the creation process or through data sharing. This risk outweighs the other reasons listed, as protecting sensitive information is critical to maintaining the company's competitive edge and legal compliance. This rationale is discussed in the sections on risk management and data privacy in the IAPP AIGP Body of Knowledge.

NEW QUESTION 31

- (Topic 2)

Which of the following would be the least likely step for an organization to take when designing an integrated compliance strategy for responsible AI?

- A. Conducting an assessment of existing compliance programs to determine overlaps and integration points.
- B. Employing a new software platform to modernize existing compliance processes across the organization.
- C. Consulting experts to consider the ethical principles underpinning the use of AI within the organization.
- D. Launching a survey to understand the concerns and interests of potentially impacted stakeholders.

Answer: B

Explanation:

When designing an integrated compliance strategy for responsible AI, the least likely step would be employing a new software platform to modernize existing compliance processes. While modernizing compliance processes is beneficial, it is not as directly related to the strategic integration of ethical principles and stakeholder concerns. More critical steps include conducting assessments of existing compliance programs to identify overlaps and integration points, consulting

experts on ethical principles, and launching surveys to understand stakeholder concerns. These steps ensure that the compliance strategy is comprehensive and aligned with responsible AI principles. Reference: AIGP Body of Knowledge on AI Governance and Compliance Integration.

NEW QUESTION 34

- (Topic 2)

CASE STUDY

Please use the following answer the next question:

A local police department in the United States procured an AI system to monitor and analyze social media feeds, online marketplaces and other sources of public information to detect evidence of illegal activities (e.g., sale of drugs or stolen goods). The AI system works by surveilling the public sites in order to identify individuals that are likely to have committed a crime. It cross-references the individuals against data maintained by law enforcement and then assigns a percentage score of the likelihood of criminal activity based on certain factors like previous criminal history, location, time, race and gender.

The police department retained a third-party consultant assist in the procurement process, specifically to evaluate two finalists. Each of the vendors provided information about their system's accuracy rates, the diversity of their training data and how their system works. The consultant determined that the first vendor's system has a higher accuracy rate and based on this information, recommended this vendor to the police department.

The police department chose the first vendor and implemented its AI system. As part of the implementation, the department and consultant created a usage policy for the system, which includes training police officers on how the system works and how to incorporate it into their investigation process.

The police department has now been using the AI system for a year. An internal review has found that every time the system scored a likelihood of criminal activity at or above 90%, the police investigation subsequently confirmed that the individual had, in fact, committed a crime. Based on these results, the police department wants to forego investigations for cases where the AI system gives a score of at least 90% and proceed directly with an arrest.

When notifying an accused perpetrator, what additional information should a police officer provide about the use of the AI system?

- A. Information about the accuracy of the AI system.
- B. Information about how the accused can oppose the charges.
- C. Information about the composition of the training data of the system.
- D. Information about how the individual was identified by the AI system.

Answer: D

Explanation:

When notifying an accused perpetrator, the police officer should provide information about how the individual was identified by the AI system. This transparency is crucial for maintaining trust and ensuring that the accused understands the basis of the charges against them. Information about the accuracy, how to oppose the charges, and the composition of the training data, while potentially relevant, do not directly address the immediate need for the accused to understand the specific process that led to their identification. Reference: AIGP Body of Knowledge on AI Transparency and Explainability.

NEW QUESTION 39

- (Topic 2)

What is the primary purpose of conducting ethical red-teaming on an AI system?

- A. To improve the model's accuracy.
- B. To simulate model risk scenarios.
- C. To identify security vulnerabilities.
- D. To ensure compliance with applicable law.

Answer: B

Explanation:

The primary purpose of conducting ethical red-teaming on an AI system is to simulate model risk scenarios. Ethical red-teaming involves rigorously testing the AI system to identify potential weaknesses, biases, and vulnerabilities by simulating real-world attack or failure scenarios. This helps in proactively addressing issues that could compromise the system's reliability, fairness, and security. Reference: AIGP Body of Knowledge on AI Risk Management and Ethical AI Practices.

NEW QUESTION 40

- (Topic 2)

A company initially intended to use a large data set containing personal information to train an AI model. After consideration, the company determined that it can derive enough value from the data set without any personal information and permanently obfuscated all personal data elements before training the model.

This is an example of applying which privacy-enhancing technique (PET)?

- A. Anonymization.
- B. Pseudonymization.
- C. Differential privacy.
- D. Federated learning.

Answer: A

Explanation:

Anonymization is a privacy-enhancing technique that involves removing or permanently altering personal data elements to prevent the identification of individuals. In this case, the company obfuscated all personal data elements before training the model, which aligns with the definition of anonymization. This ensures that the data cannot be traced back to individuals, thereby protecting their privacy while still allowing the company to derive value from the dataset. Reference: AIGP Body of Knowledge, privacy-enhancing techniques section.

NEW QUESTION 41

- (Topic 2)

What is the term for an algorithm that focuses on making the best choice achieve an immediate objective at a particular step or decision point, based on the available information and without regard for the longer-term best solutions?

- A. Single-lane.
- B. Optimized.
- C. Efficient.
- D. Greedy.

Answer: D

Explanation:

A greedy algorithm is one that makes the best choice at each step to achieve an immediate objective, without considering the longer-term consequences. It focuses on local optimization at each decision point with the hope that these local solutions will lead to an optimal global solution. However, greedy algorithms do not always produce the best overall solution for certain problems, but they are useful when an immediate, locally optimal solution is desired. Reference: AIGP Body of Knowledge, algorithm types section.

NEW QUESTION 42

- (Topic 2)

After completing model testing and validation, which of the following is the most important step that an organization takes prior to deploying the model into production?

- A. Perform a readiness assessment.
- B. Define a model-validation methodology.
- C. Document maintenance teams and processes.
- D. Identify known edge cases to monitor post-deployment.

Answer: A

Explanation:

After completing model testing and validation, the most important step prior to deploying the model into production is to perform a readiness assessment. This assessment ensures that the model is fully prepared for deployment, addressing any potential issues related to infrastructure, performance, security, and compliance. It verifies that the model meets all necessary criteria for a successful launch. Other steps, such as defining a model-validation methodology, documenting maintenance teams and processes, and identifying known edge cases, are also important but come secondary to confirming overall readiness. Reference: AIGP Body of Knowledge on Deployment Readiness.

NEW QUESTION 45

- (Topic 2)

CASE STUDY

Please use the following answer the next question:

A mid-size US healthcare network has decided to develop an AI solution to detect a type of cancer that is most likely arise in adults. Specifically, the healthcare network intends to create a recognition algorithm that will perform an initial review of all imaging and then route records a radiologist for secondary review pursuant agreed-upon criteria (e.g., a confidence score below a threshold).

To date, the healthcare network has taken the following steps: defined its AI ethical principles; conducted discovery to identify the intended uses and success criteria for the system; established an AI governance committee; assembled a broad, crossfunctional team with clear roles and responsibilities; and created policies and procedures to document standards, workflows, timelines and risk thresholds during the project.

The healthcare network intends to retain a cloud provider to host the solution and a consulting firm to help develop the algorithm using the healthcare network's existing data and de-identified data that is licensed from a large US clinical research partner.

In the design phase, what is the most important step for the healthcare network to take when mapping its existing data to the clinical research partner data?

- A. Apply privacy-enhancing technologies to the data.
- B. Identify fits and gaps in the combined data.
- C. Ensure the data is labeled and formatted.
- D. Evaluate the country of origin of the data.

Answer: B

Explanation:

In the design phase of integrating data from different sources, identifying fits and gaps is crucial. This process involves understanding how well the data from the clinical research partner aligns with the healthcare network's existing data. It ensures that the combined data set is coherent and can be effectively used for training the AI algorithm. This step helps in spotting any discrepancies, inconsistencies, or missing data that might affect the performance and accuracy of the AI model. It directly addresses the integrity and compatibility of the data, which is foundational before applying any privacy-enhancing technologies, labeling, or evaluating the origin of the data. Reference: AIGP Body of Knowledge on Data Integration and Quality.

NEW QUESTION 48

- (Topic 2)

All of the following are included within the scope of post-deployment AI maintenance EXCEPT?

- A. Ensuring that all model components are subject a control framework.
- B. Dedicating experts to continually monitor the model output.
- C. Evaluating the need for an audit under certain standards.
- D. Defining thresholds to conduct new impact assessments.

Answer: D

Explanation:

Post-deployment AI maintenance typically includes ensuring that all model components are subject to a control framework, dedicating experts to continually monitor the model output, and evaluating the need for audits under certain standards. However, defining thresholds to conduct new impact assessments is usually part of the initial deployment and ongoing governance processes rather than a maintenance activity. Maintenance focuses more on the operational aspects of the AI system rather than setting new thresholds for impact assessments.

Reference: AIGP BODY OF KNOWLEDGE, sections discussing AI lifecycle management and post-deployment activities.

NEW QUESTION 50

- (Topic 2)

You are an engineer that developed an AI-based ad recommendation tool. Which of the following should be monitored to evaluate the tool's effectiveness?

- A. Output data, assess the delta between the prediction and actual ad clicks.
- B. Algorithmic patterns, to show the model has a high degree of accuracy.

- C. Input data, to ensure the ads are reaching the target audience.
- D. GPU performance, to evaluate the tool's robustness.

Answer: A

Explanation:

To evaluate the effectiveness of an AI-based ad recommendation tool, the most relevant metric is the output data, specifically assessing the delta between the prediction and actual ad clicks. This metric directly measures the tool's accuracy and effectiveness in making accurate recommendations that lead to user engagement. While monitoring algorithmic patterns and input data can provide insights into the model's behavior and targeting accuracy, and GPU performance can indicate the robustness and efficiency of the tool, the primary indicator of effectiveness for an ad recommendation tool is how well it predicts actual ad clicks. Reference: AIGP BODY OF KNOWLEDGE, sections on AI performance metrics and evaluation methods.

NEW QUESTION 54

- (Topic 2)

Pursuant to the White House Executive Order of November 2023, who is responsible for creating guidelines to conduct red-teaming tests of AI systems?

- A. National Institute of Standards and Technology (NIST).
- B. National Science and Technology Council (NSTC).
- C. Office of Science and Technology Policy (OSTP).
- D. Department of Homeland Security (DHS).

Answer: A

Explanation:

The White House Executive Order of November 2023 designates the National Institute of Standards and Technology (NIST) as the responsible body for creating guidelines to conduct red-teaming tests of AI systems. NIST is tasked with developing and providing standards and frameworks to ensure the security, reliability, and ethical deployment of AI systems, including conducting rigorous red-teaming exercises to identify vulnerabilities and assess risks in AI systems.

Reference: AIGP BODY OF KNOWLEDGE, sections on AI governance and regulatory frameworks, and the White House Executive Order of November 2023.

NEW QUESTION 59

- (Topic 2)

All of the following are reasons to deploy a challenger AI model in addition a champion AI model EXCEPT to?

- A. Provide a framework to consider alternatives to the champion model.
- B. Automate real-time monitoring of the champion model.
- C. Perform testing on the champion model.
- D. Retrain the champion model.

Answer: D

Explanation:

Deploying a challenger AI model alongside a champion model is a strategy used to compare the performance of different models in a real-world environment. This approach helps in providing a framework to consider alternatives to the champion model, automating real-time monitoring of the champion model, and performing testing on the champion model. However, retraining the champion model is not a reason to deploy a challenger model. Retraining is a separate process that involves updating the champion model with new data or techniques, which is not related to the use of a challenger model.

Reference: AIGP BODY OF KNOWLEDGE, sections on model evaluation and management.

NEW QUESTION 63

- (Topic 2)

CASE STUDY

Please use the following answer the next question:

A mid-size US healthcare network has decided to develop an AI solution to detect a type of cancer that is most likely arise in adults. Specifically, the healthcare network intends to create a recognition algorithm that will perform an initial review of all imaging and then route records a radiologist for secondary review pursuant agreed-upon criteria (e.g., a confidence score below a threshold).

To date, the healthcare network has taken the following steps: defined its AI ethical principles: conducted discovery to identify the intended uses and success criteria for the system: established an AI governance committee; assembled a broad, crossfunctional team with clear roles and responsibilities; and created policies and procedures to document standards, workflows, timelines and risk thresholds during the project.

The healthcare network intends to retain a cloud provider to host the solution and a consulting firm to help develop the algorithm using the healthcare network's existing data and de-identified data that is licensed from a large US clinical research partner.

Which stakeholder group is most important in selecting the specific type of algorithm?

- A. The cloud provider.
- B. The consulting firm.
- C. The healthcare network's data science team.
- D. The healthcare network's AI governance committee.

Answer: C

Explanation:

In selecting the specific type of algorithm for the AI solution, the healthcare network's data science team is most important. This team possesses the technical expertise and understanding of the data, the clinical context, and the performance requirements needed to make an informed decision about which algorithm is most suitable. While the cloud provider and consulting firm can offer support and infrastructure, and the AI governance committee provides oversight, the data science team's specialized knowledge is crucial for selecting and implementing the appropriate algorithm. Reference: AIGP Body of Knowledge, AI governance and team roles section.

NEW QUESTION 67

.....

Thank You for Trying Our Product

We offer two products:

1st - We have Practice Tests Software with Actual Exam Questions

2nd - Questons and Answers in PDF Format

AIGP Practice Exam Features:

- * AIGP Questions and Answers Updated Frequently
- * AIGP Practice Questions Verified by Expert Senior Certified Staff
- * AIGP Most Realistic Questions that Guarantee you a Pass on Your FirstTry
- * AIGP Practice Test Questions in Multiple Choice Formats and Updatesfor 1 Year

100% Actual & Verified — Instant Download, Please Click
[Order The AIGP Practice Test Here](#)