

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 to 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.

Each of the following steps would support fairness testing by the compliance team during the first month in production EXCEPT?

- A. Validating a similar level of decision-making across different demographic groups.
- B. Providing the loan applicants with information about the model capabilities and limitations.
- C. Identifying if additional training data should be collected for specific demographic groups.
- D. Using tools to help understand factors that may account for differences in decision-making.

Answer: B

Explanation:

Providing the loan applicants with information about the model capabilities and limitations would not directly support fairness testing by the compliance team. Fairness testing focuses on evaluating the model's decisions for biases and ensuring equitable treatment across different demographic groups, rather than informing applicants about the model.

Reference: The AIGP Body of Knowledge outlines that fairness testing involves technical assessments such as validating decision-making consistency across demographics and using tools to understand decision factors. While transparency to applicants is important for ethical AI use, it does not contribute directly to the technical process of fairness testing.

NEW QUESTION 2

- (Topic 1)

Each of the following actors are typically engaged in the AI development life cycle EXCEPT?

- A. Data architects.
- B. Government regulators.
- C. Socio-cultural and technical experts.
- D. Legal and privacy governance experts.

Answer: B

Explanation:

Typically, actors involved in the AI development life cycle include data architects (who design the data frameworks), socio-cultural and technical experts (who ensure the AI system is socio-culturally aware and technically sound), and legal and privacy governance experts (who handle the legal and privacy aspects). Government regulators, while important, are not directly engaged in the development process but rather oversee and regulate the industry. Reference: AIGP BODY OF KNOWLEDGE and AI development frameworks.

NEW QUESTION 3

- (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 4

- (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.

To 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 5

- (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 .. 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.

During the first month when ABC monitors the model for bias, it is most important to?

- A. Continue disparity testing.
- B. Analyze the quality of the training and testing data.
- C. Compare the results to human decisions prior to deployment.
- D. Seek approval from management for any changes to the model.

Answer: A

Explanation:

During the first month of monitoring the model for bias, it is most important to continue disparity testing. Disparity testing involves regularly evaluating the model's decisions to identify and address any biases, ensuring that the model operates fairly across different demographic groups.

Reference: Regular disparity testing is highlighted in the AIGP Body of Knowledge as a

critical practice for maintaining the fairness and reliability of AI models. By continuously monitoring for and addressing disparities, organizations can ensure their AI systems remain compliant with ethical and legal standards, and mitigate any unintended biases that may arise in production.

NEW QUESTION 6

- (Topic 1)

According to the GDPR, what is an effective control to prevent a determination based solely on automated decision-making?

- A. Provide a just-in-time notice about the automated decision-making logic.
- B. Define suitable measures to safeguard personal data.
- C. Provide a right to review automated decision.
- D. Establish a human-in-the-loop procedure.

Answer: D

Explanation:

The GDPR requires that individuals have the right to not be subject to decisions based solely on automated processing, including profiling, unless specific exceptions apply. One effective control is to establish a human-in-the-loop procedure (D), ensuring human oversight and the ability to contest decisions. This goes beyond just-in- time notices (A), data safeguarding (B), or review rights (C), providing a more robust mechanism to protect individuals' rights.

NEW QUESTION 7

- (Topic 1)

CASE STUDY

Please use the following answer the next question:

Good Values Corporation (GVC) is a U.S. educational services provider that employs teachers to create and deliver enrichment courses for high school students. GVC has learned that many of its teacher employees are using generative AI to create the enrichment courses, and that many of the students are using generative AI to complete their assignments.

In particular, GVC has learned that the teachers they employ used open source large language models ("LLM") to develop an online tool that customizes study questions for individual students. GVC has also discovered that an art teacher has expressly incorporated the use of generative AI into the curriculum to enable students to use prompts to create digital art.

GVC has started to investigate these practices and develop a process to monitor any use of generative AI, including by teachers and students, going forward.

What is the best reason for GVC to offer students the choice to utilize generative AI in limited, defined circumstances?

- A. Toenable students to learn how to manage their time.
- B. Toenable students to learn about performing research.
- C. Toenable students to learn about practical applications of AI.
- D. Toenable students to learn how to use AI as a supportive educational tool.

Answer: D

Explanation:

The best reason for GVC to offer students the choice to utilize generative AI in limited, defined circumstances is to enable students to learn how to use AI as a supportive educational tool. By integrating AI in a controlled manner, students can learn the practical applications of AI and develop skills to use AI responsibly and effectively in their educational pursuits.

Reference: The AIGP Body of Knowledge highlights the importance of teaching students about AI's practical applications and the responsible use of AI technologies. This aligns with the goal of fostering a better understanding of AI's role and its potential benefits in various contexts, including education.

NEW QUESTION 8

- (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.

Which of the following measures should XYZ adopt to best mitigate its risk of reputational harm from using the AI tool?

- A. Test the AI tool pre- and post-deployment.
- B. Ensure the vendor assumes responsibility for all damages.
- C. Direct the procurement team to select the most economical AI tool.
- D. Continue to require XYZ's hiring personnel to manually screen all applicants.

Answer: A

Explanation:

To mitigate the risk of reputational harm from using an AI hiring tool, XYZ Corp should rigorously test the AI tool both before and after deployment. Pre-deployment testing ensures the tool works correctly and does not introduce bias or other issues. Post- deployment testing ensures the tool continues to operate as intended and adapts to any changes in data or usage patterns. This approach helps to identify and address potential issues proactively, thereby reducing the risk of reputational harm. Ensuring the vendor assumes responsibility for damages (B) does not address the root cause of potential issues, selecting the most economical tool (C) may compromise quality, and continuing manual screening (D) defeats the purpose of using the AI tool.

NEW QUESTION 9

- (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 10

- (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 10

- (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 11

- (Topic 1)

Under the NIST AI Risk Management Framework, all of the following are defined as characteristics of trustworthy AI EXCEPT?

- A. Tested and Effective.
- B. Secure and Resilient.
- C. Explainable and Interpretable.
- D. Accountable and Transparent.

Answer: A

Explanation:

The NIST AI Risk Management Framework outlines several characteristics of trustworthy AI, including being secure and resilient, explainable and interpretable, and accountable and transparent. While being tested and effective is important, it is not explicitly listed as a characteristic of trustworthy AI in the NIST framework. The focus is more on the system's ability to function safely, securely, and transparently in a way that stakeholders can understand and trust. Reference: AIGP Body of Knowledge, NIST AI RMF section.

NEW QUESTION 14

- (Topic 1)

A company is working to develop a self-driving car that can independently decide the appropriate route to take the driver after the driver provides an address. If they want to make this self-driving car "strong" AI, as opposed to "weak," the engineers would also need to ensure?

- A. That the AI has full human cognitive abilities that can independently decide where to take the driver.
- B. That they have obtained appropriate intellectual property (IP) licenses to use data for training the AI.
- C. That the AI has strong cybersecurity to prevent malicious actors from taking control of the car.
- D. That the AI can differentiate among ethnic backgrounds of pedestrians.

Answer: A

Explanation:

Strong AI, also known as artificial general intelligence (AGI), refers to AI that possesses the ability to understand, learn, and apply intelligence across a broad range of tasks, similar to human cognitive abilities. For the self-driving car to be classified as "strong" AI, it would need to possess full human cognitive abilities to make independent decisions beyond pre-programmed instructions. Reference: AIGP BODY OF KNOWLEDGE and AI classifications.

NEW QUESTION 17

- (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 22

- (Topic 2)

All of the following types of testing can help evaluate the performance of a responsible AI system EXCEPT?

- A. Risk probability/severity.
- B. Adversarial robustness.
- C. Statistical sampling.
- D. Decision analysis.

Answer: A

Explanation:

Risk probability/severity testing is not typically used to evaluate the performance of an AI system. While important for risk management, it does not directly assess an AI system's operational performance. Adversarial robustness, statistical sampling, and decision analysis are all methods that can help evaluate the performance of a responsible AI system by testing its resilience, accuracy, and decision-making processes under various conditions. Reference: AIGP Body of Knowledge on AI Performance Evaluation and Testing.

NEW QUESTION 26

- (Topic 2)

You are a privacy program manager at a large e-commerce company that uses an AI tool to deliver personalized product recommendations based on visitors' personal information that has been collected from the company website, the chatbot and public data the company has scraped from social media.

A user submits a data access request under an applicable U.S. state privacy law, specifically seeking a copy of their personal data, including information used to create their profile for product recommendations.

What is the most challenging aspect of managing this request?

- A. Some of the visitor's data is synthetic data that the company does not have to provide to the data subject.
- B. The data subject's data is structured data that can be searched, compiled and reviewed only by an automated tool.
- C. The data subject is not entitled to receive a copy of their data because some of it was scraped from public sources.
- D. Some of the data subject's data is unstructured data and you cannot untangle it from the other data, including information about other individuals.

Answer: D

Explanation:

The most challenging aspect of managing a data access request in this scenario is dealing with unstructured data that cannot be easily disentangled from other data, including information about other individuals. Unstructured data, such as free-text inputs or social media posts, often lacks a clear structure and may be intermingled with data from multiple individuals, making it difficult to isolate the specific data related to the requester. This complexity poses significant challenges in complying with data access requests under privacy laws. Reference: AIGP Body of Knowledge on Data Subject Rights and Data Management.

NEW QUESTION 30

- (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 34

- (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 37

- (Topic 2)

You are the chief privacy officer of a medical research company that would like to collect and use sensitive data about cancer patients, such as their names, addresses, race and ethnic origin, medical histories, insurance claims, pharmaceutical prescriptions, eating and drinking habits and physical activity. The company will use this sensitive data to build an AI algorithm that will spot common attributes that will help predict if seemingly healthy people are more likely to get cancer. However, the company is unable to obtain consent from enough patients to sufficiently collect the minimum data to train its model. Which of the following solutions would most efficiently balance privacy concerns with the lack of available data during the testing phase?

- A. Deploy the current model and recalibrate it over time with more data.
- B. Extend the model to multi-modal ingestion with text and images.
- C. Utilize synthetic data to offset the lack of patient data.
- D. Refocus the algorithm to patients without cancer.

Answer: C

Explanation:

Utilizing synthetic data to offset the lack of patient data is an efficient solution that balances privacy concerns with the need for sufficient data to train the model. Synthetic data can be generated to simulate real patient data while avoiding the privacy issues associated with using actual patient data. This approach allows for the development and testing of the AI algorithm without compromising patient privacy, and it can be refined with real data as it becomes available. Reference: AIGP Body of Knowledge on Data Privacy and AI Model Training.

NEW QUESTION 39

- (Topic 2)

A company plans on procuring a tool from an AI provider for its employees to use for certain business purposes. Which contractual provision would best protect the company's intellectual property in the tool, including training and testing data?

- A. The provider will give privacy notice to individuals before using their personal data to train or test the tool.
- B. The provider will defend and indemnify the company against infringement claims.
- C. The provider will obtain and maintain insurance to cover potential claims.
- D. The provider will warrant that the tool will work as intended.

Answer: B

Explanation:

To protect the company's intellectual property, the most pertinent contractual provision is ensuring that the AI provider will defend and indemnify the company against infringement claims. This clause means the provider will take responsibility for any intellectual property disputes that arise, thereby safeguarding the company from potential legal and financial repercussions related to the use of the tool. Other options, while beneficial, do not directly address the protection of intellectual property. This concept is detailed in the contractual best practices section of the IAPP AIGP Body of Knowledge.

NEW QUESTION 40

- (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 45

- (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.

The best human oversight mechanism for the police department to implement is that a police officer should?

- A. Explain to the accused how the AI system works.

- B. Confirm the AI recommendation prior to sentencing.
- C. Ensure an accused is given notice that the AI system was used.
- D. Consider the AI recommendation as part of the criminal investigation.

Answer: D

Explanation:

The best human oversight mechanism for the police department to implement is for a police officer to consider the AI recommendation as part of the criminal investigation. This ensures that the AI system's output is used as a tool to aid human decision-making rather than replace it. The police officer should integrate the AI's insights with other evidence and contextual information to make informed decisions, maintaining a balance between technological aid and human judgment. Reference: AIGP Body of Knowledge on AI Integration and Human Oversight.

NEW QUESTION 47

- (Topic 2)

Which of the following deployments of generative AI best respects intellectual property rights?

- A. The system produces content that is modified to closely resemble copyrighted work.
- B. The system categorizes and applies filters to content based on licensing terms.
- C. The system provides attribution to creators of publicly available information.
- D. The system produces content that includes trademarks and copyrights.

Answer: B

Explanation:

Respecting intellectual property rights means adhering to licensing terms and ensuring that generated content complies with these terms. A system that categorizes and applies filters based on licensing terms ensures that content is used legally and ethically, respecting the rights of content creators. While providing attribution is important, categorization and application of filters based on licensing terms are more directly tied to compliance with intellectual property laws. This principle is elaborated in the IAPP AIGP Body of Knowledge sections on intellectual property and compliance.

NEW QUESTION 49

.....

Thank You for Trying Our Product

We offer two products:

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

2nd - Questions 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 First Try
- * AIGP Practice Test Questions in Multiple Choice Formats and Updates for 1 Year

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