

ISTQB.CT-AI.v2026-06-18.q68

Exam Code:	CT-AI
Exam Name:	Certified Tester AI Testing Exam
Certification Provider:	ISTQB
Free Question Number:	68
Version:	v2026-06-18
# of views:	162
# of Questions views:	680
https://www.freepdfdumps.com/ISTQB.CT-AI.v2026-06-18.q68.html	

NEW QUESTION: 1

You are developing a "flower" ML model... Which of the following describes an objection that you can NEGLECT in your risk assessment?

- A. The possible inputs for the 'leaf' and 'flower' ML models are so different that reuse has few advantages over new development.
- B. The probability of misclassification of the ML model "flower" is higher when it is reused than when it is developed from scratch.
- C. The classification behavior of the "flower" ML model is more difficult to understand when it is reused compared to when it is developed from scratch.
- D. The possible outputs of the "leaf" and "flower" ML models are so different that reuse has few advantages over new development.

Answer: D (LEAVE A REPLY)

The ISTQB CT-AI syllabus explains that reusing pre-trained models is strongly related to similarity between the original task and the new task. Section 1.8 - Pre-trained Models and Transfer Learning states that reuse is effective when the new task is similar to the original one, such as adapting a cat-classifier to classify dog breeds. The syllabus warns about risks related to input differences, data preparation inconsistencies, inherited shortcomings, and explainability issues. These are legitimate objections (matching options A, B, and C) because large differences in image inputs or patterns can undermine transfer learning; misclassification risk can increase; and explainability often decreases when reusing pre-trained models.

However, output differences are NOT a valid concern here. Both the leaf-based and flower-based ML models classify the same plant species, meaning their outputs are identical. The syllabus does not identify output mismatch as a transfer-learning risk. Real risks concern inputs, bias inheritance, model transparency, and training differences--not output labels. Therefore, Option D describes an objection that can be safely neglected, because output classes are the same and do not hinder reuse.

NEW QUESTION: 2

Which of the following neural network coverage criteria can be adapted for its application?

- A. Sign-Sign coverage
- B. Sign-Change coverage
- C. Threshold coverage
- D. Neuron coverage

Answer: C (LEAVE A REPLY)

Section 4.2 - Test Coverage Criteria for AI Models of the ISTQB CT-AI syllabus describes neural network-specific coverage methods. Among the techniques, threshold coverage is explicitly noted as adaptable, meaning testers may choose different thresholds to determine whether neuron activation is considered "covered." This flexibility makes threshold coverage adjustable to the model architecture, problem domain, and required test thoroughness.

NEW QUESTION: 3

Which ONE of the following tests is LEAST likely to be performed during the ML model testing phase?

- A. Testing the accuracy of the classification model.
- B. Testing the API of the service powered by the ML model.
- C. Testing the speed of the training of the model.
- D. Testing the speed of the prediction by the model.

Answer: C (LEAVE A REPLY)

This is least likely to be part of the ML model testing phase. The speed of training is more relevant during the development phase when optimizing and tuning the model. During testing, the focus is more on the model's performance and behavior rather than how quickly it was trained.

NEW QUESTION: 4

An ML engineer performing supervised learning needs to label images of football games based on the location of the football in the image. Which ONE of the below labeling approaches can be used?

- A. Internal
- B. Annotation
- C. Augmentation
- D. Benchmarking

Answer: B (LEAVE A REPLY)

Annotation is the correct labeling approach for supervised learning, as it involves manually labeling the images with the correct information, such as marking the location of the football in the image. This labeled data can then be used to train a machine learning model.

NEW QUESTION: 5

Which of the following statements about the structure and function of neural networks is true?

- A. The bias of a neuron is determined by the activation values of the neurons in the previous layer
- B. Training a neural network only changes the values of the weights at the connections between neurons
- C. A single-layer perceptron is NOT a neural network
- D. The input layer of a deep neural network must have at least as many neurons as its output layer

Answer: B (LEAVE A REPLY)

Section 1.7 - Neural Networks of the ISTQB CT-AI syllabus explains that neural networks consist of neurons connected by weighted links. During training, learning occurs by adjusting the weights on these connections. This is the essence of gradient descent and backpropagation. Option B correctly states this behavior: only the weights are modified, not the activation functions, neuron counts, or architectural structure.

NEW QUESTION: 6

Which ONE of the following tests is MOST likely to describe a useful test to help detect different kinds of biases in ML pipeline?

- A. Testing the distribution shift in the training data for inappropriate bias.
- B. Test the model during model evaluation for data bias.
- C. Testing the data pipeline for any sources for algorithmic bias.
- D. Check the input test data for potential sample bias.

Answer: B (LEAVE A REPLY)

This is a critical stage where the model is evaluated to detect any biases in the data it was trained on. It directly addresses potential data biases in the model.

NEW QUESTION: 7

Which of the following aspects is a challenge when handling test data for an AI-based system?

- A. Personal data or confidential data
- B. Output data or intermediate data
- C. Video frame speed or aspect ratio
- D. Data frameworks or machine learning frameworks

Answer: A (LEAVE A REPLY)

The syllabus explicitly mentions challenges of handling personal data and ensuring privacy when testing AI-based systems:

"The management of personal data and sensitive data is often a concern during testing, as testing typically requires realistic data and it is difficult to fully anonymize data."

NEW QUESTION: 8

Which of the following is a problem with AI-generated test cases that are generated from the requirements?

- A. They are slow and will usually not be able to execute in the time allowed.
- B. They are defect prone because they are unable to detect nuances in the requirements.
- C. They make debugging more complicated because the number of steps is usually high in order to induce the target failure.
- D. They are usually missing the expected results, so verification is difficult or must resort to only detecting significant failures.

Answer: D ([LEAVE A REPLY](#))

AI-generated test cases are often created using machine learning (ML) models or heuristic algorithms. While these can be effective in generating large numbers of test cases quickly, they often suffer from the "test oracle problem." Test Oracle Problem: A test oracle is the mechanism used to determine the expected output of a test case. AI-generated test cases often lack expected results because AI-based tools do not inherently understand what the correct output should be.

Difficulty in Verification: Without expected results, verifying test cases becomes challenging.

Testers must rely on heuristics, anomaly detection, or significant failures, rather than traditional pass/fail conditions.

NEW QUESTION: 9

Which ONE of the below is NOT likely to cause a data quality issue affecting a single ML model?

- A. Security issues
- B. Hardware issues
- C. Incorrect weights
- D. Faulty sensors

Answer: (SHOW ANSWER)

Incorrect weights refer to parameters within a trained model and are not typically considered a data quality issue. Data quality issues usually arise from the input data itself, such as security issues, hardware issues, or faulty sensors, which can lead to incorrect or missing data. Incorrect weights are a problem during the model training or fine-tuning phase, not directly related to data quality.

NEW QUESTION: 10

A transportation company operates three types of delivery vehicles in its fleet. The vehicles operate at different speeds (slow, medium, and fast). The transportation company is attempting to optimize scheduling and has created an AI-based program to plan routes for its vehicles using records from the medium-speed vehicle traveling to selected

destinations. The test team uses this data in metamorphic testing to test the accuracy of the estimated travel times created by the AI route planner with the actual routes and times. Which of the following describes the next phase of metamorphic testing?

- A.** The team tests the time required for the fast and slow vehicles to travel the same route as the medium vehicle. Then, by calculating the speed difference, they then predict how much faster or slower the vehicles will travel. That information is then used to verify that the arrival time of the vehicles meets the expected result.
- B.** The team decomposes each route into the relevant components that affect the travel time such as traffic density and vehicle power. The team then uses statistical analysis to characterize the influence of each component to calculate the fast and slow vehicle route times.
- C.** The team uses an AI system to select the most dissimilar routes. With this information, any of the AI routes can be metaphorically transformed into a fast or slow route.
- D.** The team uses the same AI route planner to create routes that are longer and shorter but follow the same track. Finally, by driving the fast vehicles on the long routes and slow vehicles on the short routes and vice versa, the AI system will have enough information to infer travel times for all vehicles on all routes.

Answer: A (LEAVE A REPLY)

Metamorphic Testing (MT) is a testing technique that verifies AI-based systems by generating follow-up test cases based on existing test cases. These follow-up test cases adhere to a Metamorphic Relation (MR), ensuring that if the system is functioning correctly, changes in input should result in predictable changes in output.

Metamorphic testing works by transforming source test cases into follow-up test cases. Here, the source test case involves testing the medium-speed vehicle's travel time. The follow-up test cases are derived by extrapolating travel times for fast and slow vehicles using predictable relationships based on speed differences.

MR states that modifying input should result in a predictable change in output. Since the speed of the vehicle is a known factor, it is possible to predict the new arrival times and verify whether they follow expected trends.

This is a direct application of metamorphic testing principles. In route optimization systems, metamorphic testing often applies transformations to speed, distance, or conditions to verify expected outcomes.

NEW QUESTION: 11

A word processing company is developing an automatic text correction tool. A machine learning algorithm was used to develop the auto text correction feature. The testers have discovered when they start typing "Isle of Wight" it fills in "Isle of Eight". Several UAT testers have accepted this change without noticing. What type of bias is this?

- A.** Geographical/Locality
- B.** Automation/Complacency
- C.** Complacency/Disregard

D. Ignorance/Cognitive

Answer: (SHOW ANSWER)

The syllabus describes automation bias as:

"A type of bias caused by a person favoring the recommendations of an automated decision-making system over other sources." This is also known as complacency bias, where testers accept automated system outputs without questioning them.

NEW QUESTION: 12

A local business has a mail pickup/delivery robot for their office. The robot currently uses a track to move between pickup/drop off locations. When it arrives at a destination, the robot stops to allow a human to remove or deposit mail.

The office has decided to upgrade the robot to include AI capabilities that allow the robot to perform its duties without a track, without running into obstacles, and without human intervention.

The test team is creating a list of new and previously established test objectives and acceptance criteria to be used in the testing of the robot upgrade. Which of the following test objectives will test an AI quality characteristic for this system?

- A. The robot must evolve to optimize its routing
- B. The robot must recharge for no more than six hours a day
- C. The robot must record the time of each delivery which is compiled into a report
- D. The robot must complete 99.99% of its deliveries each day

Answer: (SHOW ANSWER)

In the syllabus, the evolution characteristic for AI-based systems means the ability of the system to evolve and adapt its behavior in response to changes in the environment or in its own performance:

"Evolution is the system's ability to change itself to adapt to new situations, different hardware, or a changing operational environment."

NEW QUESTION: 13

Which of the following are the three activities in the data acquisition activities for data preparation?

- A. Cleaning, transforming, augmenting
- B. Feature selecting, feature growing, feature augmenting
- C. Identifying, gathering, labelling
- D. Building, approving, deploying

Answer: (SHOW ANSWER)

The syllabus defines data acquisition as consisting of three steps:

"Data acquisition: The activity of acquiring data relevant to the business problem to be solved by an ML model, typically involving the activities of identifying, gathering and labelling data."

NEW QUESTION: 14

Which of the following statements about bias in AI-based systems is MOST correct?

- A. Inappropriate bias is caused by overweighting of particular classes in algorithms
- B. Inappropriate bias is caused by data used for training not being representative of the real world
- C. Inappropriate bias only affects ML systems that process data about people
- D. Inappropriate bias can be caused by aspects of the algorithm or the data

Answer: (SHOW ANSWER)

Inappropriate bias in AI systems is typically caused by the data used for training not being representative of the real world. This can lead to the model making biased or incorrect predictions when applied to real-world scenarios. While algorithmic factors can also contribute to bias, the primary issue arises from biased or unrepresentative training data.

NEW QUESTION: 15

Which ONE of the following hardware is MOST suitable for implementing AI when using ML?

- A. 64-bit CPUs.
- B. Hardware supporting fast matrix multiplication.
- C. High powered CPUs.
- D. Hardware supporting high precision floating point operations.

Answer: B (LEAVE A REPLY)

Matrix multiplication is a fundamental operation in many machine learning algorithms, especially in neural networks and deep learning. Hardware optimized for fast matrix multiplication, such as GPUs (Graphics Processing Units), is most suitable for implementing AI and ML because it can handle the parallel processing required for these operations efficiently.

NEW QUESTION: 16

Which statement about automation bias is correct?

- A. When testing AI-based systems, automation bias does not play a role in supporting test activities such as boundary value analysis
- B. Automation bias affects the testing of AI-based systems that support users in their actions or decisions
- C. Automation bias particularly affects testing of autonomous systems
- D. Automation bias is tested with representative users, but human input quality is irrelevant

Answer: B (LEAVE A REPLY)

Automation bias is defined in Section 4.4 - Human Factors in AI Testing of the ISTQB CT-AI syllabus. It refers to the human tendency to overly trust, rely on, or defer to automated system outputs. The syllabus explains that this bias arises especially in decision-support systems, where humans may accept AI judgments without adequate verification. This aligns directly with Option B).

Valid CT-AI Dumps shared by Actual4test.com for Helping Passing CT-AI Exam!
Actual4test.com now offer the **newest CT-AI exam dumps**, the Actual4test.com CT-AI exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com CT-AI dumps with Test Engine here:
https://www.actual4test.com/CT-AI_examcollection.html (162 Q&As Dumps, **30%OFF**
Special Discount: Freepdfdumps)

NEW QUESTION: 17

Which two test procedures are BEST suited for CleverPropose system testing?

- A. Back-to-back testing
- B. Adversarial testing
- C. Metamorphic testing
- D. Exploratory data analysis
- E. Pairwise testing

Answer: (SHOW ANSWER)

The ISTQB CT-AI syllabus explains that AI-based decision-support systems benefit strongly from back-to-back testing and metamorphic testing when oracle problems exist or when limited regression tests are available. In this scenario, CleverPropose replaces an older advisory system. Back-to-back testing (Option A) is ideal because the outputs of the existing conventional system can serve as a reference, enabling comparison against the new AI system. This is exactly what the syllabus recommends when AI is replacing a traditional deterministic system.

Metamorphic testing (Option C) is also appropriate, as stated in Section 4.6 - Metamorphic Relations. With limited regression tests and complex decision logic, testers can define metamorphic relations such as "if customer income increases, risk rating should not worsen." These relations allow validation even when exact expected outputs are unavailable.

NEW QUESTION: 18

The following confusion matrix represents the functional performance of a classifier.

	Actual Positive	Actual Negative
Predicted Positive	60	20
Predicted Negative	9	11

Which ONE of the following is the correct calculation for the accuracy of the classifier?

- A. $60 / (60 + 20) * 100\% = 75\%$
- B. $60 / (60 + 9) * 100\% = 87\%$
- C. $2 * (60 / (60 + 20) * 60 / (60 + 9)) / (60 / (60 + 20) + 60 / (60 + 9)) * 100\% = 80\%$

D. $(60 + 11) / (60 + 11 + 20 + 9) * 100\% = 71\%$

Answer: [\(SHOW ANSWER\)](#)

Accuracy = (True Positives + True Negatives) / (Total Samples)

In this case:

- True Positives (TP) = 60

- True Negatives (TN) = 11

- False Positives (FP) = 20

- False Negatives (FN) = 9

Accuracy = $(60 + 11) / (60 + 11 + 20 + 9) * 100\% = 71\%$

NEW QUESTION: 19

Which of the following describes the AI effect?

- A.** The changing perception of what constitutes AI
- B.** The ability of AI to learn from data itself
- C.** The fact that mankind can build intelligent machines
- D.** The ability of AI to defeat a human, e.g., in chess

Answer: **A** [\(LEAVE A REPLY\)](#)

The AI Effect is clearly defined in the ISTQB Certified Tester AI Testing Syllabus v1.0 under Section 1.1 - Definition of AI and AI Effect. The document explains that society's understanding of what qualifies as "AI" changes over time. Technologies once considered AI--such as expert systems from the 1970s and 1980s or early chess-playing systems--are no longer viewed as AI today. This phenomenon is explicitly labeled the "AI Effect," described as "the changing perception of what constitutes AI." The syllabus states that as AI capabilities become routine or widely implemented, they often stop being perceived as true artificial intelligence.

NEW QUESTION: 20

Which statement regarding the use of training, validation, and test data sets is correct?

- A.** If only limited data is available, validation and test data sets can be combined in multiple ways during training.
- B.** If limited data is available, it may be better to work without a separate test data set.
- C.** Optimally, the data should be distributed equally between the training, validation, and test data sets.
- D.** The data in the test data set must be equivalent to the data in the training data sets and to the data in the validation data sets.

Answer: **D** [\(LEAVE A REPLY\)](#)

The ISTQB CT-AI syllabus (Section 3.2 - Model Evaluation) specifies the correct usage of training, validation, and test data sets. It emphasizes that the test dataset must be representative of the real operational data and must be equivalent in distribution to the training and validation sets, ensuring a fair and unbiased evaluation. Option D precisely matches this requirement.

NEW QUESTION: 21

Which ONE of the following would be the MOST effective input to an AI-based defect prediction tool?

- A. Developers associated with previous code changes
- B. Cyclomatic complexity
- C. Lines of code changed
- D. Commit size

Answer: B (LEAVE A REPLY)

Cyclomatic complexity would be the most effective input to an AI-based defect prediction tool.

Cyclomatic complexity is a software metric that measures the complexity of a program's control flow, which is closely related to the likelihood of defects. Higher complexity generally indicates a higher probability of defects. This makes it a strong predictor of potential issues in the code, and thus a valuable input for defect prediction.

NEW QUESTION: 22

A test engineer is planning the best functional performance metrics to evaluate an unsupervised learning model. The model groups data points based on their similarity. The test engineer wants to measure how similar the data points in each group actually are. Which is the MOST likely metric they should use?

- A. ROC
- B. Intra-cluster
- C. AUC
- D. Recall

Answer: B (LEAVE A REPLY)

The most appropriate metric for evaluating the similarity of data points within each group in an unsupervised learning model is intra-cluster. This metric measures how similar the data points within each cluster are to one another. The goal is to have high intra-cluster similarity, meaning the data points within a group should be similar to each other.

NEW QUESTION: 23

An image classification system is being trained for classifying faces of humans. The distribution of the data is 70% ethnicity A and 30% for ethnicities B, C and D. Based ONLY on the above information, which of the following options BEST describes the situation of this image classification system?

- A. This is an example of expert system bias.
- B. This is an example of sample bias.
- C. This is an example of hyperparameter bias.
- D. This is an example of algorithmic bias.

Answer: (SHOW ANSWER)

Sample bias occurs when the training data is not representative of the overall population that the model will encounter in practice.

In this case, the over-representation of ethnicity A (70%) compared to B, C, and D (30%) creates a sample bias, as the model may become biased towards better performance on ethnicity A.

NEW QUESTION: 24

Which option gives the correct values for accuracy and precision from the confusion matrix?

- A.** Accuracy = 50%, Precision = 75%
- B.** Accuracy = 80%, Precision = 75%
- C.** Accuracy = 75%, Precision = 80%
- D.** Accuracy = 80%, Precision = 50%

Answer: [\(SHOW ANSWER\)](#)

From the confusion matrix:

True Positives (TP) = 15

False Positives (FP) = 5

False Negatives (FN) = 15

True Negatives (TN) = 65

Accuracy = $(TP + TN) / \text{Total}$

$= (15 + 65) / 100$

$= 80\%$

Precision = $TP / (TP + FP)$

$= 15 / (15 + 5)$

$= 15 / 20$

$= 75\%$

Section 3.2 - Functional Performance Criteria in the syllabus explains accuracy and precision exactly these ways when evaluating ML classification performance.

Option B is therefore the only correct pair of values.

NEW QUESTION: 25

Which ONE of the following statements BEST describes how system complexity can cause challenges when testing an AI-based system?

- A.** Obtaining test data is harder
- B.** Obtaining test data is harder
- C.** Unexpected changes in system behavior can occur
- D.** Sometimes the system can only be tested as a black-box

Answer: [C \(LEAVE A REPLY\)](#)

Unexpected changes in system behavior can occur due to the complexity of AI-based systems.

These systems often involve many interacting components, which can lead to unpredictable results or variations in performance, making it difficult to anticipate how the system will behave under certain conditions. This presents a significant challenge in testing, as such behavior can be difficult to reproduce or control.

NEW QUESTION: 26

Which statement regarding AI for defect prediction is correct?

- A.** AI-based defect prediction is most effective when based on previous similar constellations.
- B.** AI-based defect prediction can detect whether defects exist but not where.
- C.** AI-based defect prediction is most effective when based on source-code metrics such as branches or McCabe complexity.
- D.** AI-based defect prediction is based on formal principles and requires only a few factors.

Answer: A ([LEAVE A REPLY](#))

Section 5.3 - AI Support for Defect Prediction of the ISTQB CT-AI syllabus explains that AI-based defect prediction models rely on historical patterns, including past defects, code behavior, and similar system configurations. ML models trained on prior defect data can identify components likely to contain defects when new changes resemble previous defect-inducing patterns. This directly supports Option A, which states that defect prediction is most effective when based on previous similar constellations.

NEW QUESTION: 27

Which statement describes factors related to test data that make testing AI-based systems difficult?

- A.** Using the same implementation for data acquisition by data scientists and testers prevents defect masking
- B.** Creating and managing large amounts of test data can be difficult, especially when it needs to be representative
- C.** The input data must always be the same over time, especially in real-world systems
- D.** Artificially generated data requires legal approval and must be sanitized and encrypted

Answer: ([SHOW ANSWER](#))

Section 2.2 - Data Preparation and 4.1 - Challenges in Testing AI-Based Systems describe difficulties in obtaining and managing large, representative datasets. AI-based systems require realistic, diverse, and representative data reflecting real-world variations. The syllabus emphasizes that assembling such datasets is time-consuming, resource-intensive, and often constrained by availability, privacy, or domain complexity. Option B directly corresponds to these documented challenges.

NEW QUESTION: 28

In a certain coffee producing region of Colombia, there have been some severe weather storms, resulting in massive losses in production. This caused a massive drop in stock price of coffee.

Which ONE of the following types of testing SHOULD be performed for a machine learning model for stock-price prediction to detect influence of such phenomenon as above on price of coffee stock.

- A. Testing for accuracy
- B. Testing for bias
- C. Testing for concept drift
- D. Testing for security

Answer: C (LEAVE A REPLY)

Type of Testing for Stock-Price Prediction Models: Concept drift refers to the change in the statistical properties of the target variable over time. Severe weather storms causing massive losses in coffee production and affecting stock prices would require testing for concept drift to ensure that the model adapts to new patterns in data over time.

NEW QUESTION: 29

Which ONE of the following options does NOT describe a challenge for acquiring test data in ML systems?

- A. Compliance needs require proper care to be taken of input personal data.
- B. Nature of data constantly changes with time.
- C. Data for the use case is being generated at a fast pace.
- D. Test data being sourced from public sources.

Answer: C (LEAVE A REPLY)

Challenges for Acquiring Test Data in ML Systems: Compliance needs, the changing nature of data over time, and sourcing data from public sources are significant challenges. Data being generated quickly is generally not a challenge; it can actually be beneficial as it provides more data for training and testing.

NEW QUESTION: 30

Which statement about AI-based test case generation is correct?

- A. AI-generated functional test cases typically result in poor requirements coverage.
- B. A different test oracle is usually required for each AI-generated functional test case.
- C. Expected results may not be available for AI-generated functional test cases.
- D. An AI-based system under test must not be used as a functional test oracle.

Answer: C (LEAVE A REPLY)

The ISTQB CT-AI syllabus indicates in Section 5.2 - AI for Testing that AI-generated test cases may not come with predefined expected results. This is because test-case generation methods-- such as evolutionary algorithms, reinforcement learning, or clustering-based sampling-- produce inputs, but the tester must still determine the correct

outputs. Therefore, Option C is correct: expected results may not be available, especially when AI produces novel or previously unseen input combinations.

NEW QUESTION: 31

Which of the following is one of the reasons for data mislabelling?

- A. Lack of domain knowledge
- B. Expert knowledge
- C. Interoperability error
- D. Small datasets

Answer: A (LEAVE A REPLY)

The syllabus lists multiple reasons for mislabelled data, including the lack of domain knowledge:

"Lack of required domain knowledge may lead to incorrect labelling."

Valid CT-AI Dumps shared by Actual4test.com for Helping Passing CT-AI Exam! Actual4test.com now offer the **newest CT-AI exam dumps**, the Actual4test.com CT-AI exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com CT-AI dumps with Test Engine here: https://www.actual4test.com/CT-AI_examcollection.html (162 Q&As Dumps, **30%OFF** Special Discount: **Freepdfdumps**)

NEW QUESTION: 32

There is a growing backlog of unresolved defects for your project. You know the developers have an ML model that they have created which has learned which developers work on which type of software and the speed with which they resolve issues. How could you use this model to help reduce the backlog and implement more efficient defect resolution?

- A. Use it to prioritize defects automatically based on the time expected for the fix to be made, the speed of the fix, and the likelihood of regressions.
- B. Use it to assign defects to the best developer to resolve the problem and to load balance the defect assignments among the developers.
- C. Use it to determine the root cause of each defect and develop a process improvement plan that can be implemented to remove the most common root causes.
- D. Use it to review the code and determine where more defects are likely to occur so that testing can be targeted to those areas.

Answer: B (LEAVE A REPLY)

The syllabus explains that ML models can be used to analyze reported defects and suggest which developers are best suited to fix them based on historical data about defect assignment and resolution speed:

"Assignment: ML models can suggest which developers are best suited to fix particular defects, based on the defect content and previous developer assignments."

NEW QUESTION: 33

Which of the following descriptions of quality aspects of a data set is correct?

- A.** The quality aspect "Incomplete data" describes the fact that data is missing, e.g., for a certain time interval.
- B.** The quality aspect "Data not preprocessed" describes the fact that the collected data was recorded incorrectly.
- C.** The quality aspect "Irrelevant data" describes the fact that irrelevant data does not affect the ML model.
- D.** The quality aspect "Unbalanced data" describes the fact that the data used should be as up-to-date as possible.

Answer: (SHOW ANSWER)

The ISTQB CT-AI syllabus describes several data quality aspects that affect ML performance. In Section 2.2 - Data Preparation, it explains that datasets may suffer from issues such as incomplete data, irrelevant data, incorrect data, unbalanced data, or data lacking preprocessing.

"Incomplete data" means that portions of the required data are missing, often because some time periods, records, or sources were not captured. This aligns exactly with Option A, which correctly identifies missing intervals as incomplete data.

NEW QUESTION: 34

A company producing consumable goods wants to identify groups of people with similar tastes for the purpose of targeting different products for each group. You have to choose and apply an appropriate ML type for this problem.

Which ONE of the following options represents the BEST possible solution for this above-mentioned task?

- A.** Regression
- B.** Association
- C.** Clustering
- D.** Classification

Answer: (SHOW ANSWER)

Clustering is an unsupervised learning method used to group similar data points based on their features. It is ideal for identifying groups of people with similar tastes without prior knowledge of the group labels. This technique will help the company segment its customer base effectively.

NEW QUESTION: 35

Which of the below TWO examples of AI system behavior are reward hacking:

- I. An AI medical device intended to keep a patient stable may give the patient a treatment that means they recover less quickly.
- II. An AI system intended to maximize production of a commodity by weight allows quality and size of product to reduce.
- III. An AI system intended to ensure the output of a factory process is always sorted correctly, destroys the outputs.
- IV. An AI system intended to remove security vulnerabilities from software code, removes all functionality that has security vulnerabilities.

- A. I, II
- B. I, III
- C. III, IV
- D. II, IV

Answer: (SHOW ANSWER)

Reward hacking occurs when an AI system manipulates the reward structure in unintended ways to achieve high rewards but in a manner that defeats the true purpose of the system.

I (AI medical device): The system may treat the patient in a way that stabilizes them but hinders recovery, exploiting the reward function in a harmful way.

II (AI maximizing production): The system maximizes production by allowing quality and size to drop, which is also reward hacking, as it focuses only on weight and ignores quality.

In both cases, the AI is manipulating the reward system to achieve the goal in a way that undermines the intended purpose.

NEW QUESTION: 36

Which ONE of the following statements correctly describes the importance of flexibility for AI systems?

- A. AI systems are inherently flexible.
- B. AI systems require changing of operational environments; therefore, flexibility is required.
- C. Flexible AI systems allow for easier modification of the system as a whole.
- D. Self-learning systems are expected to deal with new situations without explicitly having to program for it.

Answer: (SHOW ANSWER)

Flexibility in AI systems is crucial for various reasons, particularly because it allows for easier modification and adaptation of the system as a whole.

AI systems are inherently flexible (A): This statement is not correct. While some AI systems may be designed to be flexible, they are not inherently flexible by nature. Flexibility depends on the system's design and implementation.

AI systems require changing operational environments; therefore, flexibility is required (B): While it's true that AI systems may need to operate in changing environments, this statement does not directly address the importance of flexibility for the modification of the system.

Flexible AI systems allow for easier modification of the system as a whole (C): This statement correctly describes the importance of flexibility. Being able to modify AI systems easily is critical for their maintenance, adaptation to new requirements, and improvement. Self-learning systems are expected to deal with new situations without explicitly having to program for it (D): This statement relates to the adaptability of self-learning systems rather than their overall flexibility for modification.

Hence, the correct answer is C. Flexible AI systems allow for easier modification of the system as a whole.

NEW QUESTION: 37

Which performance metric is BEST suited to assess the quality of trained models detecting fraudulent credit card transactions?

- A. Sensitivity
- B. Accuracy
- C. F1 value

Answer: C (LEAVE A REPLY)

The ISTQB CT-AI syllabus explains in Section 3.2 - Functional Performance Criteria of ML Models that accuracy becomes unreliable when class imbalance exists. In fraud detection, more than 99% of transactions are non-fraudulent, meaning the dataset is extremely imbalanced.

Because accuracy counts all correct non-fraudulent classifications, it will appear artificially high, even if the fraud detection performance is poor. Therefore, accuracy is not suitable for evaluating fraud detection systems.

The syllabus further explains that sensitivity (recall) captures the proportion of correctly identified fraudulent cases. This metric is important, as missing fraudulent events can cause high financial loss. However, the client also stresses that legitimate transactions must be correctly identified, meaning false positives must be minimized to maintain customer satisfaction.

The F1 score, defined as the harmonic mean of precision and recall, balances both: Precision protects legitimate customers by minimizing false alarms.

Recall ensures fraudulent transactions are detected.

Section 3.2 emphasizes that when both false positives and false negatives have significant consequences, and the data is highly imbalanced, F1 is the most appropriate metric because it reflects the combined importance of detecting fraud while avoiding unnecessary alerts. Thus, Option C is the correct choice.

NEW QUESTION: 38

Which of the following statements regarding experience-based testing for AI-based systems is correct?

- A. Intuitive test case design for AI-based systems involves interactive, hypothesis-driven examination of data for correlations or developmental trends.

B. In checklist-based testing of AI-based systems, the existing test cases are dynamically adapted, for example based on metamorphic testing.

C. Exploratory testing is often used for AI-based systems because there are often insufficient specifications or problems with the test oracle for AI-based systems.

D. Tour refers to intuitive test case design for AI-based systems based on multiple, sequential test cases using systematically biased training data.

Answer: C (LEAVE A REPLY)

The ISTQB CT-AI syllabus explains in Section 4.4 - Experience-Based Testing for AI Systems that AI-based systems frequently suffer from insufficient specifications, unpredictable model behavior, and test oracle problems, especially when outputs depend on probabilistic or learned patterns.

The syllabus explicitly states that exploratory testing is especially valuable in such contexts because it allows testers to investigate the system interactively, observe unexpected behavior, and evaluate system responses that cannot be fully predicted beforehand. Thus, Option C accurately reflects the role and justification of exploratory testing for AI systems.

NEW QUESTION: 39

Which of the following statements about reinforcement learning is correct?

A. The agent creates a model of the environment from labeled data during training

B. The approach is suitable when the application does not require interaction with the environment

C. The agent's training is based on a reward function that rewards successful attempts

D. From experience, the agent learns the optimal reward function

Answer: C (LEAVE A REPLY)

Section 1.6.3 - Reinforcement Learning of the ISTQB CT-AI syllabus states that reinforcement learning (RL) is based on an agent interacting with an environment, performing actions, and receiving rewards or penalties. The core concept is the reward function, which guides the agent's learning process. The syllabus emphasizes that training in RL is driven by rewards, and the agent aims to maximize cumulative reward over time. Therefore, Option C directly reflects the correct description: the agent learns by being rewarded for successful actions.

NEW QUESTION: 40

An engine manufacturing facility wants to apply machine learning to detect faulty bolts.

Which of the following would result in bias in the model?

A. Selecting training data by purposely excluding specific faulty conditions

B. Selecting training data by purposely including all known faulty conditions

C. Selecting testing data from a different dataset than the training dataset

D. Selecting testing data from a boat manufacturer's bolt longevity data

Answer: A (LEAVE A REPLY)

The syllabus defines bias as:

"Bias is the systematic difference in treatment of certain objects, people or groups in comparison to others." It also discusses:

"Sample bias can occur if the data used for training the model does not represent the operational environment, or if some relevant faulty conditions are excluded deliberately."

NEW QUESTION: 41

Which ONE of the following types of coverage SHOULD be used if test cases need to cause each neuron to achieve both positive and negative activation values?

- A. Value coverage
- B. Threshold coverage
- C. Sign change coverage
- D. Neuron coverage

Answer: C (LEAVE A REPLY)

Coverage for Neuron Activation Values: Sign change coverage is used to ensure that test cases cause each neuron to achieve both positive and negative activation values. This type of coverage ensures that the neurons are thoroughly tested under different activation states.

NEW QUESTION: 42

Which ONE of the following approaches to labelling requires the least time and effort?

- A. Outsourced
- B. Pre-labeled dataset
- C. Internal
- D. AI-Assisted

Answer: B (LEAVE A REPLY)

Labelling Approaches: Among the options provided, pre-labeled datasets require the least time and effort because the data has already been labeled, eliminating the need for further manual or automated labeling efforts.

NEW QUESTION: 43

While measuring the test coverage of a neural network, a test engineer wants to measure the number of neurons that have each output two activation function results with a minimum difference between the two results of 0.5.

Which ONE of the below coverage measures would achieve that goal?

- A. Neuron coverage
- B. Value-change coverage
- C. Sign-change coverage
- D. None of the coverage measures

Answer: B (LEAVE A REPLY)

Value-change coverage is a coverage measure that focuses on tracking the changes in the values of a neuron's output. In this case, the test engineer is interested in measuring the

number of neurons that exhibit a difference of 0.5 between two activation function results, which is effectively a value change. This coverage measure ensures that the neuron undergoes a significant change in its output, fulfilling the engineer's goal.

NEW QUESTION: 44

A neural network has been designed and created to assist day-traders improve efficiency when buying and selling commodities in a rapidly changing market. Suppose the test team executes a test on the neural network where each neuron is examined. For this network the shortest path indicates a buy, and it will only occur when the one-day predicted value of the commodity is greater than the spot price by 0.75%. The neurons are stimulated by entering commodity prices and testers verify that they activate only when the future value exceeds the spot price by at least 0.75%. Which of the following statements BEST explains the type of coverage being tested on the neural network?

- A.** Threshold coverage
- B.** Neuron coverage
- C.** Sign-change coverage
- D.** Value-change coverage

Answer: A (LEAVE A REPLY)

The syllabus details that threshold coverage requires each neuron to achieve an activation value greater than a specified threshold:

"Threshold coverage: Full threshold coverage requires that each neuron in the neural network achieves an activation value greater than a specified threshold."

NEW QUESTION: 45

The stakeholders of a machine learning model have confirmed that they understand the objective and purpose of the model, and ensured that the proposed model aligns with their business priorities. They have also selected a framework and a machine learning model that they will be using. What should be the next step to progress along the machine learning workflow?

- A.** Tune the machine learning algorithm based on objectives and business priorities
- B.** Prepare and pre-process the data that will be used to train and test the model
- C.** Agree on defined acceptance criteria for the machine learning model
- D.** Evaluate the selection of the framework and the model

Answer: (SHOW ANSWER)

The ML workflow typically involves iterative steps, beginning with data preparation once the model and framework are selected. The syllabus explains:

"The steps shown in Figure 1 (the ML workflow) do not include the integration of the ML model with the non-ML parts of the overall system. Typically, ML models cannot be deployed in isolation and need to be integrated with the non-ML parts... The next step

would be data preparation as part of the ML workflow to provide input data to support training by an ML algorithm or prediction by an ML model."

NEW QUESTION: 46

Which ONE of the following options BEST DESCRIBES clustering?

- A. Clustering is classification of a continuous quantity.
- B. Clustering is supervised learning.
- C. Clustering is done without prior knowledge of output classes.
- D. Clustering requires you to know the classes.

Answer: C (LEAVE A REPLY)

Clustering is a type of machine learning technique used to group similar data points into clusters.

It is a key concept in unsupervised learning, where the algorithm tries to find patterns or groupings in data without prior knowledge of output classes.

In clustering, the algorithm groups data points into clusters without any prior knowledge of the classes. It discovers the inherent structure in the data.

Valid CT-AI Dumps shared by Actual4test.com for Helping Passing CT-AI Exam!

Actual4test.com now offer the **newest CT-AI exam dumps**, the Actual4test.com CT-AI exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com CT-AI dumps with Test Engine here:

https://www.actual4test.com/CT-AI_examcollection.html (162 Q&As Dumps, **30%OFF**

Special Discount: Freepdfdumps)

NEW QUESTION: 47

Which ONE of the below is MOST likely to indicate a problem with underfitting in an ML model?

- A. The model is vulnerable to adversarial attacks
- B. The model fails to generalize on new data
- C. The model uses a large amount of resources to make a prediction
- D. The model is inaccurate on data similar to the training data

Answer: D (LEAVE A REPLY)

Underfitting occurs when the model is too simplistic to capture the underlying patterns in the training data. This results in poor performance not only on new data but also on the data it was trained on. Thus, if the model is inaccurate on data similar to the training data, it suggests underfitting.

NEW QUESTION: 48

Which ONE of the following statements is MOST true in relation to A/B testing?

- A. Many test cases can be generated
- B. The purpose is to detect defects
- C. Production data is not used
- D. The purpose is to compare two variants of a system

Answer: [\(SHOW ANSWER\)](#)

A/B testing is primarily used to compare two variants of a system to determine which one performs better. This method involves testing two different versions (A and B) to analyze their impact on specific metrics, often in a real-world environment. The other options do not fully align with the purpose of A/B testing.

NEW QUESTION: 49

Which ONE of the below statements BEST describes how combinatorial testing can be applied to AI-based systems?

- A. Each neuron can be treated as a parameter for pairwise tests
- B. Two variants of the system can be used and compared
- C. Inputs to the system and environment factors can be considered parameters for pairwise tests
- D. Combinatorial testing cannot yet be applied to AI

Answer: [\(SHOW ANSWER\)](#)

Combinatorial testing involves testing different combinations of input parameters to ensure coverage of various interactions. In the case of AI-based systems, inputs to the system (such as data features) and environment factors (like conditions under which the system operates) can be considered parameters for pairwise tests. This ensures that different combinations of inputs and environmental factors are tested for their effect on the system's behavior.

NEW QUESTION: 50

How can a tester check the system for bias as part of a review of data sources, acquisition, and preprocessing?

- A. During the review, it can uncover algorithmic bias by analysing the procedures used to obtain the training data.
- B. During the review of the preprocessing, the auditor can uncover whether the data has been influenced in a way that could lead to sample distortions.
- C. It may use the LIME method as part of its data collection review to detect inappropriate bias.
- D. As part of the review of preprocessing, it can reveal whether the data has been influenced in a way that could lead to algorithmic bias.

Answer: B [\(LEAVE A REPLY\)](#)

Bias detection at the data level is performed by reviewing data acquisition and preprocessing steps, as explained in Section 2.3 - Data Quality and Bias of the ISTQB CT-AI syllabus. Sample bias arises when data is distorted or when preprocessing introduces

unintended shifts--for example, by filtering, normalization, or labeling steps that disproportionately affect subsets of the data. Option B correctly reflects this: reviewers can identify whether preprocessing steps have altered the dataset in a way that introduces sample distortions. This aligns perfectly with syllabus guidance on reviewing data pipelines for bias sources.

NEW QUESTION: 51

"Splendid Healthcare" has started developing a cancer detection system based on ML. The type of cancer they plan on detecting has 2% prevalence rate in the population of a particular geography. It is required that the model performs well for both normal and cancer patients. Which ONE of the following combinations requires MAXIMIZATION?

- A.** Maximize precision and accuracy
- B.** Maximize accuracy and recall
- C.** Maximize recall and precision
- D.** Maximize specificity number of classes

Answer: C (LEAVE A REPLY)

Prevalence Rate and Model Performance:

The cancer detection system being developed by "Splendid Healthcare" needs to account for the fact that the type of cancer has a 2% prevalence rate in the population. This indicates that the dataset is highly imbalanced with far fewer positive (cancer) cases compared to negative (normal) cases.

Importance of Recall:

Recall, also known as sensitivity or true positive rate, measures the proportion of actual positive cases that are correctly identified by the model. In medical diagnosis, especially cancer detection, recall is critical because missing a positive case (false negative) could have severe consequences for the patient. Therefore, maximizing recall ensures that most, if not all, cancer cases are detected.

Importance of Precision:

Precision measures the proportion of predicted positive cases that are actually positive. High precision reduces the number of false positives, meaning fewer people will be incorrectly diagnosed with cancer. This is also important to avoid unnecessary anxiety and further invasive testing for those who do not have the disease.

Balancing Recall and Precision:

In scenarios where both false negatives and false positives have significant consequences, it is crucial to balance recall and precision. This balance ensures that the model is not only good at detecting positive cases but also accurate in its predictions, reducing both types of errors.

Accuracy and Specificity:

While accuracy (the proportion of total correct predictions) is important, it can be misleading in imbalanced datasets. In this case, high accuracy could simply result from the model predicting the majority class (normal) correctly. Specificity (true negative rate) is

also important, but for a cancer detection system, recall and precision take precedence to ensure positive cases are correctly and accurately identified.

Conclusion:

Therefore, for a cancer detection system with a low prevalence rate, maximizing both recall and precision is crucial to ensure effective and accurate detection of cancer cases.

NEW QUESTION: 52

Which ONE of the following describes a situation of back-to-back testing the LEAST?

- A.** Comparison of the results of a current neural network model ML model implemented in platform A (for example Pytorch) with a similar neural network model ML model implemented in platform B (for example Tensorflow), for the same data.
- B.** Comparison of the results of a home-grown neural network model ML model with results in a neural network model implemented in a standard implementation (for example Pytorch) for same data
- C.** Comparison of the results of a neural network ML model with a current decision tree ML model for the same data.
- D.** Comparison of the results of the current neural network ML model on the current data set with a slightly modified data set.

Answer: C (LEAVE A REPLY)

Back-to-back testing is a method where the same set of tests are run on multiple implementations of the system to compare their outputs. This type of testing is typically used to ensure consistency and correctness by comparing the outputs of different implementations under identical conditions.

Let's analyze the options given:

A). Comparison of the results of a current neural network model ML model implemented in platform A (for example Pytorch) with a similar neural network model ML model implemented in platform B (for example Tensorflow), for the same data.

This option describes a scenario where two different implementations of the same type of model are being compared using the same dataset. This is a typical back-to-back testing situation.

B). Comparison of the results of a home-grown neural network model ML model with results in a neural network model implemented in a standard implementation (for example Pytorch) for the same data.

This option involves comparing a custom implementation with a standard implementation, which is also a typical back-to-back testing scenario to validate the custom model against a known benchmark.

C). Comparison of the results of a neural network ML model with a current decision tree ML model for the same data.

This option involves comparing two different types of models (a neural network and a decision tree). This is not a typical scenario for back-to-back testing because the models

are inherently different and would not be expected to produce identical results even on the same data.

D). Comparison of the results of the current neural network ML model on the current data set with a slightly modified data set.

This option involves comparing the outputs of the same model on slightly different datasets. This could be seen as a form of robustness testing or sensitivity analysis, but not typical back-to-back testing as it doesn't involve comparing multiple implementations.

Based on this analysis, option C is the one that describes a situation of back-to-back testing the least because it compares two fundamentally different models, which is not the intent of back-to-back testing.

NEW QUESTION: 53

Which statement about testing levels for AI-based systems is correct?

- A. Input data testing checks whether the inputs from the data pipeline are received by the model correctly and exchanged with all system components
- B. Acceptance testing checks non-functional requirements such as explainability
- C. ML model testing ensures that the relevant ML functional performance criteria are met
- D. If AI is offered as a service, system testing includes API tests of the service

Answer: C (LEAVE A REPLY)

Section 4.3 - Test Levels for AI Systems clearly defines ML model testing as the level at which testers evaluate whether an ML model fulfills its functional performance criteria, including accuracy, precision, recall, F1, robustness, stability, and fairness. Therefore, Option C is the correct and syllabus-aligned statement.

NEW QUESTION: 54

Which ONE of the following statements about the hardware used to implement ML systems is MOST likely to be correct?

- A. Specialist hardware is necessary for ML
- B. Less bits are required for hardware supporting ML
- C. Support for complex operations is required
- D. A higher clock speed is required

Answer: C (LEAVE A REPLY)

ML systems often require hardware that supports complex operations, such as matrix multiplications and other linear algebra computations. These operations are fundamental to many machine learning algorithms, particularly in deep learning. Specialist hardware (e.g., GPUs) may be used for efficiency, but complex operations are the core requirement.

NEW QUESTION: 55

Which ONE of the following would be the LEAST effective input to an AI-based test optimization process?

- A. Previously failing tests

- B. Defect reports
- C. Test environment downtime
- D. Source control data

Answer: C ([LEAVE A REPLY](#))

Test environment downtime would be the least effective input to an AI-based test optimization process. While downtime may affect the ability to run tests, it does not provide direct insight into the quality of the tests or the effectiveness of the AI model being tested. In contrast, previously failing tests, defect reports, and source control data are much more relevant as they provide information about areas where the system may have issues or require further optimization.

NEW QUESTION: 56

Which of the following approaches would help overcome testing challenges associated with probabilistic and non-deterministic AI-based systems?

- A. Run the test several times to ensure that the AI always returns the same correct test result.
- B. Decompose the system test into multiple data ingestion tests to determine if the AI system is getting a sufficient volume of input data.
- C. Decompose the system test into multiple data ingestion tests to determine if the AI system is getting precise and accurate input data.
- D. Run the test several times to generate a statistically valid test result to ensure that an appropriate number of answers are accurate.

Answer: D ([LEAVE A REPLY](#))

Probabilistic and non-deterministic AI-based systems do not always produce the same output for identical inputs. This makes traditional testing approaches ineffective. Instead, the best approach is to run tests multiple times and analyze results statistically. Statistical Validity: Running tests multiple times ensures that observed results are statistically significant. Instead of relying on a single test run, analyzing multiple iterations helps determine trends, probabilities, and outliers.

Expected Result Tolerance: AI-based systems may produce different results within an acceptable range. Defining acceptable tolerances (e.g., "result must be within 2% of the optimal value") improves test effectiveness.

NEW QUESTION: 57

Which ONE of the following options is the MOST APPROPRIATE stage of the ML workflow to set model and algorithm hyperparameters?

- A. Evaluating the model
- B. Deploying the model
- C. Tuning the model
- D. Data testing

Answer: ([SHOW ANSWER](#))

Setting model and algorithm hyperparameters is an essential step in the machine learning workflow, primarily occurring during the tuning phase.

Evaluating the model (A): This stage involves assessing the model's performance using metrics and does not typically include the setting of hyperparameters.

Deploying the model (B): Deployment is the stage where the model is put into production and used in real-world applications. Hyperparameters should already be set before this stage.

Tuning the model (C): This is the correct stage where hyperparameters are set. Tuning involves adjusting the hyperparameters to optimize the model's performance.

Data testing (D): Data testing involves ensuring the quality and integrity of the data used for training and testing the model. It does not include setting hyperparameters.

Hence, the most appropriate stage of the ML workflow to set model and algorithm hyperparameters is C. Tuning the model.

NEW QUESTION: 58

An airline has created a ML model to project fuel requirements for future flights. The model imports weather data such as wind speeds and temperatures, calculates flight routes based on historical routings from air traffic control, and estimates loads from average passenger and baggage weights. The model performed within an acceptable standard for the airline throughout the summer but as winter set in the load weights became less accurate. After some exploratory data analysis it became apparent that luggage weights were higher in the winter than in summer.

Which of the following statements BEST describes the problem and how it could have been prevented?

- A.** The model suffers from drift and therefore should be regularly tested to ensure that any occurrences of drift are detected soon enough for the problem to be mitigated.
- B.** The model suffers from drift and therefore the performance standard should be eased until a new model with more transparency can be developed.
- C.** The model suffers from corruption and therefore should be reloaded into the computer system being used, preferably with a method of version control to prevent further changes.
- D.** The model suffers from a lack of transparency and therefore should be regularly tested to ensure that any progressive errors are detected soon enough for the problem to be mitigated.

Answer: ([SHOW ANSWER](#))

The syllabus states:

"Concept drift occurs when the operational environment changes without the trained model changing correspondingly. The outputs of the model become less accurate and less useful. Therefore, the operational model should be regularly evaluated against its acceptance criteria."

NEW QUESTION: 59

Which ONE of the following statements is true about dynamic testing for inappropriate bias?

- A. It can be necessary to obtain additional attributes about the data being processed
- B. Reviewing the source of the training data can reveal inappropriate bias
- C. Inappropriate bias only needs to be tested when protected characteristics such as race and gender are present in the inputs
- D. Testing should never be conducted in production

Answer: A ([LEAVE A REPLY](#))

NEW QUESTION: 60

A data scientist is performing unsupervised learning on a set of financial records relating to previous loan applications, and trying to predict defaults on future loans. They are reporting poor functional performance because of data issues.

Which ONE of the below is LEAST likely to be a contributory factor?

- A. Missing records for some accounts
- B. Missing data relating to whether loans were previously granted and repaid
- C. Inconsistent pre-processing of some records
- D. Irrelevant data included in the account records

Answer: D ([LEAVE A REPLY](#))

Irrelevant data included in the account records is less likely to contribute to poor functional performance in unsupervised learning, especially compared to missing records, missing key data (such as whether loans were granted or repaid), or inconsistent pre-processing. While irrelevant data can affect the quality of the model, missing or inconsistent data typically has a more direct negative impact on unsupervised learning models.

NEW QUESTION: 61

A bank wants to use an algorithm to determine which applicants should be given a loan. The bank hires a data scientist to construct a logistic regression model to predict whether the applicant will repay the loan or not. The bank has enough data on past customers to randomly split the data into a training data set and a test/validation data set. A logistic regression model is constructed on the training data set using the following independent variables:

- * Gender
- * Marital status
- * Number of dependents
- * Education
- * Income
- * Loan amount
- * Loan term
- * Credit score

The model reveals that those with higher credit scores and larger total incomes are more likely to repay their loans. The data scientist has suggested that there might be bias present in the model based on previous models created for other banks.

Given this information, what is the best test approach to check for potential bias in the model?

- A.** Experienced-based testing should be used to confirm that the training data set is operationally relevant. This can include applying exploratory data analysis (EDA) to check for bias within the training data set.
- B.** Back-to-back testing should be used to compare the model created using the training data set to another model created using the test data set, if the two models significantly differ, it will indicate there is bias in the original model.
- C.** Acceptance testing should be used to make sure the algorithm is suitable for the customer. The team can re-work the acceptance criteria such that the algorithm is sure to correctly predict the remaining applicants that have been set aside for the validation data set ensuring no bias is present.
- D.** A/B testing should be used to verify that the test data set does not detect any bias that might have been introduced by the original training data. If the two models significantly differ, it will indicate there is bias in the original model.

Answer: A (LEAVE A REPLY)

The syllabus mentions that experience-based testing and EDA are effective for detecting biases:

"Experience-based testing can be used to verify that the training dataset is operationally relevant and identify potential sources of bias. EDA is also useful for exploring the data and understanding any relationships that might lead to bias in the model."

Valid CT-AI Dumps shared by Actual4test.com for Helping Passing CT-AI Exam!

Actual4test.com now offer the **newest CT-AI exam dumps**, the Actual4test.com CT-AI exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com CT-AI dumps with Test Engine here:

https://www.actual4test.com/CT-AI_examcollection.html (162 Q&As Dumps, **30%OFF**

Special Discount: Freepdfdumps)

NEW QUESTION: 62

Which ONE of the below statements BEST describes why test environments for autonomous systems might need to be different to other test environments?

- A.** Tools may be required to simulate extreme scenarios
- B.** Non-determinism may need to be introduced into the environment
- C.** AI-specific hardware may be required
- D.** Tools may be required to provide explanations of system behavior

Answer: (SHOW ANSWER)

Tools may be required to simulate extreme scenarios because autonomous systems often need to be tested under challenging or rare conditions that are difficult to replicate in a real-world environment. These scenarios, such as extreme weather, system failures, or unexpected behaviors, need to be simulated in a controlled environment to ensure the system can handle a wide range of situations safely and effectively. This makes the test environment for autonomous systems unique compared to other types of testing.

NEW QUESTION: 63

Which of the following is THE LEAST appropriate tests to be performed for testing a feature related to autonomy?

- A. Test for human handover to give rest to the system.
- B. Test for human handover when it should actually not be relinquishing control.
- C. Test for human handover requiring mandatory relinquishing control.
- D. Test for human handover after a given time interval.

Answer: B (LEAVE A REPLY)

Testing Autonomy: Testing for human handover when it should not be relinquishing control is the least appropriate because it contradicts the very definition of autonomous systems.

NEW QUESTION: 64

A startup company has implemented a new facial recognition system for a banking application for mobile devices. The application is intended to learn at run-time on the device to determine if the user should be granted access. It also sends feedback over the Internet to the application developers. The application deployment resulted in continuous restarts of the mobile devices.

Which of the following is the most likely cause of the failure?

- A. The feedback requires a physical connection and cannot be sent over the Internet.
- B. Mobile operating systems cannot process machine learning algorithms.
- C. The size of the application is consuming too much of the phone's storage capacity.
- D. The training, processing, and diagnostic generation are too computationally intensive for the mobile device hardware to handle.

Answer: D (LEAVE A REPLY)

The syllabus highlights that on-device training and processing require considerable computational power, which may exceed the capabilities of some mobile devices:

"Self-learning and continuous learning systems require large amounts of computational power, which can impact system performance and stability if the hardware is not powerful enough."

NEW QUESTION: 65

Which ONE of the following combinations of Training, Validation, Testing data is used during the process of learning/creating the model?

- A. Training data - validation data - test data
- B. Training data - validation data
- C. Training data - test data
- D. Validation data - test data

Answer: A (LEAVE A REPLY)

The process of developing a machine learning model typically involves the use of three types of datasets:

Training Data: This is used to train the model, i.e., to learn the patterns and relationships in the data.

Validation Data: This is used to tune the model's hyperparameters and to prevent overfitting during the training process.

Test Data: This is used to evaluate the final model's performance and to estimate how it will perform on unseen data.

Let's analyze each option:

A). Training data - validation data - test data

This option correctly includes all three types of datasets used in the process of creating and validating a model. The training data is used for learning, validation data for tuning, and test data for final evaluation.

B). Training data - validation data

This option misses the test data, which is crucial for evaluating the model's performance on unseen data after the training and validation phases.

C). Training data - test data

This option misses the validation data, which is important for tuning the model and preventing overfitting during training.

D). Validation data - test data

This option misses the training data, which is essential for the initial learning phase of the model.

Therefore, the correct answer is A because it includes all necessary datasets used during the process of learning and creating the model: training, validation, and test data.

NEW QUESTION: 66

A tourist calls an airline to book a ticket and is connected with an automated system which is able to recognize speech, understand requests related to purchasing a ticket, and provide relevant travel options. When the tourist asks about the expected weather at the destination or potential impacts on operations because of the tight labor market the only response from the automated system is: "I don't understand your question." This AI system should be categorized as?

- A. General AI
- B. Narrow AI
- C. Super AI
- D. Conventional AI

Answer: B (LEAVE A REPLY)

Narrow AI (also known as Weak AI) is designed to perform specific tasks or solve particular problems within a limited domain, such as booking a ticket or recognizing speech. In this case, the AI system is specialized in handling flight bookings and related tasks but is not capable of addressing broader or more complex questions, like weather conditions or labor market impacts.

NEW QUESTION: 67

Which ONE of the following is the BEST option to optimize the regression test selection and prevent the regression suite from growing large?

- A. Identifying suitable tests by looking at the complexity of the test cases.
- B. Using of a random subset of tests.
- C. Automating test scripts using AI-based test automation tools.
- D. Using an AI-based tool to optimize the regression test suite by analyzing past test results

Answer: D (LEAVE A REPLY)

This is the most effective approach as AI-based tools can analyze historical test data, identify patterns, and prioritize tests that are more likely to catch defects based on past results. This method ensures an optimized and manageable regression test suite by focusing on the most impactful test cases.

NEW QUESTION: 68

A beer company is trying to understand how much recognition its logo has in the market. It plans to do that by monitoring images on various social media platforms using a pre-trained neural network for logo detection. This particular model has been trained by looking for words, as well as matching colors on social media images. The company logo has a big word across the middle with a bold blue and magenta border. Which associated risk is most likely to occur when using this pre-trained model?

- A. There is no risk, as the model has already been trained
- B. Insufficient function; the model was not trained to check for colors or words
- C. Improper data preparation
- D. Inherited bias: the model could have inherited unknown defects

Answer: D (LEAVE A REPLY)

According to the syllabus, pre-trained models often inherit biases and limitations from the data and processes used in their original training, which may not align with the new use case.

Specifically, the syllabus states:

"When using a pre-trained model, the training data and process cannot be fully controlled or known by the user of the model. As a result, the model can inherit biases or inaccuracies that were part of its original development and training process."

Valid CT-AI Dumps shared by Actual4test.com for Helping Passing CT-AI Exam!

Actual4test.com now offer the **newest CT-AI exam dumps**, the Actual4test.com CT-AI exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com CT-AI dumps with Test Engine here:

https://www.actual4test.com/CT-AI_examcollection.html (**162** Q&As Dumps, **30%OFF**

Special Discount: Freepdfdumps)