



ISTQB

Exam Questions CT-AI

Certified Tester AI Testing Exam

About ExamBible

Your Partner of IT Exam

Found in 1998

ExamBible is a company specialized on providing high quality IT exam practice study materials, especially Cisco CCNA, CCDA, CCNP, CCIE, Checkpoint CCSE, CompTIA A+, Network+ certification practice exams and so on. We guarantee that the candidates will not only pass any IT exam at the first attempt but also get profound understanding about the certificates they have got. There are so many alike companies in this industry, however, ExamBible has its unique advantages that other companies could not achieve.

Our Advances

* 99.9% Uptime

All examinations will be up to date.

* 24/7 Quality Support

We will provide service round the clock.

* 100% Pass Rate

Our guarantee that you will pass the exam.

* Unique Gurantee

If you do not pass the exam at the first time, we will not only arrange FULL REFUND for you, but also provide you another exam of your claim, ABSOLUTELY FREE!

NEW QUESTION 1

A ML engineer is trying to determine the correctness of the new open-source implementation "X", of a supervised regression algorithm implementation. R-Square is one of the functional performance metrics used to determine the quality of the model.

Which ONE of the following would be an APPROPRIATE strategy to achieve this goal? SELECT ONE OPTION

- A. Add 10% of the rows randomly and create another model and compare the R-Square scores of both the model.
- B. Train various models by changing the order of input features and verify that the R- Square score of these models vary significantly.
- C. Compare the R-Square score of the model obtained using two different implementations that utilize two different programming languages while using the same algorithm and the same training and testing data.
- D. Drop 10% of the rows randomly and create another model and compare the R-Square scores of both the models.

Answer: C

Explanation:

- ? A. Add 10% of the rows randomly and create another model and compare the R- Square scores of both the models.
- ? B. Train various models by changing the order of input features and verify that the R-Square score of these models vary significantly.
- ? C. Compare the R-Square score of the model obtained using two different implementations that utilize two different programming languages while using the same algorithm and the same training and testing data.
- ? D. Drop 10% of the rows randomly and create another model and compare the R- Square scores of both the models.

Therefore, option C is the most appropriate strategy because it directly compares the performance of the new implementation "X" with another implementation using the same algorithm and datasets, which helps in verifying the correctness of the implementation.

NEW QUESTION 2

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

SELECT ONE OPTION

- 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

Explanation:

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. The other tests are relevant to ensuring smooth operation and transitions between human and AI control.

Reference: ISTQB_CT-AI_Syllabus_v1.0, Sections on Testing Autonomous AI-Based Systems and Testing for Human-AI Interaction.

NEW QUESTION 3

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

SELECT ONE OPTION

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

Answer: C

Explanation:

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.

References:

? ISTQB CT-AI Syllabus Section 3.2 on the ML Workflow outlines the different stages of the ML process, including the tuning phase where hyperparameters are set.

? Sample Exam Questions document, Question #31 specifically addresses the stage in the ML workflow where hyperparameters are configured.

NEW QUESTION 4

Which ONE of the following models BEST describes a way to model defect prediction by looking at the history of bugs in modules by using code quality metrics of modules of historical versions as input?

SELECT ONE OPTION

- A. Identifying the relationship between developers and the modules developed by them.
- B. Search of similar code based on natural language processing.
- C. Clustering of similar code modules to predict based on similarity.
- D. Using a classification model to predict the presence of a defect by using code quality metrics as the input data.

Answer: D

Explanation:

Defect prediction models aim to identify parts of the software that are likely to contain defects by analyzing historical data and code quality metrics. The primary goal is to use this predictive information to allocate testing and maintenance resources effectively. Let's break down why option D is the correct choice:

? Understanding Classification Models:

? Input Data - Code Quality Metrics:

? Historical Data:

? Why Option D is Correct:

? Eliminating Other Options:

References:

? ISTQB CT-AI Syllabus, Section 9.5, Metamorphic Testing (MT), describes various testing techniques including classification models for defect prediction.

? "Using AI for Defect Prediction" (ISTQB CT-AI Syllabus, Section 11.5.1).

NEW QUESTION 5

Which ONE of the following options describes the LEAST LIKELY usage of AI for detection of GUI changes due to changes in test objects?

SELECT ONE OPTION

- A. Using a pixel comparison of the GUI before and after the change to check the differences.
- B. Using a computer vision to compare the GUI before and after the test object changes.
- C. Using a vision-based detection of the GUI layout changes before and after test object changes.
- D. Using a ML-based classifier to flag if changes in GUI are to be flagged for humans.

Answer: A

Explanation:

* A. Using a pixel comparison of the GUI before and after the change to check the differences.

? Pixel comparison is a traditional method and does not involve AI. It compares images at the pixel level, which can be effective but is not an intelligent approach. It is not considered an AI usage and is the least likely usage of AI for detecting GUI changes.

* B. Using computer vision to compare the GUI before and after the test object changes.

? Computer vision involves using AI techniques to interpret and process images. It is a likely usage of AI for detecting changes in the GUI.

* C. Using vision-based detection of the GUI layout changes before and after test object changes.

? Vision-based detection is another AI technique where the layout and structure of the GUI are analyzed to detect changes. This is a typical application of AI.

* D. Using a ML-based classifier to flag if changes in GUI are to be flagged for humans.

? An ML-based classifier can intelligently determine significant changes and decide if they need human review, which is a sophisticated AI application.

NEW QUESTION 6

Which ONE of the following options describes a scenario of A/B testing the LEAST? SELECT ONE OPTION

- A. A comparison of two different websites for the same company to observe from a user acceptance perspective.
- B. A comparison of two different offers in a recommendation system to decide on the more effective offer for same users.
- C. A comparison of the performance of an ML system on two different input datasets.
- D. A comparison of the performance of two different ML implementations on the same input data.

Answer: C

Explanation:

A/B testing, also known as split testing, is a method used to compare two versions of a product or system to determine which one performs better. It is widely used in web development, marketing, and machine learning to optimize user experiences and model performance. Here's why option C is the least descriptive of an A/B testing scenario:

? Understanding A/B Testing:

? Application in Machine Learning:

? Why Option C is the Least Descriptive:

? Clarifying the Other Options:

References:

? ISTQB CT-AI Syllabus, Section 9.4, A/B Testing, explains the methodology and application of A/B testing in various contexts.

? "Understanding A/B Testing" (ISTQB CT-AI Syllabus).

NEW QUESTION 7

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

SELECT ONE OPTION

- 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

Explanation:

Detecting biases in the ML pipeline involves various tests to ensure fairness and accuracy throughout the ML process.

? Testing the distribution shift in the training data for inappropriate bias (A): This involves checking if there is any shift in the data distribution that could lead to bias in the model. It is an important test but not the most direct method for detecting biases.

? Test the model during model evaluation for data bias (B): 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.

? Testing the data pipeline for any sources for algorithmic bias (C): This test is crucial as it helps identify biases that may originate from the data processing and transformation stages within the pipeline. Detecting sources of algorithmic bias ensures that the model does not inherit biases from these processes.

? Check the input test data for potential sample bias (D): While this is an important step, it focuses more on the input data and less on the overall data pipeline. Hence, the most likely useful test to help detect different kinds of biases in the ML pipeline is B. Test the model during model evaluation for data bias.

References:

? ISTQB CT-AI Syllabus Section 8.3 on Testing for Algorithmic, Sample, and Inappropriate Bias discusses various tests that can be performed to detect biases at different stages of the ML pipeline.

? Sample Exam Questions document, Question #32 highlights the importance of evaluating the model for biases.

NEW QUESTION 8

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?

SELECT ONE OPTION

- 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: B

Explanation:

? 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.

Based on the provided information, option B (sample bias) best describes the situation because the training data is skewed towards ethnicity A, potentially leading to biased model performance.

NEW QUESTION 9

In the near future, technology will have evolved, and AI will be able to learn multiple tasks by itself without needing to be retrained, allowing it to operate even in new environments. The cognitive abilities of AI are similar to a child of 1-2 years.??

In the above quote, which ONE of the following options is the correct name of this type of AI?

SELECT ONE OPTION

- A. Technological singularity
- B. Narrow AI
- C. Super AI
- D. General AI

Answer: D

Explanation:

* A. Technological singularity

? Technological singularity refers to a hypothetical point in the future when AI surpasses human intelligence and can continuously improve itself without human intervention. This scenario involves capabilities far beyond those described in the question.

* B. Narrow AI

? Narrow AI, also known as weak AI, is designed to perform a specific task or a narrow range of tasks. It does not have general cognitive abilities and cannot learn multiple tasks by itself without retraining.

* C. Super AI

? Super AI refers to an AI that surpasses human intelligence and capabilities across all fields. This is an advanced concept and not aligned with the description of having cognitive abilities similar to a young child.

* D. General AI

? General AI, or strong AI, has the ability to understand, learn, and apply knowledge across a wide range of tasks, similar to human cognitive abilities. It aligns with the description of AI that can learn multiple tasks and operate in new environments without needing retraining.

NEW QUESTION 10

In a conference on artificial intelligence (AI), a speaker made the statement, "The current implementation of AI using models which do NOT change by themselves is NOT true AI". Based on your understanding of AI, is this above statement CORRECT or INCORRECT and why?

SELECT ONE OPTION

- A. This statement is incorrec
- B. Current AI is true AI and there is no reason to believe that this fact will change over time.
- C. This statement is correc
- D. In general, what is considered AI today may change over time.
- E. This statement is incorrec
- F. What is considered AI today will continue to be AI even as technology evolves and changes.
- G. This statement is correc
- H. In general, today the term AI is utilized incorrectly.

Answer: B

Explanation:

A. This statement is incorrect. Current AI is true AI and there is no reason to believe that this fact will change over time.

? AI is an evolving field, and the definition of what constitutes AI can change as technology advances.

* B. This statement is correct. In general, what is considered AI today may change over time.

? The term AI is dynamic and has evolved over the years. What is considered AI today might be viewed as standard computing in the future. Historically, as technologies become mainstream, they often cease to be considered "AI".

* C. This statement is incorrect. What is considered AI today will continue to be AI even as technology evolves and changes.

? This perspective does not account for the historical evolution of the definition of AI.

As new technologies emerge, the boundaries of AI shift.

* D. This statement is correct. In general, today the term AI is utilized incorrectly.

? While some may argue this, it is not a universal truth. The term AI encompasses a broad range of technologies and applications, and its usage is generally consistent with current technological capabilities.

NEW QUESTION 10

??BioSearch?? is creating an AI model used for predicting cancer occurrence via examining X- Ray images. The accuracy of the model in isolation has been found to be good. However, the users of the model started complaining of the poor quality of results, especially inability to detect real cancer cases, when put to practice in the diagnosis lab, leading to stopping of the usage of the model.

A testing expert was called in to find the deficiencies in the test planning which led to the above scenario.

Which ONE of the following options would you expect to MOST likely be the reason to be discovered by the test expert?

SELECT ONE OPTION

- A. A lack of similarity between the training and testing data.
- B. The input data has not been tested for quality prior to use for testing.
- C. A lack of focus on choosing the right functional-performance metrics.
- D. A lack of focus on non-functional requirements testing.

Answer: A

Explanation:

The question asks which deficiency is most likely to be discovered by the test expert given the scenario of poor real-world performance despite good isolated accuracy.

? A lack of similarity between the training and testing data (A): This is a common issue in ML where the model performs well on training data but poorly on real- world data due to a lack of representativeness in the training data. This leads to poor generalization to new, unseen data.

? The input data has not been tested for quality prior to use for testing (B): While data quality is important, this option is less likely to be the primary reason for the described issue compared to the representativeness of training data.

? A lack of focus on choosing the right functional-performance metrics (C): Proper metrics are crucial, but the issue described seems more related to the data mismatch rather than metric selection.

? A lack of focus on non-functional requirements testing (D): Non-functional requirements are important, but the scenario specifically mentions issues with detecting real cancer cases, pointing more towards data issues.

References:

? ISTQB CT-AI Syllabus Section 4.2 on Training, Validation, and Test Datasets emphasizes the importance of using representative datasets to ensure the model generalizes well to real-world data.

? Sample Exam Questions document, Question #40 addresses issues related to data representativeness and model generalization.

NEW QUESTION 14

"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? SELECT ONE OPTION

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

Answer: C

Explanation:

? Prevalence Rate and Model Performance:

? Importance of Recall:

? Importance of Precision:

? Balancing Recall and Precision:

? Accuracy and Specificity:

? Conclusion:

: This explanation aligns with the principles outlined in the ISTQB CT-AI Syllabus, particularly sections on performance metrics for ML models and handling imbalanced datasets (Chapter 5: ML Functional Performance Metrics).

NEW QUESTION 18

Pairwise testing can be used in the context of self-driving cars for controlling an explosion in the number of combinations of parameters.

Which ONE of the following options is LEAST likely to be a reason for this incredible growth of parameters?

SELECT ONE OPTION

- A. Different Road Types
- B. Different weather conditions
- C. ML model metrics to evaluate the functional performance
- D. Different features like ADAS, Lane Change Assistance etc.

Answer: C

Explanation:

Pairwise testing is used to handle the large number of combinations of parameters that can arise in complex systems like self-driving cars. The question asks which of the given options is least likely to be a reason for the explosion in the number of parameters.

? Different Road Types (A): Self-driving cars must operate on various road types, such as highways, city streets, rural roads, etc. Each road type can have different characteristics, requiring the car's system to adapt and handle different scenarios. Thus, this is a significant factor contributing to the growth of parameters.

? Different Weather Conditions (B): Weather conditions such as rain, snow, fog, and bright sunlight significantly affect the performance of self-driving cars. The car's sensors and algorithms must adapt to these varying conditions, which adds to the number of parameters that need to be considered.

? ML Model Metrics to Evaluate Functional Performance (C): While evaluating machine learning (ML) model performance is crucial, it does not directly contribute to the explosion of parameter combinations in the same way that road types, weather conditions, and car features do. Metrics are used to measure and assess performance but are not themselves variable conditions that the system must handle.

? Different Features like ADAS, Lane Change Assistance, etc. (D): Advanced Driver Assistance Systems (ADAS) and other features add complexity to self-driving cars. Each feature can have multiple settings and operational modes, contributing to the overall number of parameters.

Hence, the least likely reason for the incredible growth in the number of parameters is C. ML model metrics to evaluate the functional performance.

References:

? ISTQB CT-AI Syllabus Section 9.2 on Pairwise Testing discusses the application of this technique to manage the combinations of different variables in AI-based systems, including those used in self-driving cars.

? Sample Exam Questions document, Question #29 provides context for the explosion in parameter combinations in self-driving cars and highlights the use of pairwise testing as a method to manage this complexity.

NEW QUESTION 20

.....

Relate Links

100% Pass Your CT-AI Exam with ExamBible Prep Materials

<https://www.exambible.com/CT-AI-exam/>

Contact us

We are proud of our high-quality customer service, which serves you around the clock 24/7.

Viste - <https://www.exambible.com/>