

iSQL.CT-AI.by.Smile.28q

Number: CT-AI
Passing Score: 800
Time Limit: 120
File Version: 3.0

Exam Code: CT-AI

Exam Name: Certified Tester AI Testing



Exam A

QUESTION 1

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

Correct Answer: D

Section:

Explanation:

A . Identifying suitable tests by looking at the complexity of the test cases.

While complexity analysis can help in selecting important test cases, it does not directly address the issue of optimizing the entire regression suite effectively.

B . Using a random subset of tests.

Randomly selecting test cases may miss critical tests and does not ensure an optimized regression suite. This approach lacks a systematic method for ensuring comprehensive coverage.

C . Automating test scripts using AI-based test automation tools.

Automation helps in running tests efficiently but does not inherently optimize the selection of tests to prevent the suite from growing too large.

D . Using an AI-based tool to optimize the regression test suite by analyzing past test results.

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.

Therefore, the correct answer is D because using an AI-based tool to analyze past test results is the best option to optimize regression test selection and manage the size of the regression suite effectively.

QUESTION 2

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?

- 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.

Correct Answer: C

Section:

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.

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.

QUESTION 3

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.

Correct Answer: C

Section:

Explanation:

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.

ISTQB CT-AI Syllabus Section 2.1 on Flexibility and Adaptability discusses the importance of flexibility in AI systems and how it enables easier modification and adaptability to new situations.

Sample Exam Questions document, Question #30 highlights the importance of flexibility in AI systems.

QUESTION 4

Written requirements are given in text documents, which ONE of the following options is the BEST way to generate test cases from these requirements?

- A. Natural language processing on textual requirements
- B. Analyzing source code for generating test cases
- C. Machine learning on logs of execution
- D. GUI analysis by computer vision

Correct Answer: A

Section:

Explanation:

When written requirements are given in text documents, the best way to generate test cases is by using Natural Language Processing (NLP). Here's why:

Natural Language Processing (NLP): NLP can analyze and understand human language. It can be used to process textual requirements to extract relevant information and generate test cases. This method is efficient in handling large volumes of textual data and identifying key elements necessary for testing.

Why Not Other Options:

Analyzing source code for generating test cases: This is more suitable for white-box testing where the code is available, but it doesn't apply to text-based requirements.

Machine learning on logs of execution: This approach is used for dynamic analysis based on system behavior during execution rather than static textual requirements.

GUI analysis by computer vision: This is used for testing graphical user interfaces and is not applicable to text-based requirements.

QUESTION 5

Upon testing a model used to detect rotten tomatoes, the following data was observed by the test engineer, based on certain number of tomato images.

Confusion Matrix	Actually Rotten	Actually Fresh
Predicted Rotten	45	8
Predicted Fresh	5	42

For this confusion matrix which combinations of values of accuracy, recall, and specificity respectively is CORRECT?

- A. 0.87,0.9, 0.84
- B. 1,0.87,0.84
- C. 1,0.9, 0.8
- D. 0.84,1,0.9

Correct Answer: A

Section:

Explanation:

To calculate the accuracy, recall, and specificity from the confusion matrix provided, we use the following formulas:

Confusion Matrix:

Actually Rotten: 45 (True Positive), 8 (False Positive)

Actually Fresh: 5 (False Negative), 42 (True Negative)

Accuracy:

Accuracy is the proportion of true results (both true positives and true negatives) in the total population.

Formula: $\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$

Calculation: $\text{Accuracy} = \frac{45 + 42}{45 + 42 + 8 + 5} = \frac{87}{100} = 0.87$

Recall (Sensitivity):

Recall is the proportion of true positive results in the total actual positives.

Formula: $\text{Recall} = \frac{TP}{TP + FN}$

Calculation: $\text{Recall} = \frac{45}{45 + 5} = \frac{45}{50} = 0.9$

Specificity:

Specificity is the proportion of true negative results in the total actual negatives.

Formula: $\text{Specificity} = \frac{TN}{TN + FP}$

Calculation: $\text{Specificity} = \frac{42}{42 + 8} = \frac{42}{50} = 0.84$

Therefore, the correct combinations of accuracy, recall, and specificity are 0.87, 0.9, and 0.84 respectively.

ISTQB CT-AI Syllabus, Section 5.1, Confusion Matrix, provides detailed formulas and explanations for calculating various metrics including accuracy, recall, and specificity.

'ML Functional Performance Metrics' (ISTQB CT-AI Syllabus, Section 5).

QUESTION 6

The activation value output for a neuron in a neural network is obtained by applying computation to the neuron.

Which ONE of the following options BEST describes the inputs used to compute the activation value?

- A. Individual bias at the neuron level, activation values of neurons in the previous layer, and weights assigned to the connections between the neurons.

- B. Activation values of neurons in the previous layer, and weights assigned to the connections between the neurons.
- C. Individual bias at the neuron level, and weights assigned to the connections between the neurons.
- D. Individual bias at the neuron level, and activation values of neurons in the previous layer.

Correct Answer: A

Section:

Explanation:

In a neural network, the activation value of a neuron is determined by a combination of inputs from the previous layer, the weights of the connections, and the bias at the neuron level. Here's a detailed breakdown:

Inputs for Activation Value:

Activation Values of Neurons in the Previous Layer: These are the outputs from neurons in the preceding layer that serve as inputs to the current neuron.

Weights Assigned to the Connections: Each connection between neurons has an associated weight, which determines the strength and direction of the input signal.

Individual Bias at the Neuron Level: Each neuron has a bias value that adjusts the input sum, allowing the activation function to be shifted.

Calculation:

The activation value is computed by summing the weighted inputs from the previous layer and adding the bias.

Formula: $z = (w_i a_i) + b_z = \sum (w_i \cdot a_i) + b_z = (w_i a_i) + b$, where w_i are the weights, a_i are the activation values from the previous layer, and b is the bias.

The activation function (e.g., sigmoid, ReLU) is then applied to this sum to get the final activation value.

Why Option A is Correct:

Option A correctly identifies all components involved in computing the activation value: the individual bias, the activation values of the previous layer, and the weights of the connections.

Eliminating Other Options:

B . Activation values of neurons in the previous layer, and weights assigned to the connections between the neurons: This option misses the bias, which is crucial.

C . Individual bias at the neuron level, and weights assigned to the connections between the neurons: This option misses the activation values from the previous layer.

D . Individual bias at the neuron level, and activation values of neurons in the previous layer: This option misses the weights, which are essential.

ISTQB CT-AI Syllabus, Section 6.1, Neural Networks, discusses the components and functioning of neurons in a neural network.

'Neural Network Activation Functions' (ISTQB CT-AI Syllabus, Section 6.1.1).

QUESTION 7

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.

Correct Answer: C

Section:

Explanation:

The question asks which test is least likely to be performed during the ML model testing phase. Let's consider each option:

Testing the accuracy of the classification model (A): Accuracy testing is a fundamental part of the ML model testing phase. It ensures that the model correctly classifies the data as intended and meets the required performance metrics.

Testing the API of the service powered by the ML model (B): Testing the API is crucial, especially if the ML model is deployed as part of a service. This ensures that the service integrates well with other systems and that the API performs as expected.

Testing the speed of the training of the model (C): 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.

Testing the speed of the prediction by the model (D): Testing the speed of prediction is important to ensure that the model meets performance requirements in a production environment, especially for real-time applications. ISTQB CT-AI Syllabus Section 3.2 on ML Workflow and Section 5 on ML Functional Performance Metrics discuss the focus of testing during the model testing phase, which includes accuracy and prediction speed but not the training speed.

QUESTION 8

A software component uses machine learning to recognize the digits from a scan of handwritten numbers. In the scenario above, which type of Machine Learning (ML) is this an example of?

- A. Reinforcement learning
- B. Regression
- C. Classification
- D. Clustering

Correct Answer: C

Section:

Explanation:

Recognizing digits from a scan of handwritten numbers using machine learning is an example of classification. Here's a breakdown:

Classification: This type of machine learning involves categorizing input data into predefined classes. In this scenario, the input data (handwritten digits) are classified into one of the 10 digit classes (0-9).

Why Not Other Options:

Reinforcement Learning: This involves learning by interacting with an environment to achieve a goal, which does not fit the problem of recognizing digits.

Regression: This is used for predicting continuous values, not discrete categories like digit recognition.

Clustering: This involves grouping similar data points together without predefined classes, which is not the case here.

QUESTION 9

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



Correct Answer: B

Section:

Explanation:

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.

Reference: ISTQB_CT-AI_Syllabus_v1.0, Section 4.5 Data Labelling for Supervised Learning, which discusses various approaches to data labeling, including pre-labeled datasets, and their associated time and effort requirements.

QUESTION 10

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

Correct Answer: C

Section:

Explanation:

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.

Reference: ISTQB_CT-AI_Syllabus_v1.0, Section 7.6 Testing for Concept Drift, which explains the need to test for concept drift in models that might be affected by changing external factors.

QUESTION 11

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

Correct Answer: C

Section:

Explanation:

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.

Reference: ISTQB_CT-AI_Syllabus_v1.0, Section 6.2 Coverage Measures for Neural Networks, which details different types of coverage measures, including sign change coverage.

QUESTION 12

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.

Correct Answer: C

Section:

Explanation:

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.

QUESTION 13

Which ONE of the following options does NOT describe an AI technology related characteristic which differentiates AI test environments from other test environments?

- A. Challenges resulting from low accuracy of the models.
- B. The challenge of mimicking undefined scenarios generated due to self-learning
- C. The challenge of providing explainability to the decisions made by the system.
- D. Challenges in the creation of scenarios of human handover for autonomous systems.

Correct Answer: D

Section:

Explanation:

AI test environments have several unique characteristics that differentiate them from traditional test environments. Let's evaluate each option:

A . Challenges resulting from low accuracy of the models.

Low accuracy is a common challenge in AI systems, especially during initial development and training phases. Ensuring the model performs accurately in varied and unpredictable scenarios is a critical aspect of AI testing.

B . The challenge of mimicking undefined scenarios generated due to self-learning.

AI systems, particularly those that involve machine learning, can generate undefined or unexpected scenarios due to their self-learning capabilities. Mimicking and testing these scenarios is a unique challenge in AI environments.

C . The challenge of providing explainability to the decisions made by the system.

Explainability, or the ability to understand and articulate how an AI system arrives at its decisions, is a significant and unique challenge in AI testing. This is crucial for trust and transparency in AI systems.

D . Challenges in the creation of scenarios of human handover for autonomous systems.

While important, the creation of scenarios for human handover in autonomous systems is not a characteristic unique to AI test environments. It is more related to the operational and deployment challenges of autonomous systems rather than the intrinsic technology-related characteristics of AI .

Given the above points, option D is the correct answer because it describes a challenge related to operational deployment rather than a technology-related characteristic unique to AI test environments.

QUESTION 14

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

Correct Answer: A

Section:

Explanation:

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.

QUESTION 15

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.

Correct Answer: C

Section:

Explanation:

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. Let's analyze each option:

A . Clustering is classification of a continuous quantity.

This is incorrect. Classification typically involves discrete categories, whereas clustering involves grouping similar data points. Classification of continuous quantities is generally referred to as regression.

B . Clustering is supervised learning.

This is incorrect. Clustering is an unsupervised learning technique because it does not rely on labeled data.

C . Clustering is done without prior knowledge of output classes.

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

D . Clustering requires you to know the classes.

This is incorrect. Clustering does not require prior knowledge of classes. Instead, it aims to identify and form the classes or groups based on the data itself.

Therefore, the correct answer is C because clustering is an unsupervised learning technique done without prior knowledge of output classes.

QUESTION 16

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.

Correct Answer: B

Section:

Explanation:

A . This is an example of expert system bias.

Expert system bias refers to bias introduced by the rules or logic defined by experts in the system, not by the data distribution.

B . This is an example of sample bias.

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.

C . This is an example of hyperparameter bias.

Hyperparameter bias relates to the settings and configurations used during the training process, not the data distribution itself.

D . This is an example of algorithmic bias.

Algorithmic bias refers to biases introduced by the algorithmic processes and decision-making rules, not directly by the distribution of training data.

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.

QUESTION 17

Which ONE of the following activities is MOST relevant when addressing the scenario where you have more than the required amount of data available for the training?

- A. Feature selection
- B. Data sampling
- C. Data labeling
- D. Data augmentation

Correct Answer: B

Section:

Explanation:

A . Feature selection

Feature selection is the process of selecting the most relevant features from the data. While important, it is not directly about handling excess data.

B . Data sampling

Data sampling involves selecting a representative subset of the data for training. When there is more data than needed, sampling can be used to create a manageable dataset that maintains the statistical properties of the full dataset.

C . Data labeling

Data labeling involves annotating data for supervised learning. It is necessary for training models but does not address the issue of having excess data.

D . Data augmentation

Data augmentation is used to increase the size of the training dataset by creating modified versions of existing data. It is useful when there is insufficient data, not when there is excess data.

Therefore, the correct answer is B because data sampling is the most relevant activity when dealing with an excess amount of data for training.

QUESTION 18

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?

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

Correct Answer: B

Section:

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.

QUESTION 19

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.

Correct Answer: B

Section:

Explanation:

A . 64-bit CPUs.

While 64-bit CPUs are essential for handling large amounts of memory and performing complex computations, they are not specifically optimized for the types of operations commonly used in machine learning.

B . Hardware supporting fast matrix multiplication.

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.

C . High powered CPUs.

High powered CPUs are beneficial for general-purpose computing tasks and some aspects of ML, but they are not as efficient as specialized hardware like GPUs for matrix multiplication and other ML-specific tasks.

D . Hardware supporting high precision floating point operations.

High precision floating point operations are important for scientific computing and some specific AI tasks, but for many ML applications, fast matrix multiplication is more critical than high precision alone.

Therefore, the correct answer is B because hardware supporting fast matrix multiplication, such as GPUs, is most suitable for the parallel processing requirements of machine learning.

QUESTION 20

Which ONE of the following statements is a CORRECT adversarial example in the context of machine learning systems that are working on image classifiers.

- A. Black box attacks based on adversarial examples create an exact duplicate model of the original.
- B. These attack examples cause a model to predict the correct class with slightly less accuracy even though they look like the original image.
- C. These attacks can't be prevented by retraining the model with these examples augmented to the training data.
- D. These examples are model specific and are not likely to cause another model trained on same task to fail.

Correct Answer: D

Section:

Explanation:

A . Black box attacks based on adversarial examples create an exact duplicate model of the original.

Black box attacks do not create an exact duplicate model. Instead, they exploit the model by querying it and using the outputs to craft adversarial examples without knowledge of the internal workings.

B . These attack examples cause a model to predict the correct class with slightly less accuracy even though they look like the original image.

Adversarial examples typically cause the model to predict the incorrect class rather than just reducing accuracy. These examples are designed to be visually indistinguishable from the original image but lead to incorrect classifications.

C . These attacks can't be prevented by retraining the model with these examples augmented to the training data.

This statement is incorrect because retraining the model with adversarial examples included in the training data can help the model learn to resist such attacks, a technique known as adversarial training.

D . These examples are model specific and are not likely to cause another model trained on the same task to fail.

Adversarial examples are often model-specific, meaning that they exploit the specific weaknesses of a particular model. While some adversarial examples might transfer between models, many are tailored to the specific model they were generated for and may not affect other models trained on the same task.

Therefore, the correct answer is D because adversarial examples are typically model-specific and may not cause another model trained on the same task to fail.

QUESTION 21

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

Correct Answer: C

Section:

Explanation:

A . Regression

Regression is used to predict a continuous value and is not suitable for grouping people based on similar tastes.

B . Association

Association is used to find relationships between variables in large datasets, often in the form of rules (e.g., market basket analysis). It does not directly group individuals but identifies patterns of co-occurrence.

C . Clustering

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.

D . Classification

Classification is a supervised learning method used to categorize data points into predefined classes. It requires labeled data for training, which is not the case here as we want to identify groups without predefined labels. Therefore, the correct answer is C because clustering is the most suitable method for grouping people with similar tastes for targeted product marketing.

QUESTION 22

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



Correct Answer: C

Section:

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.

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.

QUESTION 23

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.

Correct Answer: B

Section:

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.

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.

QUESTION 24

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?

- 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.

Correct Answer: D

Section:

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:

Classification models are a type of supervised learning algorithm used to categorize or classify data into predefined classes or labels. In the context of defect prediction, the classification model would classify parts of the code as either 'defective' or 'non-defective' based on the input features.

Input Data - Code Quality Metrics:

The input data for these classification models typically includes various code quality metrics such as cyclomatic complexity, lines of code, number of methods, depth of inheritance, coupling between objects, etc. These metrics help the model learn patterns associated with defects.

Historical Data:

Historical versions of the code along with their defect records provide the labeled data needed for training the classification model. By analyzing this historical data, the model can learn which metrics are indicative of defects.

Why Option D is Correct:

Option D specifies using a classification model to predict the presence of defects by using code quality metrics as input data. This accurately describes the process of defect prediction using historical bug data and quality metrics.

Eliminating Other Options:

A . Identifying the relationship between developers and the modules developed by them: This does not directly involve predicting defects based on code quality metrics and historical data.

B . Search of similar code based on natural language processing: While useful for other purposes, this method does not describe defect prediction using classification models and code metrics.

C . Clustering of similar code modules to predict based on similarity: Clustering is an unsupervised learning technique and does not directly align with the supervised learning approach typically used in defect prediction models.

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).

QUESTION 25

Which ONE of the following options is an example that BEST describes a system with AI-based autonomous functions?

- A. A system that utilizes human beings for all important decisions.
- B. A fully automated manufacturing plant that uses no software.
- C. A system that utilizes a tool like Selenium.
- D. A system that is fully able to respond to its environment.

Correct Answer: D

Section:

Explanation:

AI-Based Autonomous Functions: An AI-based autonomous system is one that can respond to its environment without human intervention. The other options either involve human decisions or do not use AI at all.

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

QUESTION 26

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.

Correct Answer: B

Section:

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.

QUESTION 27

'AllerEgo' is a product that uses self-learning to predict the behavior of a pilot under combat situation for a variety of terrains and enemy aircraft formations. Post training the model was exposed to the real-world data and the model was found to be behaving poorly. A lot of data quality tests had been performed on the data to bring it into a shape fit for training and testing.

Which ONE of the following options is least likely to describe the possible reason for the fall in the performance, especially when considering the self-learning nature of the AI system?

- A. The difficulty of defining criteria for improvement before the model can be accepted. Defining criteria for improvement is a challenge in the acceptance of AI models, but it is not directly related to the performance drop in real-world scenarios. It relates more to the evaluation and deployment phase rather than affecting the model's real-time performance post-deployment.
- B. The fast pace of change did not allow sufficient time for testing. This can significantly affect the model's performance. If the system is self-learning, it needs to adapt quickly, and insufficient testing time can lead to incomplete learning and poor performance.
- C. The unknown nature and insufficient specification of the operating environment might have caused the poor performance. This is highly likely to affect performance. Self-learning AI systems require detailed specifications of the operating environment to adapt and learn effectively. If the environment is insufficiently specified, the model may fail to perform accurately in real-world scenarios.
- D. There was an algorithmic bias in the AI system. Algorithmic bias can significantly impact the performance of AI systems. If the model has biases, it will not perform well across different scenarios and data distributions. Given the context of the self-learning nature and the need for real-time adaptability, option A is least likely to describe the fall in performance because it deals with acceptance criteria rather than real-time performance issues.

Correct Answer: A

Section:

QUESTION 28

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?

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

Correct Answer: D

Section:

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.

