

AIP-210 Pass4sure Dumps & AIP-210 Sichere Praxis Dumps



Cisco- 210-260

IMPLEMENTING CISCO NETWORK SECURITY
Pass Pass4sure 210-260 exam in just 24
HOURS With 100% Guarantee

Top 100% REAL EXAM QUESTIONS ANSWERS

Get All PDF with Complete Questions Answers File from

www.pass4surebraindumps.com/210-260.html

100% Exam Passing Guarantee & Money Back Assurance



Wir sind der Schnellste, der das CertNexus AIP-210 Zertifikat erhält; wir sind noch der höchste, der Ihre Interessen schützt. Wir sind Fast2test. Fast2test kann Ihnen versprechen, dass die Testaufgaben von CertNexus AIP-210 Zertifizierungsprüfung 100% richtig und ganz umfassend sind. Nachdem Sie die Testfragen zur CertNexus AIP-210 Zertifizierung gekauft haben, werden Sie kostenlos die einjährige Aktualisierung genießen.

Die Produkte von Fast2test sind zuverlässig und von guter Qualität. Sie können im Internet teilweise die Demo zur CertNexus AIP-210 Zertifizierungsprüfung kostenlos als Probe herunterladen. Nach dem Benutzen, meine ich, werden Sie mit unseren Produkten zufrieden sein. Weshalb zögern Sie noch, wenn es so gute Produkte zum Bestehen der CertNexus AIP-210 Prüfung gibt. Schicken Sie doch schnell die Produkte von Fast2test in den Warenkorb.

>> AIP-210 Prüfungssimulationen <<

AIP-210 Vorbereitungsfragen - AIP-210 Fragenpool

Die echten und originalen Prüfungsfragen und Antworten zu AIP-210 Zertifizierung (CertNexus Certified Artificial Intelligence Practitioner (CAIP)) bei Fast2test wurden verfasst von unseren IT-Experten mit den Informationen von AIP-210 Prüfungen (CertNexus Certified Artificial Intelligence Practitioner (CAIP)) aus dem Testcenter wie PROMETRIC oder VUE.

CertNexus Certified Artificial Intelligence Practitioner (CAIP) AIP-210 Prüfungsfragen mit Lösungen (Q24-Q29):

24. Frage

You train a neural network model with two layers, each layer having four nodes, and realize that the model is underfit. Which of the actions below will NOT work to fix this underfitting?

- **A. Get more training data**
- B. Increase the complexity of the model
- C. Train the model for more epochs
- D. Add features to training data

Antwort: A

Begründung:

Underfitting is a problem that occurs when a model learns too little from the training data and fails to capture the underlying complexity or structure of the data. Underfitting can result from using insufficient or irrelevant features, a low complexity of the model, or a lack of training data. Underfitting can reduce the accuracy and generalization of the model, as it may produce oversimplified or inaccurate predictions. Some of the ways to fix underfitting are:

* Add features to training data: Adding more features or variables to the training data can help increase the information and diversity of the data, which can help the model learn more complex patterns and relationships.

* Increase the complexity of the model: Increasing the complexity of the model can help increase its expressive power and flexibility, which can help it fit better to the data. For example, adding more layers or nodes to a neural network can increase its complexity.

* Train the model for more epochs: Training the model for more epochs can help increase its learning ability and convergence, which can help it optimize its parameters and reduce its error.

Getting more training data will not work to fix underfitting, as it will not change the complexity or structure of the data or the model.

Getting more training data may help with overfitting, which is when a model learns too much from the training data and fails to generalize well to new or unseen data.

25. Frage

Which three security measures could be applied in different ML workflow stages to defend them against malicious activities? (Select three.)

- **A. Use Secrets Manager to protect credentials.**
- B. Disable logging for model access.
- **C. Use data encryption.**
- **D. Launch ML Instances In a virtual private cloud (VPC).**
- E. Monitor model degradation.
- F. Use max privilege to control access to ML artifacts.

Antwort: A,C,D

Begründung:

Security measures can be applied in different ML workflow stages to defend them against malicious activities, such as data theft, model tampering, or adversarial attacks. Some of the security measures are:

* Launch ML Instances In a virtual private cloud (VPC): A VPC is a logically isolated section of a cloud provider's network that allows users to launch and control their own resources. By launching ML instances in a VPC, users can enhance the security and privacy of their data and models, as well as restrict the access and traffic to and from the instances.

* Use data encryption: Data encryption is the process of transforming data into an unreadable format using a secret key or algorithm. Data encryption can protect the confidentiality, integrity, and availability of data at rest (stored in databases or files) or in transit (transferred over networks). Data encryption can prevent unauthorized access, modification, or leakage of sensitive data.

* Use Secrets Manager to protect credentials: Secrets Manager is a service that helps users securely store, manage, and retrieve secrets, such as passwords, API keys, tokens, or certificates. Secrets Manager can help users protect their credentials from unauthorized access or exposure, as well as rotate them automatically to comply with security policies.

26. Frage

What is the open framework designed to help detect, respond to, and remediate threats in ML systems?

- A. MITRE ATT&CK Matrix
- **B. Adversarial ML Threat Matrix**
- C. OWASP Threat and Safeguard Matrix
- D. Threat Susceptibility Matrix

Antwort: B

Begründung:

Explanation

The Adversarial ML Threat Matrix is an open framework designed to help detect, respond to, and remediate threats in ML systems. The Adversarial ML Threat Matrix is inspired by the MITRE ATT&CK Matrix¹, which is a framework for describing cyberattacks across various stages of an attack lifecycle. The Adversarial ML Threat Matrix adapts this framework to address specific threats and vulnerabilities in ML systems, such as data poisoning, model stealing, model evasion, or model inversion². The Adversarial ML Threat Matrix provides a structured way to organize and classify adversarial techniques, tactics, procedures, examples, and mitigations for ML systems².

27. Frage

Which of the following is NOT an activation function?

- A. ReLU
- B. Sigmoid
- **C. Additive**
- D. Hyperbolic tangent

Antwort: C

Begründung:

Explanation

An activation function is a function that determines the output of a neuron in a neural network based on its input. An activation function can introduce non-linearity into a neural network, which allows it to model complex and non-linear relationships between inputs and outputs. Some of the common activation functions are:

Sigmoid: A sigmoid function is a function that maps any real value to a value between 0 and 1. It has an S-shaped curve and is often used for binary classification or probability estimation.

Hyperbolic tangent: A hyperbolic tangent function is a function that maps any real value to a value between -1 and 1. It has a similar shape to the sigmoid function but is symmetric around the origin. It is often used for regression or classification problems.

ReLU: A ReLU (rectified linear unit) function is a function that maps any negative value to 0 and any positive value to itself. It has a piecewise linear shape and is often used for hidden layers in deep neural networks.

Additive is not an activation function, but rather a term that describes a property of some functions. Additive functions are functions that satisfy the condition $f(x+y) = f(x) + f(y)$ for any x and y . Additive functions are linear functions, which means they have a constant slope and do not introduce non-linearity.

28. Frage

A market research team has ratings from patients who have a chronic disease, on several functional, physical, emotional, and professional needs that stay unmet with the current therapy. The dataset also captures ratings on how the disease affects their day-to-day activities.

A pharmaceutical company is introducing a new therapy to cure the disease and would like to design their marketing campaign such that different groups of patients are targeted with different ads. These groups should ideally consist of patients with similar unmet needs.

Which of the following algorithms should the market research team use to obtain these groups of patients?

- A. Logistic regression
- B. Naive-Bayes
- **C. k-means clustering**
- D. k-nearest neighbors

Antwort: C

Begründung:

Explanation

k-means clustering is an algorithm that should be used by the market research team to obtain groups of patients with similar unmet needs. k-means clustering is an unsupervised learning technique that partitions the data into k clusters based on the similarity of the features. The algorithm iteratively assigns each data point to the cluster with the nearest centroid and updates the centroid until convergence. k-means clustering can help identify patterns and segments in the data that may not be obvious or intuitive. References: [K-means clustering - Wikipedia], [How to Run K-Means Clustering in Python]

• • • • •

AIP-210 Vorbereitungsfragen: <https://de.fast2test.com/AIP-210-premium-file.html>

AIP-210 Studienmaterialien: CertNexus Certified Artificial Intelligence Practitioner (CAIP) & AIP-210 Zertifizierungstraining

Unser Fast2test verspricht, dass Sie nur AIP-210 einmal die Prüfung bestehen und das Zertifikat von den Experten bekommen können.

- [illegible]