

UiPath.UiPath-SAIv1.v2026-02-16.q92

Exam Code:	UiPath-SAIv1
Exam Name:	UiPath Specialized AI Associate Exam (2023.10)
Certification Provider:	UiPath
Free Question Number:	92
Version:	v2026-02-16
# of views:	114
# of Questions views:	920
https://www.exam-tests.com/UiPath-SAIv1-exam/UiPath.UiPath-SAIv1.v2026-02-16.q92.html	

NEW QUESTION: 1

When designing a taxonomy in UiPath Communications Mining, what is the similarity between labels and general fields?

- A. They can be pre-trained or trained from scratch.
- B. They are structured in hierarchies to add levels of specificity.
- C. They are always assigned at the message level.
- D. They are entirely rule-based and follow a particular format.

Answer: (SHOW ANSWER)

Comprehensive and Detailed Explanation From Exact Extract:

In UiPath Communications Mining, both labels (used for classification) and general fields (used for data extraction) can be:

- * Trained from scratch using manually labeled data
- * Or based on pre-trained models or imported training sets

This makes both elements part of the supervised training process, and both contribute to improving model accuracy over time.

* UiPath Documentation Reference: Taxonomy Management - Communications Mining

NEW QUESTION: 2

What is the recommended number of documents per vendor to train the initial dataset?

- A. 5
- B. 10
- C. 15
- D. 20

Answer: B (LEAVE A REPLY)

According to the UiPath documentation, the recommended number of documents per vendor to train the initial dataset is 10. This means that for each vendor that provides a specific type of

document, such as invoices or receipts, you should have at least 10 samples of their documents in your training dataset. This helps to ensure that the dataset is balanced and representative of the real-world data, and that the machine learning model can learn from the variations and features of each vendor's documents. Having too few documents per vendor can lead to poor model performance and accuracy, while having too many documents from a single vendor can cause overfitting and bias¹.

References: 1: Document Understanding - Training High Performing Models

NEW QUESTION: 3

How long does the typical Machine Learning model deployment process take in UiPath AI Center?

- A. Less than 5 minutes.
- B. Between 5 and 10 minutes.
- C. Between 10 and 15 minutes.
- D. More than 15 minutes.

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

The typical machine learning model deployment process in UiPath AI Center usually takes between 10-15 minutes¹. This process involves wrapping the model in UiPath's serving framework and deploying it within a namespace on AI Fabric's Kubernetes cluster that is only accessible by your tenant¹. Please note that the actual time may vary depending on the complexity of the model and other factors.

AI Center - Managing ML Skills (uipath.com)

NEW QUESTION: 4

While training a UiPath Communications Mining model, the Search feature was used to pin a certain label on a few communications. After retraining, the new model version starts to predict the tagged label but infrequently and with low confidence.

According to best practices, what would be the correct next step to improve the model's predictions for the label, in the "Explore" phase of training?

- A. Use the "Rebalance" training mode to pin the label to more communications.
- B. Use the "Teach" training mode to pin the label to more communications.
- C. Use the "Low confidence" training mode to pin the label to more communications.
- D. Use the "Search" feature to pin the label to more communications.

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

According to the UiPath documentation, the 'Teach' training mode is used to improve the model's predictions for a specific label by pinning it to more communications that match the label's criteria. This helps the model learn from more examples and increase its confidence and accuracy. The 'Teach' mode also allows you to unpin the label from communications that do not match it, which helps the model avoid false positives. The other training modes are not as effective for this purpose, as they either focus on different aspects of the model performance or do not provide enough feedback to the model.

References:

Model training and labelling best practice

Overview of the model training process

Model Training FAQs

NEW QUESTION: 5

What is Document Understanding?

- A. A solution for combining different approaches to extract information from workflows.
- B. A solution that offers the ability to digitize, extract, validate, and train data from documents.
- C. A solution for the processing of Excel files for extracting data tables.
- D. A solution for combining different approaches to extract entities from Word documents such as contracts and agreements.

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

Document Understanding is a tool that helps you create and manage documents for your automation scenarios in the UiPath ecosystem. It allows you to process and extract data from multiple document types in an open, extensible, and versatile environment. The framework consists of components such as Taxonomy, Digitization, Data Extraction, Data Extraction Validation, Data Extraction Training, and Data Consumption, and enables you to customize and train your own algorithms¹².

References: Intelligent Document Processing - Document Understanding | UiPath, Document Understanding - Introduction - UiPath Documentation Portal

NEW QUESTION: 6

When should a UiPath Communications Mining taxonomy be imported?

- A. Before starting the model training.
- B. When new labels must be added to the taxonomy.
- C. As part of the "Increase coverage" phase.
- D. After pruning and reorganizing a taxonomy.

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

In UiPath Communications Mining, importing a taxonomy should be done before starting model training. The taxonomy, which includes labels and categories, defines how the data will be classified and structured during the training process. It is essential to have a well-defined taxonomy to ensure accurate predictions and classifications. Importing the taxonomy before training allows the model to learn from it, enhancing its performance. Changes to the taxonomy can be made later, but the initial import is crucial at the start of the training phase to guide the model effectively.

(Source: UiPath Docs on Communications Mining)

NEW QUESTION: 7

What is the correct order to Configure Extractor Wizard?

Instructions: Drag the Description found on the left and drop on the correct Step found on the right.

Description	Order of Steps
Place one or more extractors.	Step 1
Add a Data Extraction Scope activity to the workflow.	Step 2
Click on the Configure Extractors button.	Step 3
Select the checkboxes next to each field for the extractor type that should be activated.	Step 4
Get capabilities (if needed).	Step 5
Click on the Save button.	Step 6

Answer:

Description	Order of Steps
Place one or more extractors.	Step 1
Add a Data Extraction Scope activity to the workflow.	Step 2
Click on the Configure Extractors button.	Step 3
Select the checkboxes next to each field for the extractor type that should be activated.	Step 4
Get capabilities (if needed).	Step 5
Click on the Save button.	Step 6

Explanation:

Here is the correct order to configure the Extractor Wizard in UiPath Document Understanding:

Step 1: Add a Data Extraction Scope activity to the workflow.
Step 2: Place one or more extractors.
Step 3:

Click on the Configure Extractors button. Step 4: Select the checkboxes next to each field for the extractor type that should be activated. Step 5: Get capabilities (if needed). Step 6: Click on the Save button.

This sequence ensures that the Extractor Wizard is correctly configured to work with the Document Understanding workflow.

NEW QUESTION: 8

What are all the types of ML (Machine Learning) models supported by AI Center?

- A. Out-of-the-box models from UiPath and UiPath technology partners.
- B. Out-of-the-box models from UiPath technology partners and custom models.
- C. Out-of-the-box models from UiPath, UiPath technology partners, open-source models from the community, and custom models.
- D. Out-of-the-box models from UiPath technology partners, and open-source models from the community.

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

In UiPath AI Center, the platform supports several types of machine learning (ML) models, including:

Out-of-the-box models from UiPath: Pre-built models designed for common automation tasks.

Models from UiPath technology partners: External models developed by UiPath's partners.

Open-source models: Community-contributed models that can be used and adapted for various use cases.

Custom models: Models that users build and train specifically for their projects using their datasets.

This flexibility in model support ensures that organizations can leverage a wide range of machine learning capabilities to suit different automation needs.

For more details, refer to:

UiPath AI Center Documentation: AI Center Models

Machine Learning Model Types: Types of Models in UiPath AI Center

NEW QUESTION: 9

Which Reports tab do we use to get a high-level overview and statistics on all labels in UiPath Communications Mining?

- A. Dashboards.
- B. Label Summary.
- C. Trends.
- D. Segments.

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

In UiPath Communications Mining, the Label Summary report tab provides a high-level overview and statistics on all labels. This summary gives insights into the performance and distribution of labels across the dataset, helping users understand how well the model is classifying different types of communications.

For more details, refer to:

UiPath Communications Mining Documentation: Label Summary Reports

NEW QUESTION: 10

Which of the following scenarios is a good candidate for using Document Understanding Cloud APIs with synchronous calls?

- A.** You need to process documents larger than five pages.
- B.** You need real-time interaction and you are only processing documents with maximum five pages.
- C.** You need to handle multiple operations simultaneously, allowing for concurrent processing and avoiding idle time.
- D.** You have a large dataset that needs long-term processing.

Answer: ([SHOW ANSWER](#))

The synchronous API calls for UiPath Document Understanding Cloud are ideal when you need real-time interaction and fast feedback, such as when processing small documents (with up to five pages). This approach is beneficial in use cases requiring low latency, but it is less suitable for large datasets or documents due to limitations in speed and page count

NEW QUESTION: 11

Which of the following is a best practice when choosing a UiPath ML (Machine Learning) Extractor?

- A.** The popularity of the ML Extractor among other UiPath users should be the primary factor when choosing a UiPath ML Extractor. Opt for the ML Extractor that has the highest number of downloads or positive reviews.
- B.** Consider the document types, language, and data quality when choosing an ML Extractor. It is important to select one that is specifically trained or optimized for the document types being processed.

It is also important to take into account the quality and diversity of the training data used to train the ML Extractor to ensure accurate and reliable extraction results.

- C.** The cost of the ML Extractor should be the main consideration when choosing an ML Extractor. Select the ML Extractor that offers the lowest price, regardless of its performance or suitability for the specific document understanding needs.
- D.** The size of the ML Extractor is the most important factor to consider when choosing an ML Extractor. Bigger models always perform better and provide more accurate extraction results because the development team invested time and effort into creating the algorithm, which in turn will result in better performance for the trained model.

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

The ML Extractor is a data extraction tool that uses machine learning models provided by UiPath to identify and extract data from documents. The ML Extractor can work with predefined document types, such as invoices, receipts, purchase orders, and utility bills, or with custom

document types that are trained using the Data Manager and the Machine Learning Classifier Trainer¹².

According to the best practice, the ML Extractor should be chosen based on the document types, language, and data quality of the documents being processed. It is important to select an ML Extractor that is specifically trained or optimized for the document types that are relevant for the use case, as different document types may have different layouts, fields, and formats. It is also important to take into account the language of the documents, as some ML Extractors may support only certain languages or require specific language settings. Moreover, it is important to consider the quality and diversity of the training data used to train the ML Extractor, as this may affect the accuracy and reliability of the extraction results. The training data should be representative of the real-world data, and should cover various scenarios, variations, and exceptions³.

References: 1: Machine Learning Extractor - UiPath Activities 2: Machine Learning Classifier Trainer - UiPath Document Understanding 3: Data Extraction - UiPath Document Understanding

NEW QUESTION: 12

What does Data Extraction do?

- A. Identifies and extracts specific information that should be processed.
- B. Applies rules for validating that the information extracted from a document is correct.
- C. Digitizes the document that should be processed.
- D. Identifies words and their coordinates from images and PDFs.

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

Data Extraction identifies and extracts specific information from documents to be processed by automations.

This is a key part of UiPath's Document Understanding Framework, enabling bots to process structured and unstructured documents effectively.

Reference: UiPath Document Understanding - Data Extraction

NEW QUESTION: 13

Which filter option should be used for the For Each File in Folder activity to iterate between all the Microsoft Word documents in a local folder?

- A. *.doc*
- B. *.doc, *.docx
- C. *.doc
- D. Microsoft Word

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

Comprehensive and Detailed Explanation From Exact Extract:

The correct syntax for filtering both .doc and .docx Word documents is to use the wildcard *.doc*.

This matches:

- *.doc
- *.docx

- * Any Word file variant starting with .doc
- * UiPath Studio uses .NET wildcards for file filters in activities like "For Each File in Folder".
- * UiPath Documentation Reference: For Each File in Folder - UiPath Docs

NEW QUESTION: 14

Which of the following functionalities does UiPath Assistant provide?

- A. Analyzing processes to determine optimal automation solutions.
- B. Developing automation workflows in UiPath Studio.
- C. Running, managing, and organizing automation workflows on the user's machine.
- D. Scheduling and monitoring robot processes in Orchestrator.

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

Reference: UiPath Assistant

NEW QUESTION: 15

Which of the below is the correct definition of "recall" in UiPath Communications Mining?

- A. For a given concept what % of cases will the model incorrectly predict.
- B. For a given concept, what % of cases will the model not detect.
- C. For a given concept what % of cases will the model correctly predict.
- D. For a given concept, what % of cases will the model detect.

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

Recall is a metric that measures the proportion of all possible true positives that the model was able to identify for a given concept¹. A true positive is a case where the model correctly predicts the presence of a concept in the data. Recall is calculated as the ratio of true positives to the sum of true positives and false negatives, where a false negative is a case where the model fails to predict the presence of a concept in the data. Recall can be interpreted as the sensitivity or completeness of the model for a given concept². For example, if there are 100 verbatims that should have been labelled as 'Request for information', and the model detects 80 of them, then the recall for this concept is 80% ($80 / (80 + 20)$). A high recall means that the model is good at finding all the relevant cases for a concept, while a low recall means that the model misses many of them.

References: 1: Recall 2: Precision and Recall

NEW QUESTION: 16

What is a function of unattended robots?

- A. Unattended robots can run independently without human interaction.
- B. Unattended robots can only work if they are not connected to Orchestrator.
- C. Unattended robots must be triggered manually.
- D. Unattended robots only run on a workstation operated by a human.

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

Comprehensive and Detailed Explanation From Exact Extract:

Unattended robots are designed to run in the background without human supervision. They can:

- * Be triggered via Orchestrator (manually, on a schedule, or via a queue)
- * Run on virtual machines or remote servers
- * Handle end-to-end automation tasks autonomously

These robots are ideal for back-office processes and scalable enterprise deployments.

* UiPath Documentation Reference: Unattended Robots - UiPath Docs

Valid UiPath-SAIv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIv1 exam dumps**, the BraindumpsPass.com UiPath-SAIv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIv1 dumps with Test Engine here: <https://www.braindumpsPass.com/UiPath/UiPath-SAIv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 17

What will be the behavior of the process if, during design time, the property ValidateUnconnectedNodes is set to True on a flowchart and a Log Message activity from this flowchart is not connected to any other node?

- A. The flowchart undergoes successful validation with no errors displayed, but an exception will occur during the workflow's runtime.
- B. The flowchart displays an error indicating that some activities are not connected to the others within the flowchart.
- C. The flowchart will display an error indicating the presence of unconnected activities, but only when validated using the Workflow Analyzer.
- D. A warning message is displayed in the Output section of Studio, indicating the presence of unconnected activities.

Answer: B (LEAVE A REPLY)

Reference: UiPath Flowchart Validation

NEW QUESTION: 18

Which feature should be used to inspect the available versions for activities employed within a workflow?

- A. Activities Panel
- B. Variables Panel
- C. Manage Packages
- D. Outline Panel

Answer: C (LEAVE A REPLY)

Reference: UiPath Manage Packages

NEW QUESTION: 19

What does the Export stage of the Document Understanding Framework do?

- A.** Converts the result of extraction to a dataset or to a customized format.
- B.** Allows a human to validate and correct the extracted data.
- C.** Classifies the document as one of the predefined document types.
- D.** Extracts the text out of the image document using OCR (Optical Character Recognition).

Answer: ([SHOW ANSWER](#))

According to the UiPath documentation portal¹, the Export stage of the Document Understanding Framework is the final stage of the document processing workflow, where the extracted data is converted to a dataset or to a customized format, such as Excel, JSON, or XML. The Export stage enables you to easily export data for training ML models, using the Export files dialog box in Data Manager or Document Manager. The Export stage also allows you to export the schema of the fields and their configurations, which can be imported into a different session. The Export stage supports different export options, such as current search results, all labeled, schema, or all. The Export stage also provides validation rules, such as requiring at least 10 labeled documents or pages for each field, and at least one document for each classification option¹. Therefore, option A is the correct answer, as it describes the main function and benefit of the Export stage. Option B is incorrect, as it refers to the Data Validation stage, which is the previous stage of the document processing workflow, where a human can review and correct the extracted data using the Validation Station or the Present Validation Station activities². Option C is incorrect, as it refers to the Classification stage, which is the second stage of the document processing workflow, where the document is classified as one of the predefined document types using the Classify Document Scope activity and a classifier of choice³. Option D is incorrect, as it refers to the Digitization stage, which is the first stage of the document processing workflow, where the text is extracted out of the image document using OCR (Optical Character Recognition) using the Digitize Document activity and an OCR engine of choice.

References: 1 Document Understanding - Export Documents 2 Document Understanding - Data Validation 3 Document Understanding - Classification Document Understanding - Digitization

NEW QUESTION: 20

Which of the following use cases is best suited for tone analysis instead of label sentiment analysis in UiPath Communications Mining?

- A.** Analyzing customer complaints in a B2C organization.
- B.** Analyzing employee engagement survey responses.
- C.** Analyzing customer satisfaction survey responses.
- D.** Monitoring "Quality of Service" in an operations-focused shared mailbox in a B2B organization.

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

Tone analysis is better suited for monitoring situations like "Quality of Service" in shared mailboxes, where the focus is on evaluating emotional tone in communications that may not always have clear-cut positive or negative sentiments. This contrasts with label sentiment analysis, which is better for datasets with explicit feedback (e.g., customer satisfaction surveys).

In operations-focused environments, tone analysis provides more nuanced insights into service quality

NEW QUESTION: 21

Which activity is used to validate and correct automatic classification outputs?

- A. Present Validation Station activity
- B. Digitize Document activity
- C. Classify Document activity
- D. Present Classification Station activity

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

Comprehensive and Detailed Explanation From Exact Extract:

- * The Present Classification Station activity is specifically used to review and correct the classification results performed by the "Classify Document" activity.
- * It allows a human to validate which document type has been assigned before data extraction begins.
- * UiPath Documentation Reference: Classification Station - Document Understanding

NEW QUESTION: 22

What is a characteristic of an Orchestrator Asset?

- A. All values from any Asset type are encrypted.
- B. Asset types can be modified after the asset is created.
- C. Asset values can be specified for each user.
- D. Assets can store complete DataTable variables.

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

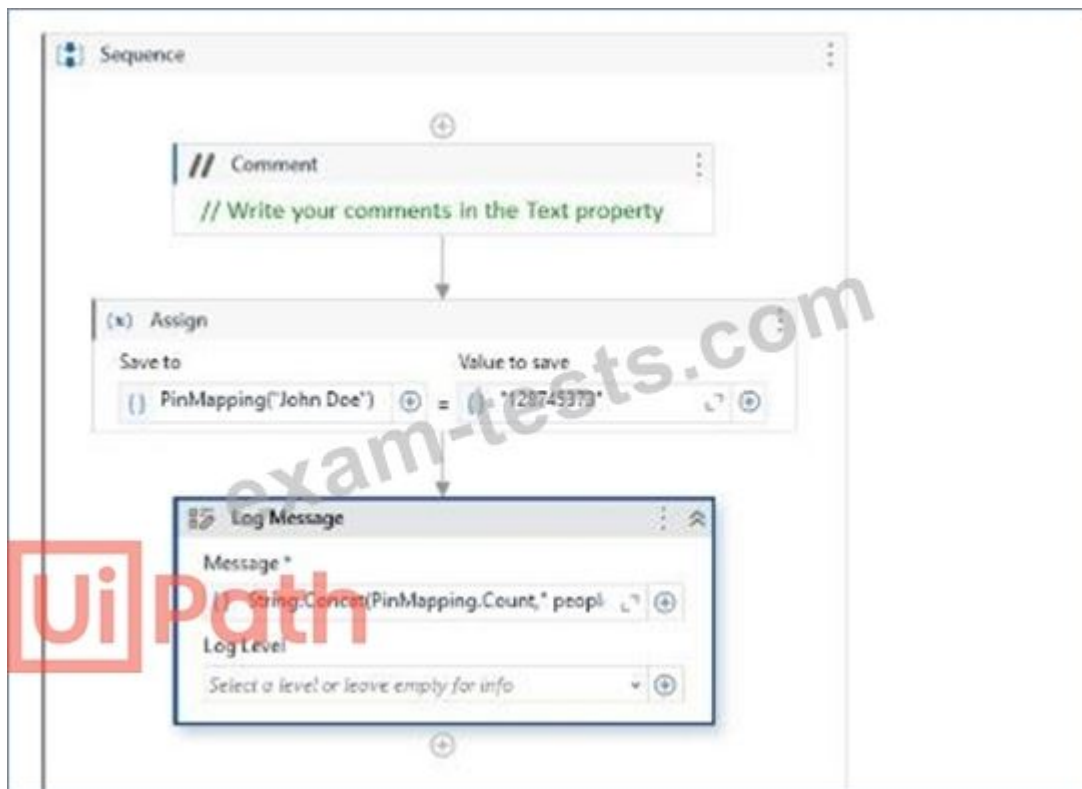
Comprehensive and Detailed Explanation From Exact Extract:

UiPath Orchestrator assets support per-user values, allowing different users (robots) to retrieve customized values for the same asset key. This is commonly used for credentials, URLs, or configurations that vary per user or environment.

- * While most asset types (text, integer, credential) support secure handling, only Credential assets are encrypted by default.
- * UiPath Documentation Reference: Assets in Orchestrator

NEW QUESTION: 23

A developer needs to create a process for the Human Resources team. During the development, they try to run the workflow containing the following dictionary variable:



What is the possible cause of the error?

- A. The assign's set value syntax should be `PinMapping["John Doe"]`.
- B. The "John Doe" key was not present in the dictionary.
- C. The Dictionary was not initialized.
- D. The assign's set value syntax should be `PinMapping<"John Doe">`.

Answer: C (LEAVE A REPLY)

The most likely cause of the error is that the dictionary was not initialized. In UiPath, a dictionary must be initialized before assigning values to its keys. If you attempt to add a key-value pair to a dictionary that has not been initialized, you will encounter a runtime error.

The correct initialization can be done as follows:

```
PinMapping = New Dictionary(Of String, String)
```

Explanation of Other Options:

A: `PinMapping["John Doe"]`: This is the correct syntax for accessing or assigning a value to a dictionary key.

There is no issue with this syntax.

B: The "John Doe" key not present: While it is true that the key might not exist, this specific error would occur only during a Get operation, not an assignment.

D: `PinMapping<"John Doe">`: This is incorrect syntax for working with dictionaries in UiPath.

Reference: UiPath Variables and Data Types - Dictionaries

NEW QUESTION: 24

Under what condition can a dataset be deleted in UiPath AI Center?

- A. If it has never been used in a pipeline.
- B. If it is not being used in any active pipeline.

- C. If it is not associated with any ML skills.
- D. If it has not been modified within the last 24 hours.

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

In UiPath AI Center, datasets can only be deleted if they are not being utilized in any active pipelines. This ensures that datasets tied to running or scheduled operations are not inadvertently deleted, which could disrupt ongoing processes. Once the dataset is confirmed to be unused in any active pipeline, it can safely be removed

NEW QUESTION: 25

What are the UiPath Action Center action statuses?

- A. Unassigned, Assigned, Modified, Completed
- B. Unassigned, Assigned, Completed
- C. Unassigned, Pending, Edited, Completed
- D. Unassigned, Pending, Completed

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

The valid Action Center statuses are:

Unassigned: The action is not assigned to any user.

Pending: The action is assigned and awaiting user response.

Reference: UiPath Action Center

NEW QUESTION: 26

Why should using the Search feature be limited when training in UiPath Communications Mining?

- A. It could increase model coverage.
- B. It could increase model bias.
- C. It could decrease model coverage.
- D. It could decrease model bias.

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

In UiPath Communications Mining, over-reliance on the Search feature during the training process can lead to an increase in model bias. This happens because using search-based filtering to identify and label examples might not represent the full diversity of the data. The model could be trained on a skewed subset of the data, causing it to favor certain patterns or keywords, and thus biasing the model towards specific types of data rather than learning to generalize effectively across all data.

Limiting the use of search ensures that the training process considers a broader and more representative sample of the data, which reduces the risk of introducing bias into the model and helps it generalize better to new, unseen communications.

For more details, refer to:

UiPath Communications Mining Documentation: Model Training and Avoiding Bias

NEW QUESTION: 27

The "Train" stage from Document Understanding Framework usually comes after?

- A. Digitization.
- B. Classification.
- C. Validation.
- D. Extraction.

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

Reference: UiPath Document Understanding Workflow

NEW QUESTION: 28

What is the purpose of field rules in Taxonomy Manager?

- A. To generate custom data fields for extraction.
- B. To optimize extraction results and automatically validate them.
- C. To create complex data manipulation operations.
- D. To handle exceptions during the extraction process.

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

In UiPath's Taxonomy Manager, field rules are used to optimize extraction results and perform automatic validation of extracted data. These rules can be applied to specific fields within the document understanding process to ensure that extracted data meets certain predefined conditions or constraints, such as formats, patterns, or value ranges. This allows for higher accuracy in the extraction process and reduces errors, as invalid data can be flagged or corrected based on these rules.

Field rules help automate the validation of data by setting criteria that extracted data must meet, thus enhancing the extraction quality and ensuring that only valid and structured data is processed further in the automation workflow.

For more details, refer to:

UiPath Taxonomy Manager Documentation: Field Rules in UiPath

UiPath Document Understanding Framework: Taxonomy Manager Overview

NEW QUESTION: 29

What is the correct execution order of the Document Understanding template stages?

Instructions: Drag the stages found on the "Left" and drop them on the "Right" in the correct order.

Template Stages	Order of Execution
☐ Taxonomy	First 0
☐ Classify	Second 0
☐ Digitize	Third 0
☐ Extract	Fourth 0
☐ Extraction Validation	Fifth 0
☐ Export	Sixth 0

Answer:

Template Stages	Order of Execution
☐ Taxonomy	First 0
☐ Classify	Second 0
☐ Digitize	Third 0
☐ Extract	Fourth 0
☐ Extraction Validation	Fifth 0
☐ Export	Sixth 0

Explanation:

The correct execution order of the Document Understanding template stages is:

Taxonomy

Digitize

Classify

Extract

Extraction Validation

Export

Comprehensive and Detailed Explanation: The Document Understanding template stages are based on a document processing flowchart that follows these steps:

First, you need to define the Taxonomy of the document types and fields that you want to process and extract information from. This is done using the Taxonomy Manager in UiPath Studio¹.

Next, you need to Digitize the input documents, which can be in various formats such as PDF, image, or text. This is done using the Digitize Document activity, which converts the documents into a machine- readable format and performs OCR if needed².

Then, you need to Classify the digitized documents into the predefined document types in your taxonomy. This is done using the Classify Document Scope activity, which can use various classifiers such as Keyword Based Classifier, Machine Learning Classifier, or Intelligent Form Extractor³.

After that, you need to Extract the relevant information from the classified documents based on the fields in your taxonomy. This is done using the Data Extraction Scope activity, which can use various extractors such as Regex Based Extractor, Machine Learning Extractor, or Form Extractor.

Next, you need to perform Extraction Validation to review and correct the extracted information, either manually or automatically. This is done using the Present Validation Station activity, which can use either the Validation Station or the Action Center for human-in-the-loop validation.

Finally, you need to Export the validated information to the desired output location, such as a file, a database, or a queue. This is done using the Export Extraction Results activity, which can use various exporters such as Excel Exporter, CSV Exporter, or Queue Item Exporter.

References:

UiPath Studio - Taxonomy Manager

UiPath Activities - Digitize Document

UiPath Activities - Classify Document Scope

NEW QUESTION: 30

Why is the Shuffle training mode important in the "Explore" phase while working with UiPath Communications Mining?

- A. Because it helps create a balanced model by focusing on outlier labels with sufficient training examples but low confidence.
- B. Because it helps create a balanced model free from labelling bias by introducing random variation into the training data.
- C. Because it creates a high precision model, focusing on labels with low precision and a low number of examples.
- D. Because it creates a new model version with shuffled labels and general fields and lets the user explore to understand the state of the model

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

The Shuffle training mode in the "Explore" phase of UiPath Communications Mining helps to introduce randomness into the training data. This prevents labeling bias by ensuring that the model does not overly focus on certain labels or examples, leading to a more balanced and generalized model. By shuffling the data, the model is exposed to a diverse range of examples, which enhances its ability to make accurate predictions across different labels and datasets.

(Source: UiPath Communications Mining documentation)

NEW QUESTION: 31

What is the primary function of the Wait for Classification Validation Task and Resume activity In UiPath's Document Understanding Framework?

- A.** It prioritizes actions in Action Center based on document classification results, optimizing task management and allocation according to the importance of document classifications.
- B.** It initiates the classification process for documents across different platforms, ensuring consistent and accurate document organization and categorization.
- C.** It automatically validates classified data without human intervention, expediting document processing by removing the need for manual review and correction.
- D.** It suspends the workflow until a specified document validation action is completed, ensuring human review and correction.

Answer: ([SHOW ANSWER](#))

The "Wait for Classification Validation Task and Resume" activity in UiPath's Document Understanding Framework is primarily used to halt or suspend the workflow until a specified document classification validation task is completed by a human. This activity is part of the broader workflow to ensure that when automatic classification of documents cannot be confidently achieved, a human-in-the-loop (HITL) approach is followed to validate or correct classifications. Once the validation is performed in UiPath's Action Center by a human, the workflow is resumed, ensuring the proper handling of documents that require review and correction.

This is aligned with the design of the Action Center, which is integrated into UiPath's Document Understanding Framework. When dealing with document classification or extraction confidence issues, manual human validation tasks are often required, which is what this activity manages. It facilitates human oversight, preventing the automation from proceeding with potentially incorrect classifications.

Reference from UiPath documentation:

UiPath Action Center explains how humans are involved in validation tasks to handle cases where classification or extraction needs manual review.

Wait for Task and Resume Activity in UiPath Documentation explains how it waits for a task (such as document validation) to be completed in the Action Center before resuming the workflow.

For more details, you can consult the official UiPath documents:

UiPath Document Understanding Framework

Wait for Classification Validation Task and Resume

This functionality ensures that incorrect data processing due to automation can be caught and rectified by a human, improving accuracy in document handling workflows.

Valid UiPath-SAIv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIv1 exam dumps**, the BraindumpsPass.com UiPath-SAIv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIv1 dumps with Test Engine here: <https://www.braindumpsPASS.com/UiPath/UiPath-SAIv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 32

Which of the following is a best practice when choosing a UiPath ML (Machine Learning) Extractor?

- A.** The popularity of the ML Extractor among other UiPath users should be the primary factor. Opt for the ML Extractor that has the highest number of downloads or positive reviews.
- B.** Consider the document types, language, and data quality. It is important to select one that is specifically trained or optimized for the document types being processed. It is also important to take into account the quality and diversity of the training data used to train the ML Extractor to ensure accurate and reliable extraction results.
- C.** The size of the ML Extractor is the most important factor to consider. Bigger models always perform better and provide more accurate extraction results because the development team invested time and effort into creating the algorithm, which in turn will result in better performance for the trained model.
- D.** The cost of the ML Extractor should be the main consideration. Select the ML Extractor that offers the lowest price, regardless of its performance or suitability for the specific document understanding needs.

Answer: ([SHOW ANSWER](#))

The best practice is to select an ML Extractor based on document types, language, and data quality. Choosing a model specifically optimized for the type of document being processed ensures higher accuracy and reliability. The quality and diversity of the training data used to develop the model play a significant role in its performance.

Reference: UiPath ML Extractors

NEW QUESTION: 33

What are the main components of a digital business process?

- A.** Inputs, Process flows, Assignees, Outputs
- B.** Inputs, Process flows, Outputs
- C.** Inputs, Source applications, Assignees
- D.** Inputs, Process flows, Source applications, Outputs

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

Comprehensive and Detailed Explanation From Exact Extract:

In the context of automation and digital workflows, a digital business process consists of:

- * Inputs: The raw data required
- * Process flows: The sequence of tasks or actions
- * Outputs: The final results generated

Assignees and source applications are supporting components, but not universally required in every digital process.

* UiPath Academy Reference: Automation Hub Training # Defining Business Processes
UiPath Process Mining - Fundamentals

NEW QUESTION: 34

What is the definition of Deep Learning?

- A.** A sub-field of artificial intelligence that enables systems to learn from data. Systems learn from previous experience and information to deduce and predict future information. To do this they use algorithms that learn to perform a specific task without being explicitly programmed.
- B.** The theory and development of computer systems that are able to perform tasks that normally require human intelligence and decision making.
- C.** A field of artificial intelligence that enables computers to gain high-level understanding from digital images or videos. If AI is the brain, then this is the eye that enables the computer to observe and understand. It works the same as the human eye.
- D.** An area of machine learning concerned with artificial neural networks. These are a series of algorithms that aim to recognize relationships in a set of data through a process that mimics biological neural networks.

Answer: D (LEAVE A REPLY)

Deep learning is a subset of machine learning that uses multiple layers of artificial neural networks to learn from data and perform complex tasks. The term "deep" refers to the number of layers in the network, which can range from a few to hundreds or even thousands. Each layer consists of a set of nodes that perform mathematical operations on the input data and pass the output to the next layer. The network learns by adjusting the weights of the connections between the nodes based on the feedback from the desired output.

Deep learning can handle various types of data, such as images, text, speech, or video, and can automatically extract features and patterns from them without human intervention. Deep learning is behind many applications of artificial intelligence, such as computer vision, natural language processing, speech recognition, and generative models¹²³.

References: 1: What is Deep Learning? | IBM 2: What Is Deep Learning? Definition, Examples, and Careers | Coursera 3: Deep learning - Wikipedia

NEW QUESTION: 35

How are UiPath RPA and AI Center used for process improvement?

- A.** UiPath RPA and AI Center have overlapping functionalities, so they are not designed to work together in process improvement. Each tool serves different purposes and is used independently to achieve automation and AI capabilities, respectively.
- B.** UiPath RPA and AI Center work together in process improvement by utilizing RPA bots to execute models deployed in Action Center that can be monitored in Insights. The ML models process data and generate insights, which are then used by the RPA bots to make informed decisions and carry out automated actions.
- C.** UiPath RPA and AI Center are independent tools that do not work together in process improvement.

RPA focuses on task automation, while AI Center is primarily used for training and deploying machine learning models.

- D.** UiPath RPA and AI Center work together in process improvement by leveraging RPA capabilities to automate repetitive tasks and integrating AI Center to enhance decision-making

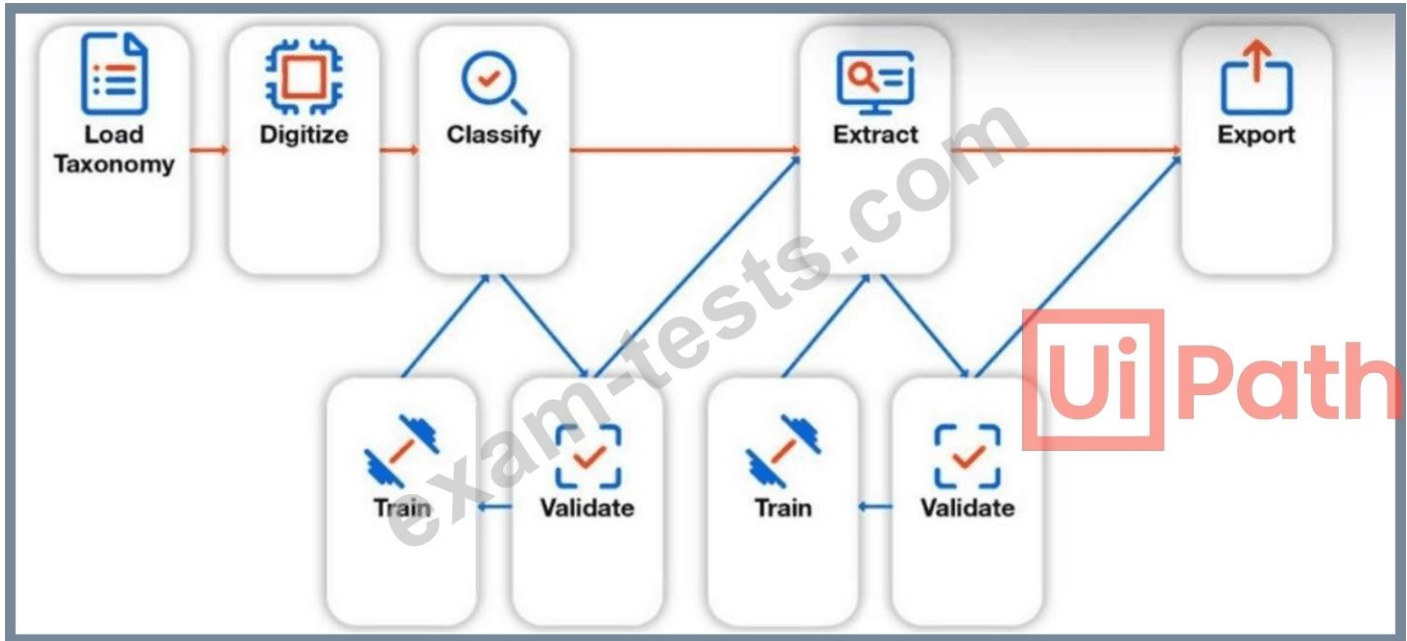
and cognitive abilities. RPA bots can collect data, perform actions, and feed information to ML models in AI Center, which in turn can analyze the data and make predictions.

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

Reference: UiPath RPA and AI Center Integration

NEW QUESTION: 36

Which component from the image answers the question "Is the extracted information correct?"



- A. Classify
- B. Validate (Extract)
- C. Train (Extract)
- D. Extract

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

Comprehensive and Detailed Explanation From Exact Extract:

The Validate (Extract) component is responsible for allowing a human to review, correct, or approve the data extracted by the robot. It typically uses the Present Validation Station activity in the workflow.

- * This step ensures data quality before submission or processing.
- * The extracted data is compared against what a human user confirms as accurate.
- * UiPath Documentation Reference: Present Validation Station - UiPath Docs

NEW QUESTION: 37

What does the darker shading of a label prediction represent in Explore in UiPath Communications Mining?

- A. An incorrect prediction.
- B. A lower confidence score.
- C. Multiple label predictions.
- D. A higher confidence score.

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

In UiPath Communications Mining, the shading of label predictions in the Explore section indicates confidence levels. Darker shading represents a higher confidence score, meaning the model is more certain of the prediction made for that particular label

NEW QUESTION: 38

What are the out-of-the-box packages types available in AI Center?

- A. Pre-trained, fine-tunable, and reviewed.
- B. Pre-trained, custom training, and fine-tunable.
- C. Custom training, fine-tunable, and reviewed.
- D. Pre-trained, custom training, and reviewed.

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

UiPath AI Center offers three primary package types: pre-trained, custom training, and fine-tunable models.

These are essential for various use cases, from leveraging existing models to training custom ones based on specific data

NEW QUESTION: 39

Which of the following options is accepted as a Column field name in Document Manager?

- A. first_n@me
- B. first name
- C. f1rst-name
- D. First_name123

Answer: ([SHOW ANSWER](#))

According to the UiPath documentation, the field name for a column field in Document Manager does not accept uppercase letters. It can only contain lowercase letters, numbers, underscore _ and dash -. Therefore, the only option that meets these criteria is D. First_name123. The other options are invalid because they either contain uppercase letters, spaces, or @ symbols, which are not allowed.

References: 1: Document Understanding - Create and Configure Fields 2: Document Understanding - Create & Configure Fields

NEW QUESTION: 40

When is it recommended to use an ML (Machine Learning) model solution?

- A. For structured or semi-structured documents in which layouts of different document providers vary greatly.
- B. For fixed-layout documents for data extraction, including handwriting recognition and signature detection.
- C. For simple use cases in which data is always found in a strict, predictable format and context.
- D. When documents have little to no variation in the document layouts.

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

Comprehensive and Detailed Explanation From Exact Extract:

ML models are best suited for complex, semi-structured or unstructured documents where rules or templates can't handle layout variability. For example, invoices from many vendors with different formats are ideal for ML.

* In contrast, fixed-format documents are better handled with regex, form, or template-based extractors.

* UiPath Documentation Reference: [Choosing the Right Extractor - Document Understanding](#)

NEW QUESTION: 41

What does adding missed labels help improve in UiPath Communications Mining?

- A. Label bias warnings.
- B. Increases data security.
- C. Increases the taxonomy coverage.
- D. Label precision and recall.

Answer: ([SHOW ANSWER](#))

Adding missed labels helps improve the label precision and recall in UiPath Communications Mining.

Precision is the percentage of correctly labeled verbatims out of all the verbatims that have the label applied, while recall is the percentage of correctly labeled verbatims out of all the verbatims that should have the label applied. By adding missed labels, you are increasing the recall of the label, as you are reducing the number of false negatives (verbatims that should have the label but do not). This also improves the precision of the label, as you are reducing the noise in the data and making the label more informative and consistent. Adding missed labels is one of the recommended actions that the platform suggests to improve the model rating and performance of the labels.

References: [Communications Mining - Training using 'Check label' and 'Missed label'](#),
[Communications Mining - Model Rating](#)

NEW QUESTION: 42

When creating a training dataset, what is the recommended number of samples for the Classification fields?

- A. 5-10 document samples from each class.
- B. 10-20 document samples from each class.
- C. 20-50 document samples from each class.
- D. 50-200 document samples from each class.

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

According to the UiPath documentation, the recommended number of samples for the classification fields depends on the number of document types and layouts that you want to classify. The more document types and layouts you have, the more samples you need to cover the diversity of your data. However, a general guideline is to have at least 20-50 document

samples from each class, as this would provide enough data for the classifiers to learn from¹². A large number of samples per layout is not mandatory, as the classifiers can generalize from other layouts as well³.

References: 1: Document Classification Training Overview 2: Document Classification Training Related Activities 3: Training High Performing Models

NEW QUESTION: 43

Which of the following statements is true regarding reviewing and applying entities in UiPath Communications Mining?

- A. A single entity value can be split across multiple paragraphs.
- B. If the entity value is correctly predicted, but the entity type is wrong, it cannot be changed.
- C. All of the entities within a paragraph should be reviewed.
- D. All of the entities in a communication must be reviewed.

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

According to the UiPath Communications Mining documentation, reviewing and applying entities is a crucial step for improving the accuracy and performance of the entity extraction models.

When reviewing entities, users should check all of the predicted entities within a paragraph, as well as any missing or incorrect ones.

Users can accept, reject, edit, or create entities using the platform's interface or keyboard shortcuts. Users can also change the entity type if the value is correct but the type is wrong. Reviewing and applying entities helps the platform learn from the user feedback and refine its predictions over time. It also helps users assess the automation potential and benefit of the communications data.

References:

Communications Mining - Reviewing and applying entities

Communications Mining - Improving entity performance

NEW QUESTION: 44

What information should be filled in when adding an entity label for the OOB (Out Of the Box) labeling template?

- A. Name. Data Type. Attribute name, and Color.
- B. Name, Data Type. Attribute name. Shortcut, and Color.
- C. Name, Shortcut, and Color.
- D. Name. Input to be labeled. Attribute name. Shortcut, and Color.

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

The OOB labeling template is a predefined template that you can use to label your text data for entity recognition models. The template comes with some preset labels and text components, but you can also add your own labels using the General UI or the Advanced Editor. When you add an entity label, you need to fill in the following information:

Name: the name of the new label. This is how the label will appear in the labeling tool and in the exported data.

Input to be labeled: the text component that you want to label. You can choose from the existing text components in the template, such as Date, From, To, CC, and Text, or you can add your own text components using the Advanced Editor. The text component determines the scope of the text that can be labeled with the entity label.

Attribute name: the name of the attribute that you want to extract from the text. You can use this to create attributes such as customer name, city name, telephone number, and so on. You can add more than one attribute for the same label by clicking on + Add new.

Shortcut: the hotkey that you want to assign to the label. You can use this to label the text faster by using the keyboard. Only single letters or digits are supported.

Color: the color that you want to assign to the label. You can use this to distinguish the label from the others visually.

References: AI Center - Managing Data Labels, Data Labeling for Text - Public Preview

NEW QUESTION: 45

What is the benefit of making an ML Skill public?

- A. It allows access from outside of the UiPath environment.
- B. It provides additional security measures for the ML Skill.
- C. It enables automatic updates and enhancements to the ML Skill without user intervention.
- D. It allows access from inside of the UiPath environment.

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

Making an ML Skill public in UiPath enables it to be accessed externally from the UiPath ecosystem. This can be beneficial if the ML Skill needs to be utilized in external applications, systems, or services beyond UiPath's automation environment. Public access expands the usability of the skill, allowing integration with other systems while maintaining security through managed endpoints.

NEW QUESTION: 46

Which role consumes ML Skills within customized workflows in Studio using the ML Skill activity from the UiPath.MLServices.Activities package?

- A. Data Scientist.
- B. Administrator.
- C. RPA Developer.
- D. Process Controller

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

According to the UiPath documentation portal¹, the RPA Developer is the role that consumes ML Skills within customized workflows in Studio using the ML Skill activity from the UiPath.MLServices.Activities package. The RPA Developer is responsible for designing, developing, testing, and deploying automation workflows using UiPath Studio and other UiPath products. The RPA Developer can use the ML Skill activity to retrieve and call all ML Skills available on the AI Center service and request them within the automation workflows. The ML Skill activity allows the RPA Developer to pass data to the input of the skill, test the skill, and receive the output of the skill as JSON response, status code, and headers². Therefore, option C

is the correct answer, as it describes the role and the activity that are related to consuming ML Skills in Studio. Option A is incorrect, as the Data Scientist is the role that creates and trains ML models using AI Center or other tools, and publishes them as ML Packages or OS Packages¹. Option B is incorrect, as the Administrator is the role that manages the AI Center service, such as configuring the infrastructure, setting up the permissions, and monitoring the usage and performance¹. Option D is incorrect, as the Process Controller is the role that deploys ML Packages or OS Packages as ML Skills, and manages the versions, the endpoints, and the API keys of the skills¹.

References: 1 AI Center - User Personas 2 Activities - ML Skill

Valid UiPath-SAIAv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIAv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIAv1 exam dumps**, the BraindumpsPass.com UiPath-SAIAv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIAv1 dumps with Test Engine here: <https://www.braindumpsPass.com/UiPath/UiPath-SAIAv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 47

What are all the ways to deploy AI Center?

- A.** In cloud availability, on-premises air-gapped, on-premises online, hybrid mode (AI Center on-premise + Orchestrator on-premise). and Automation Suite.
- B.** In cloud availability, on-premises air-gapped, on-premises online, and hybrid mode {cloud AI Center + Orchestrator on-premise}.
- C.** In cloud availability, on-premises, hybrid mode (AI Center on-premise + cloud Orchestrator), and Automation Suite.
- D.** In cloud availability, on-premises air-gapped, on-premises online, hybrid mode (cloud AI Center + Orchestrator on-premise). and Automation Suite.

Answer: A (LEAVE A REPLY)

UiPath AI Center can be deployed in multiple ways to meet different organizational needs and infrastructures.

The available deployment options include:

Cloud availability: Using UiPath's cloud services.

On-premises air-gapped: A fully isolated, offline environment for organizations with strict security requirements.

On-premises online: Deployed on-premise but with internet connectivity.

Hybrid mode: Combining on-premises AI Center with on-premises Orchestrator for flexibility.

Automation Suite: A comprehensive deployment of UiPath tools, including AI Center.

For more details, refer to:

NEW QUESTION: 48

How has UiPath improved the experience for handling failed queue transactions in the 2023.10 release?

- A. Enhancing the logging capabilities to provide more detailed insights into each transaction.
- B. Implementing a sophisticated automatic retry mechanism to handle failed transactions without manual intervention.
- C. Recording and storing failed queue transactions for debugging purposes.
- D. Introducing real-time error alerts to notify users immediately when a queue transaction fails.

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

In the UiPath 2023.10 release, a key improvement introduced was the ability to record and store failed queue transactions for debugging purposes. This feature helps users by capturing the failure details and making them accessible directly in Orchestrator. The recording of failed transactions is enabled at the process level, within the Video section of the Job recording settings. It is intended to facilitate easier troubleshooting of failed automations by providing insights into the moments leading to the failure, which can be reviewed later for debugging.

NEW QUESTION: 49

What function in the train.py file is responsible for persisting the trained model?

- A. evaluate(self, evaluation_directory)
- B. save(self)
- C. train(self, trainmg_directory)
- D. __init__(self)

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

In the context of a machine learning pipeline, specifically within the train.py file in UiPath AI models, the function responsible for persisting the trained model is typically named save(self). This function is invoked after the model has been trained, and it ensures that the model's state, weights, and configurations are saved to disk or another persistent storage medium. The saved model can then be reused for inference or further fine-tuning in future executions.

For more details, refer to:

UiPath AI Center Documentation: Training Models

Model Persistence: Saving Machine Learning Models in AI Center

NEW QUESTION: 50

What action should be performed to have the Taxonomy Manager button appear on the Ribbon, in the Wizard section?

- A. Install UiPath.IntelligentOCR.Activities package.
- B. Install a newer version of Studio.
- C. Install UiPath.Documentunderstanding.ML.Activities package.
- D. Nothing to do, it should be already there.

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

To enable the Taxonomy Manager button in the Ribbon, the UiPath.IntelligentOCR.Activities package must be installed. This package includes all activities related to Document Understanding, including the Taxonomy Manager.

Reference: UiPath Taxonomy Manager Setup

NEW QUESTION: 51

What fields are available when creating an AI Center project?

- A. Name, description, and labels.
- B. Name, description, and permissions.
- C. Name and description.
- D. Name and labels.

Answer: ([SHOW ANSWER](#))

When creating an AI Center project in UiPath, the fields available to input are the project's name and description. These fields allow you to clearly label and describe the purpose of the AI project within the UiPath platform. Permissions and labels can be managed separately after the project is created

NEW QUESTION: 52

To determine the number of characters scraped from a website in an "ExtractedText" String variable, excluding leading and trailing white-space characters, what should a developer use?

- A. ExtractedText.Trim.Chars
- B. ExtractedText.Length
- C. ExtractedText.Trim.Length
- D. ExtractedText.Chars

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

Comprehensive and Detailed Explanation From Exact Extract:

To get the character count excluding leading and trailing spaces, .Trim() is used to remove whitespace and .

Length provides the character count. So the correct expression is ExtractedText.Trim.Length.

* Trim: Removes all leading and trailing white-space characters.

* Length: Returns the number of characters in the string.

* UiPath Documentation Reference: String Manipulations in VB.NET - Microsoft Docs

* Also validated in UiPath Academy: Developer Foundation Course - String Manipulation Module

NEW QUESTION: 53

What is the difference between the Document Understanding Process and the Document Understanding Framework?

- A. The Document Understanding Framework contains the activities that can be used in a Library, while the Document Understanding Process is the template that can be found in Studio.

- B.** The Document Understanding Framework contains the activities that can be used in a Process, while the Document Understanding Process is the template that can be found in Studio.
- C.** The Document Understanding Process contains the activities that can be used in a Library, while the Document Understanding Framework is the template that can be found in Studio.
- D.** The Document Understanding Process contains the activities that can be used in a Process, while the Document Understanding Framework is the template that can be found in Studio.

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

According to the UiPath documentation portal¹, the Document Understanding Process is a fully functional UiPath Studio project template based on a document processing flowchart. It provides logging, exception handling, retry mechanisms, and all the methods that should be used in a Document Understanding workflow, out of the box. The Document Understanding Process is preconfigured with a series of basic document types in a taxonomy, a classifier configured to distinguish between these classes, and extractors to showcase how to use the Data Extraction capabilities of the framework. It is meant to be used as a best practice example that can be adapted to your needs while displaying how to configure each of its components¹. The Document Understanding Framework, on the other hand, is a set of activities that can be used to build custom document processing workflows. The framework facilitates the processing of incoming files, from file digitization to extracted data validation, all in an open, extensible, and versatile environment. The framework enables you to combine different approaches to extract information from multiple document types. The framework consists of several components, such as Taxonomy, Digitization, Classification, Data Extraction, Data Validation, and Data Consumption². Therefore, option D is the correct answer, as it describes the difference between the Document Understanding Process and the Document Understanding Framework.

References: 1 Document Understanding Process: Studio Template 2 Document Understanding - Introduction

NEW QUESTION: 54

What is the relationship between AI Center and UiPath Document Understanding?

- A.** AI Center is the infrastructure on top of which UiPath Document Understanding digitization runs.
- B.** Document Understanding is the infrastructure on which AI Center digitization runs.
- C.** AI Center is the infrastructure on top of which UiPath Document Understanding machine learning models run.
- D.** Document Understanding is the infrastructure on which AI Center machine learning models run.

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

Reference: UiPath AI Center and Document Understanding

NEW QUESTION: 55

Which of the following extractors can be used for Data Extraction Scope activity?

A. Intelligent Form Extractor, Machine Learning Extractor. Logic Extractor, and Regex Based Extractor.

B. Full Extractor. Machine Learning Extractor, Intelligent Form Extractor, and Regex Based Extractor.

C. Form Extractor Incremental Extractor Machine Learning Extractor and Intelligent Form Extractor

D. Regex Based Extractor. Form Extractor. Intelligent Form Extractor, and Machine Learning Extractor.

Answer: D (LEAVE A REPLY)

The Data Extraction Scope activity provides a scope for extractor activities, enabling you to configure them according to the document types defined in your taxonomy. The output of the activity is stored in an ExtractionResult variable, containing all automatically extracted data, and can be used as input for the Export Extraction Results activity. This activity also features a Configure Extractors wizard, which lets you specify exactly what fields from the document types defined in the taxonomy you want to extract¹.

The extractors that can be used for Data Extraction Scope activity are:

Regex Based Extractor: This extractor enables you to use regular expressions to extract data from text documents. You can define your own expressions or use the predefined ones from the Regex Based Extractor Configuration wizard².

Form Extractor: This extractor enables you to extract data from semi-structured documents, such as invoices, receipts, or purchase orders, based on the position and relative distance of the fields. You can define the templates for each document type using the Form Extractor Configuration wizard³.

Intelligent Form Extractor: This extractor enables you to extract data from semi-structured documents, such as invoices, receipts, or purchase orders, based on the labels and values of the fields. You can define the fields for each document type using the Intelligent Form Extractor Configuration wizard.

Machine Learning Extractor: This extractor enables you to extract data from any type of document, using a machine learning model that is trained on your data. You can use the predefined models from UiPath or your own custom models hosted on AI Center or other platforms. You can configure the fields and the model for each document type using the Machine Learning Extractor Configuration wizard.

References: 1: Data Extraction Scope 2: Regex Based Extractor 3: Form Extractor 4: Intelligent Form Extractor

5: Machine Learning Extractor

NEW QUESTION: 56

In which of the following scenarios, the ML Classifier is the only recommended classifier to be used, according to best practice?

A. When the custom document types are very similar and file splitting is not necessary.

B. When the custom document types are not similar and file splitting is not necessary.

- C. When the custom document types are not similar and file splitting is necessary.
- D. When the custom document types are very similar and file splitting is necessary.

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

The ML Classifier is a document classifier that uses a machine learning model deployed as an ML Skill in AI Center to perform document classification tasks. The ML Classifier can work by default with Invoices, Purchase Orders, Receipts, and Utility Bills, or with custom document types that are trained using the Data Manager and the Machine Learning Classifier Trainer¹².

According to the best practice, the ML Classifier is the only recommended classifier to be used when the custom document types are very similar and file splitting is not necessary. This is because the ML Classifier can handle complex and ambiguous cases where the document types are hard to distinguish by rules or keywords, and can also learn from feedback and improve over time. File splitting is not necessary when the documents are single-page or have a consistent number of pages per document type³.

The other options are not correct because they are scenarios where other classifiers, such as the Keyword Based Classifier or the Intelligent Keyword Classifier, can be used in combination with the ML Classifier or instead of it. These classifiers are based on rules or keywords that can identify the document types based on their content or metadata, and can also perform file splitting if the documents are multi-page or have a variable number of pages per document type³.

References: 1: Machine Learning Classifier - UiPath Activities 2: Machine Learning Classifier Trainer - UiPath Document Understanding 3: Document Classification - UiPath Document Understanding

NEW QUESTION: 57

Which of the following statements best defines Dashboards in UiPath Communications Mining?

- A. Dashboards are pre-set visual displays for each dataset that can be modified only by admins.
- B. Dashboards are used for monitoring label and entity performance.
- C. Dashboards are fully customizable pages containing relevant charts and visuals for a specific dataset.
- D. Dashboards are a tool used for monitoring the performance of live communications channel integrations.

Answer: ([SHOW ANSWER](#))

In UiPath Communications Mining, Dashboards are fully customizable pages that allow users to display relevant charts, visualizations, and insights for a specific dataset. These dashboards provide critical insights into the performance of models, labels, and other metrics relevant to the communication mining process.

Users can create and adjust these dashboards to reflect specific data points, making them a powerful tool for tracking and improving model performance over time.

For more details, refer to:

UiPath Communications Mining Dashboards: Customizing Dashboards

UiPath AI Center Documentation: Dashboards and Visualization

NEW QUESTION: 58

Which technology enables UiPath Communications Mining to analyze and enable action on messages?

- A. Natural Language Processing (NLP)
- B. Virtual Reality.
- C. Cloud Computing.
- D. Robotic Process Automation

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

UiPath Communications Mining is a new capability to understand and automate business communications. It uses state-of-the-art AI models to turn business messages-from emails to tickets-into actionable data. It does this in real time and on all major business communications channels¹. Natural Language Processing (NLP) is the branch of AI that deals with analyzing, understanding, and generating natural language. NLP enables UiPath Communications Mining to extract the most important data from any message, such as reasons for contact, data fields, and sentiment². NLP also allows UiPath Communications Mining to deploy custom AI models in hours, not weeks, by using automatic labeling and annotation².

References: 2 Communications Mining - Automate Business Communications | UiPath 1
Introducing UiPath Communications Mining | UiPath

NEW QUESTION: 59

How do the prediction mechanisms for labels and general fields differ in the UiPath Communications Mining platform?

- A. Label predictions are made based on the text of the message as well as some metadata properties, while general fields are learnt from the assigned span of text and the context of the text surrounding that span.
- B. Label predictions rely solely on metadata properties, general fields are learnt from the presence of certain key phrases in the text.
- C. Labels predictions are based on assigned span of text and entities are predicted solely from metadata properties.
- D. Labels and general fields are both predicted based on the text of the message, with no consideration given to metadata properties or contextual surrounding text.

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

Reference: UiPath Communications Mining

NEW QUESTION: 60

How do you choose the appropriate document processing methodology?

- A. Based on the text language and the number of pages.
- B. Based on the document structure and data points to be extracted.
- C. Based on the document type and length.
- D. Based on the number of fields to be extracted.

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

Reference: UiPath Document Understanding

NEW QUESTION: 61

Which generic ML Package should be used when the document type you are using is not part of the out of the box models?

- A. DocumentClassifier
- B. Document
- C. DocumentUnderstanding
- D. DocumentMLPackage

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

The DocumentUnderstanding ML package is a generic, retrainable model designed to handle various document types that are not covered by out-of-the-box models. It allows for the extraction of data from structured and semi-structured documents by building a model from scratch through training. This package is highly flexible and can be tailored to fit different document formats, making it ideal when specific document types are not pre-configured in UiPath's out-of-the-box offerings. (Source: UiPath Documentation on ML Packages)

Valid UiPath-SAIAv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIAv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIAv1 exam dumps**, the BraindumpsPass.com UiPath-SAIAv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIAv1 dumps with Test Engine here: <https://www.braindumpsPASS.com/UiPath/UiPath-SAIAv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF** Special Discount: **Exam-Tests**)

NEW QUESTION: 62

Which of the following is true when creating an ML Package in UiPath AI Center?

- A. The package name cannot use any Python reserved keywords.
- B. The package name cannot contain any spaces or hyphens.
- C. The package name cannot exceed 10 characters in length.
- D. The package name cannot include any special characters such as "@", "S", or "#".

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

When creating an ML Package in UiPath AI Center, the package name must adhere to Python naming conventions, meaning it cannot include reserved keywords such as class, break, finally, etc. This ensures that the package can be successfully deployed without conflicts in the Python environment.

NEW QUESTION: 63

In UiPath Communications Mining, what does the Reports section contain?

- A. Tools for evaluating model performance.

- B. Tools for evaluating label and entity performance.
- C. Tools for comparing different model versions.
- D. Tools for message analytics and monitoring.

Answer: ([SHOW ANSWER](#))

In UiPath Communications Mining, the Reports section provides a variety of tools for analyzing and monitoring messages within a dataset. This includes functionalities like label summaries, trends, and message segmentation, which allow users to gain insights into message volumes, sentiment trends, and other analytical metrics. The Reports section is integral for tracking the performance of messages and the overall dataset, making it useful for monitoring communication channels

NEW QUESTION: 64

Which of the following are the two key categories that use cases for UiPath Communications Mining typically fall into?

- A. Research and Development
- B. Communication and Collaboration
- C. Customer Support and Marketing
- D. Analytics and Automation

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

Comprehensive and Detailed Explanation From Exact Extract:

Use cases in UiPath Communications Mining typically fall into two major categories:

- * Analytics- Understand themes, sentiment, and intent in communications.
 - * Automation- Trigger workflows based on classified and extracted data from messages.
- These capabilities help organizations act on insights and automate responses efficiently.
- * UiPath Documentation Reference: Communications Mining Use Cases

NEW QUESTION: 65

When designing the Taxonomy for document types, what should be a primary consideration?

- A. Creating taxonomies that focus solely on scanned documents.
- B. Designing taxonomies without considering the need for post-processing rules.
- C. Focusing on creating separate taxonomies for each document type to avoid confusion.
- D. Grouping as many document types under the same taxonomy as possible.

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

When designing a taxonomy for document types in UiPath, a key consideration is to structure it in a way that maximizes efficiency and reusability. Grouping related document types under the same taxonomy helps to simplify processing and reduce redundancy. This approach ensures that similar document types are treated consistently, making it easier to apply extraction methods and post-processing rules across different but related document types. Over-segmentation into separate taxonomies for each document type can lead to unnecessary complexity and confusion, making management and scaling of automation workflows more difficult. The goal is to create a cohesive structure that can handle various document types effectively.

(Source: UiPath Document Understanding and Communications Mining documentation)

NEW QUESTION: 66

How many types of synchronization mechanisms exist in the Document Understanding Process to prevent multiple robots to write in a file at the same time?2

- A. 2
- B. 3
- C. 4
- D. 5

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

The Document Understanding Process uses two types of synchronization mechanisms to prevent multiple robots from writing in a file at the same time: file locks and queues. File locks are used to ensure that only one robot can access a file at a time, while queues are used to store the information extracted from the documents and to avoid data loss or duplication. The process uses the following activities to implement these mechanisms:

File Lock Scope: This activity creates a lock on a file or folder and executes a set of activities within it. The lock is released when the scope ends or when an exception occurs. This activity ensures that only one robot can read or write a file or folder at a time, and prevents other robots from accessing it until the lock is released.

Add Queue Item: This activity adds an item to a queue in Orchestrator, along with some relevant information, such as a reference, a priority, or a deadline. The item can be a JSON object, a string, or any serializable .NET type. This activity ensures that the information extracted from the documents is stored in a queue and can be retrieved by another robot or process later.

Get Queue Items: This activity retrieves a collection of items from a queue in Orchestrator, based on some filters, such as status, reference, or creation time. The items can be processed by the robot or passed to another activity, such as Update Queue Item or Delete Queue Item. This activity ensures that the information stored in the queue can be accessed and manipulated by the robot or process.

References: Document Understanding Process: Studio Template, File Lock Scope, Add Queue Item, [Get Queue Items]

NEW QUESTION: 67

What are the available options for Scoring in Document Manager, that apply to string fields only?

- A. Exact match and Auto.
- B. Exact match and First span.
- C. Exact match and Native string search.
- D. Exact match and Levenshtein.

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

In UiPath Document Manager, for string fields, the available options for scoring include:

Exact match: This checks whether the extracted string exactly matches the expected value.

Levenshtein: This method computes the "distance" between two strings, which is the number of single-character edits (insertions, deletions, or substitutions) required to change one string into the other. It is used to assess the similarity between the extracted string and the expected value. These scoring methods help evaluate the accuracy of string extraction in the document understanding process.

For more details, refer to:

UiPath Document Understanding Framework: String Field Scoring

Levenshtein Distance in UiPath: Scoring String Fields

NEW QUESTION: 68

What is the difference between OCR (Optical Character Recognition) and IntelligentOCR?

- A.** OCR (Optical Character Recognition) is a method that reads text from images, recognizing each character and its position, while IntelligentOCR is an enhanced version of it that can also work with noisier input data.
- B.** IntelligentOCR is simply a rebranding of the OCR (Optical Character Recognition), both of them being methods that read text from images, recognizing each character and its position.
- C.** OCR (Optical Character Recognition) is a UiPath Studio activity package that contains IntelligentOCR as an activity used to read text from images, recognizing each character and its position. OCR is widely used in Document Understanding processes.
- D.** IntelligentOCR is a UiPath Studio activity package that contains all the activities needed to enable information extraction, while OCR (Optical Character Recognition) is a method that reads text from images, recognizing each character and its position.

Answer: ([SHOW ANSWER](#))

According to the UiPath documentation and web search results, OCR (Optical Character Recognition) is a method that reads text from images, recognizing each character and its position. OCR is used to digitize documents and make them searchable and editable. OCR can be performed by different engines, such as Tesseract, Microsoft OCR, Microsoft Azure OCR, OmniPage, and Abbyy. OCR is a basic step in the Document Understanding Framework, which is a set of activities and services that enable the automation of document processing workflows. IntelligentOCR is a UiPath Studio activity package that contains all the activities needed to enable information extraction from documents. Information extraction is the process of identifying and extracting relevant data from documents, such as fields, tables, entities, and labels.

IntelligentOCR uses different components, such as classifiers, extractors, validators, and trainers, to perform information extraction.

IntelligentOCR also supports different formats, such as PDF, PNG, JPG, TIFF, and BMP.

IntelligentOCR is an advanced step in the Document Understanding Framework, which builds on the OCR output and provides more functionality and flexibility.

References:

About the IntelligentOCR Activities Package

OCR Activities

OCR Feature Comparison: UiPath Community vs UiPath Licensed OCR

NEW QUESTION: 69

A developer has created a string array variable as shown below:

UserNames = {"Jane", "Jack", "Jill", "John"}

Which expression should the developer use in a Log Message activity to print the elements of the array separated by the string ","?

- A. String.Concat(UserNames, ",")
- B. String.Concat(",", UserNames)
- C. String.Join(UserNames, ",")
- D. String.Join(",", UserNames)

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

Reference: UiPath String Manipulations

NEW QUESTION: 70

Under what condition can a project be deleted in UiPath AI Center?

- A. If it does not have any pipeline data.
- B. If it does not have any running pipelines.
- C. If it does not have any deployed packages.
- D. If it does not have any scheduled pipelines.

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

A project in UiPath AI Center is an isolated group of resources (datasets, pipelines, packages, skills, and logs) that you use to build a specific ML solution. You can create, edit, or delete projects from the Projects page or the project's Dashboard page. However, you can only delete a project if it does not have any package currently deployed in a skill. A package is a versioned and deployable unit of an ML model or an OS script that can be used to create an ML skill. A skill is a consumer-ready, live deployment of a package that can be used in RPA workflows in Studio. If a project has a package deployed in a skill, you need to undeploy the skill first before deleting the project. This is to ensure that you do not accidentally delete a project that is being used by a skill¹².

References: 1: AI Center - Managing Projects 2: AI Center - Managing ML Skills

NEW QUESTION: 71

Which are all the options for managing ML Skills?

- A. ML skills can be created, stopped, redeployed, updated to a new package version, rolled back to a previous package version, modified to use or not use GPU. modified to use or not use AI units, made public or private, or deleted.
- B. ML skills can be created, stopped, redeployed, updated to a new package version, rolled back to a previous package version, made public or private, or deleted.

C. ML skills can be created, stopped, redeployed, updated to a new package version, rolled back to a previous package version, modified to use or not use GPU. made public or private, or deleted.

D. ML skills can be created, updated to a new package version, rolled back to a previous package version, modified to use or not use GPU. made public or private, or deleted.

Answer: ([SHOW ANSWER](#))

In UiPath AI Center, ML Skills can be managed in various ways, allowing users to customize and control how these skills are deployed and used. The management options include:

Creating a new ML skill.

Stopping a deployed skill.

Redeploying an ML skill.

Updating to a new package version.

Rolling back to a previous version if needed.

Modifying GPU usage.

Modifying the use of AI units.

Making the skill public or private.

Deleting an ML skill when no longer needed.

This provides flexibility for both managing the ML infrastructure and optimizing resources in real-time.

For more details, refer to:

UiPath AI Center Documentation: Managing ML Skills

ML Skill Management Options: Managing Machine Learning Skills in AI Center

NEW QUESTION: 72

Who is responsible for devising a strategy to prioritize processes during the Business Case and Technical Validation phase?

A. Project Manager

B. Solution Architect

C. Automation Developer

D. Business Analyst

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

The Solution Architect is responsible for devising a strategy to prioritize processes during the Business Case and Technical Validation phase. Their role involves assessing technical feasibility, scalability, and business value to determine process prioritization.

Reference: UiPath Solution Architect Role

NEW QUESTION: 73

What should a UiPath Communications Mining taxonomy contain when it is being imported?

A. Label predictions.

B. Entity descriptions.

C. Entity predictions.

D. Label descriptions.

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

According to the UiPath documentation, a UiPath Communications Mining taxonomy is a collection of all the labels applied to the verbatims in a dataset, structured in a hierarchical manner. Labels are the concepts and intents that you want to capture in the dataset to suit your specific objectives. When you import your taxonomy from a spreadsheet, you need to provide the label descriptions, which are the names of the labels and their level in the hierarchy. Label predictions and entity predictions are the outputs of the model training process, and they are not part of the taxonomy. Entity descriptions are the definitions of the entities that you want to extract from the verbatims, and they are not part of the taxonomy either.

References:

Communications Mining - Taxonomies

Communications Mining - Importing your taxonomy

Communications Mining - Building your taxonomy structure

NEW QUESTION: 74

How can a Pipeline be scheduled?

A. A Pipeline can be scheduled at multiple future dates or with a recurring schedule.

B. A Pipeline can be scheduled at a single future date.

C. A Pipeline can be scheduled with a recurring schedule.

D. A Pipeline can be scheduled at a single future date or with a recurring schedule.

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

In UiPath's AI Center, a Pipeline can be scheduled for execution either at a specific future date or with a recurring schedule. This allows for flexibility in automating the retraining or execution of machine learning models based on predefined time intervals or specific dates. For example, you can schedule a pipeline to run once on a given date or set it to run daily, weekly, or monthly, depending on the project's needs.

For more information, refer to:

UiPath AI Center Documentation: Pipeline Scheduling

Automation Pipelines in AI Center: Pipeline Execution and Scheduling

NEW QUESTION: 75

Which of the following document type is part of UiPath's out-of-the-box models?

A. Emails.

B. Medical prescriptions.

C. Invoices.

D. Letters.

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

UiPath offers several pre-trained out-of-the-box models, and Invoices is one of the most commonly supported document types. This model is optimized to extract relevant fields such as

invoicenumbers, dates, amounts, and more from structured invoices. Other common out-of-the-box models include receipts and purchase orders.

NEW QUESTION: 76

What is one of the purposes of the Config file in the UiPath Document Understanding Template?

- A. It contains the configuration settings for the UiPath Robot and Orchestrator integration.
- B. It stores the API keys and authentication credentials for accessing external services.
- C. It specifies the output file path and format for the processed documents.
- D. It defines the input document types and formats supported by the template.

Answer: B (LEAVE A REPLY)

The Config file in the UiPath Document Understanding Template is a JSON file that contains various parameters and values that control the behavior and functionality of the template. One of the purposes of the Config file is to store the API keys and authentication credentials for accessing external services, such as the Document Understanding API, the Computer Vision API, the Form Recognizer API, and the Text Analysis API. These services are used by the template to perform document classification, data extraction, and data validation tasks. The Config file also allows the user to customize the template according to their needs, such as enabling or disabling human-in-the-loop validation, setting the retry mechanism, defining the custom success logic, and specifying the taxonomy of document types.

References: Document Understanding Process: Studio Template, Automation Suite - Document Understanding configuration file

Valid UiPath-SAIv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIv1 exam dumps**, the BraindumpsPass.com UiPath-SAIv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIv1 dumps with Test Engine here: <https://www.braindumpsPASS.com/UiPath/UiPath-SAIv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 77

What happens during the Classify stage of the Document Understanding Framework?

- A. The OCR engine is used to extract text from the image document.
- B. The extracted data is exported as a dataset.
- C. The target fields are extracted from the document and sent to Action Center for human validation.
- D. The documents are included in one of the taxonomy document types or skipped.

Answer: D (LEAVE A REPLY)

According to the UiPath documentation, the Classify stage of the Document Understanding Framework is used to automatically determine what document types are found within a digitized

file. The document types are defined in the project taxonomy, which is a collection of all the labels and fields applied to the documents in a dataset. The Classify stage uses one or more classifiers, which are algorithms that assign document types to files based on their content and structure. The classifiers can be configured and executed using the Classify Document Scope activity, which also allows for document type filtering, taxonomy mapping, and minimum confidence threshold settings. The Classify stage outputs the classification information in a unified manner, irrespective of the source of classification. The documents that are classified are then sent to the next stage of the framework, which is Data Extraction. The documents that are not classified or skipped are either excluded from further processing or sent to Action Center for human validation and correction.

References:

Document Understanding - Document Classification Overview

Document Understanding - Introduction

Generative Extraction & Classification using Document Understanding in Cross-Platform Projects (Public Preview)

NEW QUESTION: 78

What is the Machine Learning Extractor?

- A.** A specialized model that can recognize multiple languages in the same document using API calls to a Hugging Face model with over 250 languages.
- B.** An extraction model that can be enabled and trained in AI Center. For better accuracy. 25 documents per model are recommended to train the model.
- C.** A tool using machine learning models to identify and report on data targeted for data extraction.
- D.** A tool that helps extract data from different document structures, and is particularly useful when the same document has multiple formats.

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

The Machine Learning Extractor utilizes machine learning models to effectively extract data from documents, especially when dealing with varying structures or formats within the same document type. This capability is crucial in scenarios where documents do not follow a strict template and have variations in their layout or content organization. The extractor can be trained to understand these variations and accurately extract the needed information.

The Machine Learning Extractor is a data extraction tool that uses machine learning models to extract data from various types of documents, such as invoices, receipts, or forms. It is especially useful when the same document type has multiple layouts or formats, as it can learn and infer the values for the targeted fields, even from documents and layouts it has never seen before¹.

The Machine Learning Extractor can be used with one of UiPath's public Document Understanding endpoints, which provide generic models for certain document types, or with custom trained machine learning models hosted in AI Center, which can be tailored to specific use cases. The Machine Learning Extractor can be configured and trained using the Data Extraction Scope activity in UiPath Studio².

References:

Document Understanding - Machine Learning Extractor

UiPath Activities - Data Extraction Scope

NEW QUESTION: 79

What is the main difference between an array and a list in UiPath?

- A.** An array is a fixed-size collection of elements of the same type while a list is a dynamic-sized collection of elements of different types.
- B.** An array is a fixed-size collection of elements of different types while a list is a dynamic-sized collection of elements of the same type.
- C.** An array is a dynamic-sized collection of elements of the same type while a list is a fixed-size collection of elements of the same type.
- D.** An array is a fixed-size collection of elements of the same type while a list is a dynamic-sized collection of elements of the same type.

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

Comprehensive and Detailed Explanation From Exact Extract:

- * Arrays in UiPath (VB.NET) are fixed-size and must be initialized with a defined number of elements of the same type.
- * Lists (List<T>) are dynamic and allow you to add/remove elements at runtime, but still enforce type safety (same type elements).
- * UiPath Documentation Reference: Collections in UiPath Academy # RPA Developer Foundation # Data Manipulation

NEW QUESTION: 80

What happens when multiple users try to label the same document concurrently?

- A.** The changes made by one user override the changes made by others.
- B.** The changes made by all users are saved successfully.
- C.** Concurrent labeling is not allowed.
- D.** A warning message is displayed to the other user(s) indicating unsuccessful changes.

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

According to the UiPath documentation, data labeling is a process that involves uploading raw data, annotating text data in the labeling tool, and using the labeled data to train ML models¹. Data labeling is performed by human labelers, who can be either internal or external to the organization². However, concurrent labeling is not supported by the UiPath Data Labeling tool, which means that only one user can label a document at a time³. If multiple users try to label the same document concurrently, they will encounter an error message that says "The document is locked by another user. Please try again later.". Therefore, the correct answer is C.

References:

1: About Data Labeling 2: Data Labeling Roles 3: Data Labeling Limitations : Data Labeling Error Messages

NEW QUESTION: 81

Which of the following is an indicator that sufficient training has been completed for a model in UiPath Communications Mining?

- A. A model rating of 30-40.
- B. A model rating of 40-50.
- C. A model rating of 50-60.
- D. A model rating of 70-90 or better.

Answer: ([SHOW ANSWER](#))

The model rating is a proprietary score that assesses the overall health and performance of a model in UiPath Communications Mining. It considers four main factors: balance, underperforming labels, coverage, and all labels. The model rating is a score from 0 to 100, which equates to a rating of 'Poor' (0-49), 'Average' (50-69), 'Good' (70-89) or 'Excellent' (90-100). A model rating of 70-90 or better indicates that the model has sufficient training and performs well in all of the most important areas. A model rating of 70-90 or better also means that the model has a balanced and representative training data, a low number of labels with performance issues or warnings, a high coverage of the dataset by informative labels, and a high average precision of all labels.

References: Communications Mining - Model Rating, Communications Mining - Understanding and improving model performance

NEW QUESTION: 82

If Label X in UiPath Communications Mining has 80% precision at a given confidence threshold, what output should this provide?

- A. For every 100 messages, 80 would be labelled correctly as 'Label X', and 20 would be labelled incorrectly.
- B. For every 100 messages which should have been labelled as 'Label X', 80 were labelled and 20 were missed.
- C. For every 100 messages, 20 would be labelled correctly as 'Label X', and 80 would be labelled incorrectly.
- D. For every 100 messages which should have been labelled as 'Label X', 20 were labelled and 80 were missed.

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

If Label X has 80% precision at a given confidence threshold, this means that out of every 100 messages predicted to have Label X, 80 are correctly labelled, and 20 are incorrectly assigned the label. Precision measures the correctness of predictions made, so in this case, 80% of the predictions were correct, while 20% were false positives

NEW QUESTION: 83

Which of the following best describes UiPath Document Understanding?

- A. A suite of tools for automating document processing tasks.
- B. A solution for managing cloud infrastructure.

- C. A software for creating machine learning models.
- D. A platform for managing robotic process automation (RPA) workflows.

Answer: A (LEAVE A REPLY)

Reference: UiPath Document Understanding

NEW QUESTION: 84

What are the two main data extraction methodologies used in document understanding processes?

- A. Hybrid and manual data extraction.
- B. Rule-based and model-based data extraction.
- C. Rule-based and hybrid data extraction.
- D. Manual and model-based data extraction.

Answer: B (LEAVE A REPLY)

According to the UiPath documentation, there are two common types of data extraction methodologies used in document understanding processes: rule-based data extraction and model-based data extraction¹². Rule-based data extraction targets structured documents, such as forms, invoices, or receipts, that have a fixed layout and a predefined set of fields. Rule-based data extraction uses predefined rules, such as regular expressions, keywords, or coordinates, to locate and extract the relevant data from the documents¹. Model-based data extraction is used to process semi-structured and unstructured documents, such as contracts, emails, or reports, that have a variable layout and a diverse set of fields. Model-based data extraction uses machine learning models, such as neural networks, to learn from examples and extract the relevant data from the documents¹. Both methodologies have their advantages and limitations, and depending on the use case, they can be used separately or in combination, in a hybrid approach².

References: 1: Data Extraction Overview 2: Document Processing with Improved Data Extraction

NEW QUESTION: 85

Which of the following OCR (Optical Character Recognition) engines is not free of charge?

- A. Tesseract.
- B. Microsoft Azure OCR.
- C. OmniPage.
- D. Microsoft OCR.

Answer: (SHOW ANSWER)

According to the UiPath documentation, OmniPage is a paid OCR engine that requires a license to use. It is one of the most accurate and reliable OCR engines available, and it supports over 200 languages. The other OCR engines listed are free of charge, but they may have different features, limitations, and performance levels. For example, Tesseract is an open-source OCR engine that supports over 100 languages, but it may not be as accurate as OmniPage. Microsoft Azure OCR and Microsoft OCR are both cloud-based OCR engines that use Microsoft's technology, but they have different capabilities and pricing models. Microsoft Azure OCR can process both printed and handwritten text, and it uses a pay-as-you-go model based on the

number of transactions. Microsoft OCR can only process printed text, and it is included in the UiPath Studio license.

References:

Document Understanding - OCR Engines

Automation Pricing - Complete UiPath Enterprise Solution

NEW QUESTION: 86

Which is a high-level view of the tabs within an AI Center project?

- A. Dashboard. Datasets. ML Packages. ML Training. ML Evaluation, and ML Logs.
- B. Datasets, Data Labeling. ML Packages, ML Training, ML Evaluation, ML Skills, and ML Logs.
- C. Datasets. Data Labeling. ML Packages. Pipelines, and ML Skills.
- D. Dashboard. Datasets, Data Labeling. ML Packages. Pipelines, ML Skills, and ML Logs.

Answer: (SHOW ANSWER)

A high-level view of the tabs within an AI Center project is as follows:

Dashboard: This tab provides an overview of the project's status, such as the number of datasets, pipelines, packages, skills, and logs, as well as the AI Units consumption and quota.

Datasets: This tab enables you to upload, view, and manage the datasets that are used for training and evaluating the ML models within the project. A dataset is a folder of storage containing arbitrary files and sub- folders¹.

Data Labeling: This tab enables you to upload raw data, annotate text data in the labeling tool (for classification or entity recognition), and use the labeled data to train ML models. It is also used by the human reviewer to re-label incorrect predictions as part of the feedback process².

ML Packages: This tab enables you to upload, view, and manage the ML packages and package versions within the project. An ML package is a group of package versions of the same package type, and a package version is a trained model that can be deployed to a skill³.

Pipelines: This tab enables you to create, view, and manage the pipelines and pipeline runs within the project. A pipeline is a description of an ML workflow, including the functions and their order of execution, and a pipeline run is an execution of a pipeline based on code provided by the user⁴.

ML Skills: This tab enables you to deploy, view, and manage the ML skills within the project. An ML skill is a live deployment of a package version, which can be consumed by an RPA workflow using an ML skill activity in UiPath Studio⁵.

ML Logs: This tab enables you to view and filter the logs related to the project, such as the events, messages, and errors that occurred during the pipeline runs, skill deployments, and skill executions⁶.

References:

1: About Datasets 2: About Data Labeling 3: About ML Packages 4: About Pipelines 5: About ML Skills 6: About ML Logs

NEW QUESTION: 87

When using UiPath Studio's publishing options, which location(s) can automation projects be published to?

- A. Orchestrator, Locally, and Git repository.
- B. Orchestrator, Locally, and SharePoint.
- C. Orchestrator, Locally, and Custom NuGet feed.
- D. Custom NuGet feed, Cloud-based storage, and SharePoint.

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

Comprehensive and Detailed Explanation From Exact Extract:

When publishing a process from UiPath Studio, it can be directed to:

- * Orchestrator(via Tenant or Personal Workspace feed)
- * Locally to a file path
- * Custom NuGet Feed(configured in NuGet.config)

Git repositories and SharePoint are not supported as direct publish targets.

- * UiPath Documentation Reference: Publishing Projects - UiPath Studio

NEW QUESTION: 88

What additional information can be included in the exported data, apart from the extraction results?

- A. The number of occurrences and the extraction confidence.
- B. The page number from which the field was extracted and the exact position on the page.
- C. The extraction confidence and the digitization confidence.
- D. The position on the page.

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

The exported data from the UiPath Document Understanding Template contains the extraction results in a JSON format, along with some additional information that can be useful for debugging or analysis purposes.

One of the additional information that can be included is the page number from which the field was extracted and the exact position on the page, represented by the coordinates of the bounding box. This information can help to locate the field on the original document image and to verify the accuracy of the extraction. The additional information can be enabled or disabled by setting the IncludeMetadata parameter to true or false in the Config file of the template.

References: Document Understanding Process: Studio Template, Export Results

NEW QUESTION: 89

How can you build custom models supported by AI Center?

- A. Using the AI Center IDE (Integrated Development Environment).
- B. Using the AI Center model builder.
- C. Using a Python IDE (Integrated Development Environment) or an AutoML platform.
- D. Using a C/C++ IDE (Integrated Development Environment), then upload the code to AI Center IDE.

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

To build custom models supported by AI Center, you can use a Python IDE or an AutoML platform of your choice. A Python IDE is a software application that provides tools and features for writing, editing, debugging, and running Python code. An AutoML platform is a service that automates the process of building and deploying machine learning models, such as data preprocessing, feature engineering, model selection, hyperparameter tuning, and model evaluation. Some examples of Python IDEs are PyCharm, Visual Studio Code, and Jupyter Notebook. Some examples of AutoML platforms are Google Cloud AutoML, Microsoft Azure Machine Learning, and DataRobot.

To use a Python IDE, you need to install the required Python packages and dependencies, write the code for your model, and test it locally. Then, you need to package your model as a zip file that follows the AI Center ML Package structure and requirements. You can then upload the zip file to AI Center and create an ML Skill to deploy and consume your model.

To use an AutoML platform, you need to sign up for the service, upload your data, configure your model settings, and train your model. Then, you need to export your model as a zip file that follows the AI Center ML Package structure and requirements. You can then upload the zip file to AI Center and create an ML Skill to deploy and consume your model.

References: AI Center - Building ML Packages, AI Center - ML Package Structure, AI Center - Creating ML Skills

NEW QUESTION: 90

Which of the following time periods can be selected when viewing Trends in UiPath Communications Mining?

Which of the following time periods can be selected when viewing Trends in UiPath Communications Mining?

- A. Daily, Monthly, Quarterly, Yearly.
- B. Daily, Weekly, Monthly, Yearly.
- C. Daily, Bi-weekly, Monthly, Yearly.
- D. Daily, Bi-weekly, Quarterly, Yearly.

Answer: ([SHOW ANSWER](#))

According to the UiPath Communications Mining documentation, the Trends tab in the Reports page displays charts for verbatim volume, label volume, and sentiment over a selected time period. Users can choose the time period for the data in the filter bar, and the time sequencing of the chart (i.e. daily, weekly, etc.) in the top right dropdown menu. The available options for the time sequencing are Daily, Weekly, Monthly, and Yearly. These options allow users to see how the trends change over different time intervals and identify patterns or anomalies.

References:

Communications Mining - Using Reports

Communications Mining - Trends

NEW QUESTION: 91

What does a UiPath Communications Mining taxonomy include?

- A. General fields and datasets.
- B. Labels and sources.
- C. Messages, labels, and general fields.
- D. Labels and general fields.

Answer: D (LEAVE A REPLY)

Reference: UiPath Communications Mining

Valid UiPath-SAIAv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIAv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIAv1 exam dumps**, the BraindumpsPass.com UiPath-SAIAv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIAv1 dumps with Test Engine here: <https://www.braindumpsPass.com/UiPath/UiPath-SAIAv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF** Special Discount: **Exam-Tests**)

NEW QUESTION: 92

A developer utilized the Add Data Row activity to insert a row into a data table called "dt_Reports".

However, during runtime, UiPath Studio encounters an exception, "Add Data Row: Object reference not set to an instance of an object," because the data table has not been initialized. To rectify this issue, what should the developer include in an Assign before the Add Data Row activity?

- A. Assign dt_Reports = New System.Data.DataRow
- B. Assign dt_Reports = New System.Data.DataTable
- C. Assign dt_Reports = New List(Of DataRow)
- D. Assign New System.Data.DataTable = dt_Reports

Answer: B (LEAVE A REPLY)

Comprehensive and Detailed Explanation From Exact Extract:

The error "Object reference not set to an instance of an object" occurs when the DataTable variable (dt_Reports) hasn't been initialized.

To fix this, use:

vb

CopyEdit

Assign dt_Reports = New System.Data.DataTable

Then, add necessary columns before adding rows.

* Note: You cannot assign a new row (DataRow) without first initializing the DataTable.

* UiPath Documentation Reference: DataTable Initialization - UiPath Academy Add Data Row - UiPath Docs

Valid UiPath-SAIAv1 Dumps shared by BraindumpsPass.com for Helping Passing UiPath-SAIAv1 Exam! BraindumpsPass.com now offer the **newest UiPath-SAIAv1 exam dumps**, the BraindumpsPass.com UiPath-SAIAv1 exam **questions have been updated** and **answers have been corrected** get the **newest** BraindumpsPass.com UiPath-SAIAv1 dumps with Test Engine here: <https://www.braindumpspass.com/UiPath/UiPath-SAIAv1-practice-exam-dumps.html> (250 Q&As Dumps, **40%OFF** Special Discount: **Exam-Tests**)