

DUMPS ARENA

**UiPath Certified Professional Specialized AI
Professional v1.0**

UiPath UiPath-SAlv1

Version Demo

Total Demo Questions: 23

Total Premium Questions: 231

Buy Premium PDF

<https://dumpsarena.co>

sales@dumpsarena.co

sales@dumpsarena.co
dumpsarena.co

Topic Break Down

Topic	No. of Questions
Topic 1, AI Center	36
Topic 2, Document Understanding	109
Topic 3, Communications Mining	58
Topic 4, Generative AI	1
Topic 5, Mix Questions	27
Total	231

QUESTION NO: 1

For what type of documents is it recommended to use the RegEx Based Extractor?

- A. For structured or semi-structured documents in which layouts of different document providers have little to no variation.
- B. For structured or semi-structured documents in which layouts of different document providers vary greatly.
- C. For structured or semi-structured documents where the extraction method is based on a set of FlexiCapture definition files.
- D. For structured or semi-structured documents where for certain fields, data is always found in a strict, predictable format and context.

ANSWER: D

Explanation:

For structured or semi-structured documents where for certain fields, data is always found in a strict, predictable format and context is the recommended use case for the RegEx Based Extractor. This extractor identifies target values by applying regular-expression patterns, so it is most dependable when a field has a stable syntax and can be located using consistent surrounding text or context. Typical examples include identifiers with a fixed structure, dates in a known format, email addresses, tax IDs, purchase-order numbers, and invoice numbers. In a document-processing workflow, the regular expression defines the expected value pattern while an anchor or context helps constrain the search to the relevant area or label. This makes the approach especially useful when only particular fields are highly standardized, even if the overall document is structured or semi-structured. Designing the expression around the real field format and configuring reliable context improves extraction precision and reduces accidental matches elsewhere in the document. UiPath documents the extractor as a rule-based extraction capability intended for predictable data patterns and contextual extraction scenarios. See the [UiPath RegEx Based Extractor documentation](#) and the [UiPath extraction methods overview](#).

QUESTION NO: 2

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

Explanation:

Use the "Teach" training mode to pin the label to more communications. Once a newly introduced label begins to appear in predictions, even rarely and with low confidence, the model has enough initial signal to support targeted teaching. Teach mode is designed for building the model's understanding of a particular label by presenting communications that are relevant to that label and confirming the label on them. Adding a broader and more representative set of correctly labelled communications supplies additional positive training examples, enabling the model to learn the linguistic patterns, context, and variation associated with the label. After these examples are added and the model is retrained, prediction frequency and confidence can improve because the model has stronger evidence for recognizing the label across the dataset. This iterative workflow—seed a label with examples, retrain, then use teaching to expand high-quality examples—is a core Communications Mining training practice. UiPath's documentation describes the training workflow and the role of targeted label training in improving model performance: [Model training](#) and [Model training best practices](#).

QUESTION NO: 3

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: D

Explanation:

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. OCR is the document digitization capability that converts text visible in an image or scanned document into machine-readable text, commonly with positional information. This output can then be consumed by later automation steps.

IntelligentOCR extends beyond recognizing text. Its activities support an end-to-end document-processing workflow: digitizing documents, classifying document types, extracting structured and semi-structured information such as fields and tables, validating extracted results, and training components where applicable. Thus, OCR can be a component used within an IntelligentOCR/Document Understanding workflow, but it is not itself the complete information-extraction framework. The distinction is important when designing automations: a process that only needs readable text may use OCR directly, whereas a process that needs business data such as invoice numbers, dates, line items, or customer details uses the broader IntelligentOCR capabilities.

UiPath describes the IntelligentOCR package and its document-processing activities in its [IntelligentOCR Activities Package documentation](#). The wider workflow and its stages are covered in the [Document Understanding introduction](#).

QUESTION NO: 4

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

Explanation:

If it is not being used in any active pipeline. is correct because AI Center protects dataset resources that are currently referenced by pipeline activity. A dataset can be an input or output dependency for training, evaluation, or other pipeline runs; deleting it while that dependency is active could prevent the pipeline from accessing required data and compromise the operation. Consequently, before deleting a dataset, the user must ensure that no active pipeline is using it. This requirement is based on the dataset's current usage, rather than on whether it was used at some point in the past. Once no active pipeline references the dataset, it can be safely removed through AI Center dataset management. This safeguard supports reliable execution of ML workflows and prevents accidental removal of data needed by an in-progress process. UiPath's AI Center documentation describes dataset management and the relationship between datasets and pipeline operations in its

[Managing datasets](#) guidance. For the broader process in which datasets are consumed by AI Center operations, see UiPath's [Pipelines](#) documentation.

QUESTION NO: 5 - (SIMULATION)

What is the order of steps for automatically retraining and deploying a Document Understanding ML Model in AI Center with data from Document Validation Action?

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

ANSWER: See the explanation for the answer

Explanation:

The correct order is:

1. **Send Human-in-the-Loop data to the used Dataset.**
2. **Use the Export feature from Document Manager using Scheduled Export.**
3. **Run a Training/Full Pipeline in AI Center using a Scheduled Pipeline with Auto-Retraining enabled.**
4. **Enable the Auto Update feature in the ML Skill.**

This sequence first captures the human validation/correction data, then exports it on a schedule for use in the AI Center dataset. The scheduled training or full pipeline retrains the model using the updated data, and ML Skill **Auto Update** deploys the newly trained model version automatically.

QUESTION NO: 6

What is the definition of a UiPath Communications Mining data source?

- A. A collection of raw unlabeled communications data of a similar type, that can be associated with up to 10 datasets.
- B. The model that we create when training the platform to understand the data in those sources.
- C. A permissioned storage area within the platform which contains communications and labels.
- D. A user-permissioned project containing a taxonomy with labels and entities.

ANSWER: A

Explanation:

A collection of raw unlabeled communications data of a similar type, that can be associated with up to 10 datasets. This is the correct definition because a source is the intake-level container for communications that share a common format or origin, such as emails, chat conversations, or support tickets. Sources hold the original communications data before it is interpreted through a dataset's taxonomy and model-training workflow. Keeping similar raw data together in a source enables the platform to process it consistently and allows the same underlying communications to support more than one analytical or automation use case. UiPath permits a single source to be associated with up to ten datasets, so teams can create separate datasets for distinct business objectives, label taxonomies, or model configurations while reusing the same imported communications. This separation is important: the source represents the raw input data, whereas datasets provide the context in which users define labels, train understanding, and analyze results. It also supports controlled reuse without requiring duplicate imports of the original communications. UiPath's Communications Mining documentation defines sources in these terms and describes their relationship to datasets: [UiPath Communications Mining — Sources](#).

QUESTION NO: 7

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

Explanation:

Out-of-the-box models from UiPath, UiPath technology partners, open-source models from the community, and custom models is correct because AI Center is designed to support ML packages from multiple sources rather than limiting users to a single model catalog. UiPath provides ready-to-use, prebuilt ML packages for common intelligent automation scenarios. Organizations can also use packages supplied by UiPath technology partners, enabling integration of specialized third-party AI capabilities into automations.

AI Center additionally supports open-source ML packages, which can be packaged and deployed for use in the platform. This gives teams flexibility to adopt community-developed frameworks and models where appropriate. Finally, teams can build, train, package, and deploy custom models for business-specific requirements and proprietary data. Together, these sources let an organization select an existing model where it fits or operationalize a tailored model when a use case requires it. AI Center provides the deployment and management capabilities that make these models available to UiPath automation workflows. See the [UiPath AI Center ML packages documentation](#) and the [AI Center overview](#).

QUESTION NO: 8

Which of the following file types are supported for the DocumentPath property in the Classify Document Scope activity?

- A.** .bmp, .pdf, .jpe, .psd
- B.** .png, .gif, .jpe, .tiff
- C.** .pdf, .jpeg, .raw, tif
- D.** .jpe, .eps, .jpg, .tiff

ANSWER: B

Explanation:

The correct selection is **.png, .gif, .jpe, .tiff**. The `DocumentPath` property of UiPath's Classify Document Scope identifies the input document that the activity sends through the document-classification workflow. UiPath supports common raster-image formats for this input, including PNG, GIF, JPE/JPEG, TIFF/TIF, BMP, and JPG, as well as PDF documents. Every extension in this selection is therefore within the documented supported set. This matters because the scope must be able to load the source document into the Intelligent Document Processing pipeline before applying the configured classifier. The property itself is populated with a string path or a String variable containing the file location; the file at that location must use one of the supported document formats. UiPath's activity documentation lists the supported extensions in the property details for the activity, while the broader Document Understanding documentation describes the use of image and PDF files as document inputs. See [UiPath Classify Document Scope documentation](#) and [UiPath Document Understanding documentation](#).

QUESTION NO: 9

What will be displayed in the Output panel after running the workflow below?



- A. Status not valid.
- B. Status is Unknown.
- C. Status is Pending.
- D. Status is New.

ANSWER: A

Explanation:

The displayed output is **Status not valid**. The workflow first assigns the string value "Approved" to the `Status` variable. It then evaluates that variable in the status-checking control-flow activity. The explicitly configured cases handle the values "Pending", "New", and "Unknown". Because "Approved" does not correspond to any of those case values, execution follows the Default path. The Write Line activity placed in that path writes `Status not valid.` to the Output panel.

This is the intended use of a Default branch in UiPath conditional branching: it provides a defined execution route for values that do not match a configured case. As a result, the workflow can handle an unexpected or unsupported status without requiring a separate condition for every possible string. UiPath's control-flow documentation describes using branching activities to direct execution according to evaluated values, while the Output panel displays messages generated by activities such as Write Line. See [UiPath Control Flow documentation](#) and [UiPath Output Panel documentation](#).

QUESTION NO: 10

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: D

Explanation:

First_name123 is accepted because Document Manager column field names can use letters, numbers, and underscores, while beginning with a letter. This naming pattern provides a valid identifier that Document Understanding can reliably use when defining and referencing fields in a document type. The underscore is supported as a separator within the name, and digits may be included after the initial letter, so the field name remains both readable and compliant. Defining valid, consistent field names is important because these names become part of the document type schema and are used throughout the extraction workflow, including labeling, model training, validation, and downstream automation. A descriptive identifier such as **First_name123** also avoids ambiguity while complying with the character and starting-character requirements applied by Document Manager. For guidance on creating and managing document types and fields in Document Manager, see the [UiPath Document Manager documentation](#). UiPath's broader documentation for configuring extraction schemas and fields is available in the [Taxonomy Manager guide](#).

QUESTION NO: 11

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

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

ANSWER: E

Explanation:

The correct filter is `"*.doc"`. The **For Each File in Folder** activity's **Filter by** property accepts wildcard patterns to determine which files the activity processes. In this pattern, the asterisk before the period allows any filename, while `.doc*` matches extensions beginning with `.doc`. Consequently, it includes both legacy Microsoft Word files such as `Report.doc` and modern Word files such as `Report.docx`, regardless of the filename that precedes the extension. This is useful when a folder can contain Word documents created in either format and the workflow should handle them in one iteration without requiring separate activities or additional filtering logic. The filter is applied while the activity enumerates the local folder's contents, so only matching files are supplied to the loop body. UiPath documents that this activity iterates through files in a specified folder and provides a filtering property for narrowing the files processed; wildcard matching is the appropriate mechanism for file-name and extension patterns. See the [UiPath For Each File in Folder activity documentation](#) and the [.NET Directory.GetFiles wildcard search-pattern reference](#).

QUESTION NO: 12

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

Explanation:

Before starting the model training. In Communications Mining, a taxonomy establishes the label structure that the model uses to organize, learn from, and classify communications. Importing it before training makes that structure available from the outset, so the dataset and its model can be trained against the intended labels and their hierarchy. This is particularly

useful when a team has already designed, standardized, or exported a taxonomy and wants to reuse it in another dataset rather than recreate the label structure manually. Establishing the taxonomy first also supports consistent annotation and model behavior, because users can apply the agreed business categories while reviewing communications and providing training feedback. The taxonomy is therefore a foundational configuration for the dataset's classification objective, not merely an artifact produced after model training has taken place. UiPath's Communications Mining documentation describes taxonomies as the organization of labels used to classify data and provides guidance for working with them in datasets. See the [UiPath Taxonomy overview](#) and the [UiPath dataset creation guidance](#).

QUESTION NO: 13

What do entities represent in UiPath Communications Mining?

- A. Structured data points.
- B. Concepts, themes, and intents.
- C. Thread properties.
- D. Metadata properties.

ANSWER: A

Explanation:

Structured data points. is correct because entities in UiPath Communications Mining identify specific, extractable pieces of information within a communication's text. They turn unstructured verbatims, such as emails, chats, and call transcripts, into usable structured values. For example, an entity can capture a date, monetary amount, currency, person's name, organization, email address, URL, account identifier, or another business-specific value. This lets an automation team locate and use the precise information mentioned in a message rather than merely recognizing its general subject.

Entity extraction supports operational analysis and downstream automation because the extracted values can be reviewed, filtered, searched, and incorporated into workflows or reporting. For instance, a team can identify messages mentioning a particular supplier, find communications containing dates in a defined period, or extract reference numbers needed to process a request. UiPath distinguishes these concrete data values from the broader categorization of communications, ensuring that models can both understand message content and capture the important data embedded in it. See UiPath's [Entities documentation](#) and [Using entities in your application](#).

QUESTION NO: 14

What can be done in the Reports section of the dataset navigation bar in UiPath Communication Mining?

- A. Train models using unsupervised learning.
- B. View, save, and modify dataset model versions.
- C. Access detailed, queryable charts, statistics, and customizable dashboards.
- D. Monitor model performance and receive recommendations.

ANSWER: C

Explanation:

Access detailed, queryable charts, statistics, and customizable dashboards is correct because the Reports area in UiPath Communication Mining is designed for analyzing the information extracted from a dataset and presenting it in an actionable form. It provides reporting views that help users explore trends, volumes, label and field outcomes, and other dataset-level measures. The available charts can be queried and filtered, enabling users to focus analysis on relevant time periods, sources, labels, or other dataset attributes rather than relying on a fixed summary alone.

Reports also support dashboard-based analysis, so teams can organize useful visualizations for recurring operational, process, or business review. This makes the section appropriate for understanding what the dataset contains, identifying

patterns in communications, and sharing a tailored view of results with stakeholders. The reporting capability is therefore centered on exploration, measurement, and visualization of dataset insights. UiPath describes the purpose and available reporting functionality in its [Communication Mining Reports documentation](#). For broader context on working with datasets and their navigation areas, see the [UiPath Communication Mining Datasets documentation](#).

QUESTION NO: 15

Which of the following options contains the correct list of Default actions that can be found in Workflow Analyzer Settings?

- A. Error, Critical, Info, Warning, Verbose
- B. Exception, Info, Warning, Verbose
- C. Error, Info, Warning, Verbose
- D. Fatal, Error, Warning, Info, Trace

ANSWER: C

Explanation:

Error, Info, Warning, Verbose is the correct list of default actions available in UiPath Workflow Analyzer settings. Workflow Analyzer evaluates automation projects against configured rules and reports each rule's result using an action level. The default action determines the severity assigned when that rule is triggered, helping teams distinguish issues that must be addressed from informational findings. An **Error** represents a validation issue that can prevent a project from passing analysis according to the selected enforcement configuration. **Warning** highlights a potentially undesirable implementation that should be reviewed, while **Info** communicates a lower-severity recommendation or observation. **Verbose** provides the lowest-priority level for detailed reporting. These action levels can be configured for rules in the Workflow Analyzer settings, allowing organizations to apply consistent development standards without treating every finding as a blocking defect. UiPath's Workflow Analyzer documentation describes configuring rule actions and using analysis results to enforce project quality policies. See the [UiPath Workflow Analyzer documentation](#) and the [Workflow Analyzer Settings documentation](#).

QUESTION NO: 16

What is the role of the Taxonomy Manager?

- A. To select which extractors are trained for each document type and field.
- B. To create and edit a Taxonomy file specific to the current automation project.
- C. To select the type of ML that can be used in the project.
- D. To present a document processing specific user interface for validating and correcting automatic classification outputs.

ANSWER: B

Explanation:

Taxonomy Manager is used to create and maintain the document taxonomy for an automation project. A taxonomy defines the business structure that Document Understanding uses when processing files: document types, groups, categories, and the fields that belong to each document type. For example, a project can define an Invoice document type with fields such as invoice number, invoice date, supplier name, and total amount. This consistent structure provides the metadata that downstream Document Understanding activities use to classify documents and extract the intended information.

The taxonomy is stored in a taxonomy file, typically named `taxonomy.json`, associated with the project. Taxonomy Manager provides a visual way to add, edit, organize, and validate those definitions rather than manually constructing the file. The resulting taxonomy can then be supplied to activities such as Taxonomy Manager, Classify Document Scope, Data Extraction Scope, and validation workflows, ensuring that all stages of the automation use the same document and field definitions. UiPath documents Taxonomy Manager as the project-level tool for creating and editing this taxonomy and saving it for use in the automation. See the [UiPath Taxonomy Manager documentation](#) and the [UiPath taxonomy overview](#).

QUESTION NO: 17

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

Explanation:

A higher confidence score. is correct. In UiPath Communications Mining Explore, predicted labels are presented with visual confidence cues so users can quickly assess how strongly the model associates a label with a message. Darker shading indicates that the model has assigned a higher probability, or confidence, to that label prediction. This helps users prioritize their review: prominently shaded predictions represent labels the model considers more certain based on the patterns it learned from the training data. The shading is therefore a compact visual representation of prediction confidence, rather than a statement about whether a prediction has already been verified or whether more than one label is present. Understanding this cue is useful when exploring message sets, validating model behavior, and deciding where human review may add the most value. UiPath documents the Explore experience and the use of predicted labels in its Communications Mining product documentation: [UiPath Communications Mining: Explore](#). For additional context on how the platform applies machine learning to classify communications and generate predictions, see [UiPath Communications Mining overview](#).

QUESTION NO: 18

Can you use Queues in the Document Understanding Process?

- A. The Document Understanding Process can't use Queues because items waiting for Human Validation for more than 10 days will be marked as Abandoned.
- B. The Document Understanding Process can use Queues but the Auto Retry Functionality should be disabled.
- C. The Document Understanding Process can use Queues but the Auto Retry Functionality should be enabled.
- D. The Document Understanding Process can't use Queues because items waiting for Human Validation for more than 24h will be marked as Abandoned.

ANSWER: B

Explanation:

The Document Understanding Process can use Queues but the Auto Retry Functionality should be disabled. Queue-based orchestration is supported by the Document Understanding Process Framework and is useful when documents must be processed as independent transaction items, with queue metadata providing traceability and scalable workload distribution. A document can also require human validation through an Action Center task before downstream processing can continue. In that situation, the process must preserve the transaction's state and wait for the validation outcome rather than treating the pending interaction as a failed transaction to be automatically reprocessed. Disabling automatic retry prevents duplicate attempts and helps ensure that the document, its extracted data, and its associated validation task remain part of one controlled processing path. This configuration is particularly important for reliable exception handling and for avoiding repeated processing while human input is pending. UiPath describes queue-based implementation and configuration in its [Document Understanding Process Framework documentation](#). General queue retry behavior and its implications for transaction processing are documented in the [UiPath Orchestrator queues documentation](#).

QUESTION NO: 19

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

Explanation:

2 is correct. The Document Understanding Process template uses two synchronization approaches when several robots may work concurrently: file locking for access to shared files and Orchestrator queues for safely coordinating work items. A file lock provides exclusive access while the protected operation runs, so only one robot can read from or write to the targeted file or folder at a given time. This prevents simultaneous updates from overwriting each other or producing inconsistent file contents. Queue-based coordination complements this approach by moving document-related work and extracted data through Orchestrator queue items, which can be independently retrieved and processed by available robots. Queues provide a durable, centrally managed work-distribution mechanism and support transaction processing across multiple robot executions. Together, these two mechanisms address shared-resource synchronization and scalable workload coordination in the template's document-processing flow. UiPath documents the locking behavior in its [File Lock Scope activity documentation](#), and describes queue-item processing and Orchestrator-backed transactions in its [Add Queue Item activity documentation](#).

QUESTION NO: 20 - (SIMULATION)

What is the order of steps for automatically retraining and deploying a Document Understanding ML Model in AI Center with data from Document Validation Action?

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

Steps	Order of Steps
☐ Enable the Auto Update feature in the ML Skill.	First
☐ Use the Export feature from Document Manager using Scheduled Export.	Second
☐ Run Training/Full Pipeline in AI Center using Scheduled Pipeline with Auto-Retraining enabled.	Third
☐ Send Human in the Loop data to the used Dataset.	Fourth

ANSWER: See the explanation for the answer

Explanation:

The correct order is:

1. Send Human in the Loop data to the used Dataset.
2. Use the Export feature from Document Manager using Scheduled Export.
3. Run Training/Full Pipeline in AI Center using Scheduled Pipeline with Auto-Retraining enabled.

4. Enable the Auto Update feature in the ML Skill.

Human-validated/corrected documents are first made available in the dataset. Document Manager then exports the data on a schedule. The scheduled AI Center training (or full) pipeline retrains the model using that data, and **Auto Update** on the ML Skill automatically deploys the newly trained model version.

QUESTION NO: 21

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

Explanation:

DocumentUnderstanding is the appropriate generic ML Package when the required document type is not covered by UiPath's prebuilt, out-of-the-box document models. This package is intended for custom document-processing use cases: it provides a trainable framework that can be adapted with labeled examples from the organization's own documents. Teams can use it to build extraction capabilities for fields and layouts specific to their business, including structured and semi-structured document formats that do not match a dedicated prebuilt model.

The key benefit is that the model can learn from the document samples and labels supplied during the training workflow, rather than relying on a fixed schema designed for a particular common document type. This supports scenarios in which an organization uses proprietary forms, regional variants, industry-specific records, or supplier-specific layouts. After the model is trained and deployed through AI Center, it can be consumed in a Document Understanding workflow to extract the relevant information at runtime. UiPath documents the available ML packages and their deployment lifecycle in [AI Center ML Packages](#), and describes how custom models fit into document-processing implementations in the [Document Understanding documentation](#).

QUESTION NO: 22

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: D

Explanation:

A Pipeline can be scheduled at a single future date or with a recurring schedule. In AI Center, pipeline scheduling supports both one-time and repeated executions, allowing teams to automate activities such as model training, evaluation, or deployment without manually starting each run. A one-time schedule is appropriate when the pipeline must execute at a particular upcoming date and time, such as after a planned data refresh. A recurring schedule is appropriate when the same process must run repeatedly according to a defined cadence, helping keep machine-learning workflows current as new data becomes available. This scheduling capability is configured when creating or managing the pipeline run schedule, so the execution is initiated automatically by AI Center at the selected time or recurrence pattern. The key distinction is that a future schedule represents one planned execution, while a recurring schedule establishes an ongoing automated series of executions. This makes the stated combination the complete description of the supported scheduling approaches. See

QUESTION NO: 23

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

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

ANSWER: C

Explanation:

Validation. is correct because the Document Understanding Framework commonly uses a human-in-the-loop validation step to review and correct extraction results before those corrected results are used to improve a model. During validation, a user confirms document data and fixes any inaccurate fields, producing reliable, labeled information. This feedback is valuable training data: it connects the document content with the expected extracted values and can be incorporated into a retraining workflow for supported machine-learning extractors.

In practice, the automation sends documents that require review to a validation action, waits for completion, and retrieves the validated extraction results. Those reviewed results can then be exported or otherwise prepared as a dataset for model training. Training after validation helps ensure that a model learns from verified labels rather than unreviewed predictions, supporting progressively better extraction quality as more representative documents are processed. UiPath documents the overall Document Understanding process and its validation capabilities in the [Document Understanding Framework documentation](#). For details on the human review component that captures corrected data, see UiPath's [Validation Station documentation](#).