

UiPath Specialized AI Associate Exam (2023.10)

UiPath UiPath-SAIAv1

Version Demo

Total Demo Questions: 27

Total Premium Questions: 273

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, AI Center	42
Topic 2, Document Understanding	117
Topic 3, Communications Mining	66
Topic 4, Mix Questions	48
Total	273

QUESTION NO: 1

How is the Taxonomy component used in the Document Understanding Template?

- A. To define the document types and the pieces of information targeted for data extraction (fields) for each document type.
- B. To assign predefined document categories based on content similarity, simplifying classification tasks for easier document organization.
- C. To convert scanned documents into machine-readable formats, utilizing taxonomy rules to enhance digitized content accuracy.
- D. To automatically extract structured data fields from documents, leveraging taxonomy mappings for precise data extraction.

ANSWER: A

Explanation:

The Taxonomy component is used to define the document-processing structure: the document types that the automation will handle and the fields to be extracted for each type. In UiPath Document Understanding, a taxonomy is the central schema that describes document groups, document types, and their associated fields, including field names and relevant configuration such as data type or whether a field is required. The Document Understanding Template uses this definition as the common contract for the subsequent stages of the workflow. Classifiers use the defined document types to identify incoming documents, while extractors and validation processes use the defined fields as the target data model. This lets the solution process several kinds of business documents consistently while ensuring that extracted results have a predictable structure. For example, an invoice document type can contain fields such as invoice number, invoice date, supplier name, and total amount. The taxonomy is created and managed through Taxonomy Manager and can be imported into the Document Understanding process, making it reusable and maintainable as document requirements evolve. UiPath describes Taxonomy Manager as the place to define document types and fields for Document Understanding projects. See the [UiPath Taxonomy Manager documentation](#) and the [Document Understanding Framework overview](#).

QUESTION NO: 2

What new capability has been introduced for processing unstructured documents in the 2023.10 release?

- A. Real-time document translation.
- B. Only using pre-trained extraction models.
- C. Manual classification of document types.
- D. Generative capabilities for classifying various document types.

ANSWER: D

Explanation:

Generative capabilities for classifying various document types is correct. In the 2023.10 release, UiPath introduced Generative Classifier functionality in Document Understanding to address unstructured-document scenarios. Rather than requiring a fixed, preconfigured set of document layouts or relying solely on traditional classification approaches, the generative AI capability can use natural-language descriptions of document types to identify and classify incoming documents. This is particularly valuable where document formats vary substantially, where the contents are more narrative than structured, or where an organization needs to process multiple document categories without creating extensive rule sets first.

Within a Document Understanding workflow, classification is an important early step because it routes each document to the appropriate downstream extraction and validation process. Generative classification expands this stage by applying generative AI to infer the appropriate document category from its content and the categories defined by the automation designer. UiPath's release notes describe the introduction of generative AI capabilities for unstructured documents, and the product documentation explains the purpose and configuration of the Generative Classifier. See the [UiPath 2023.10 release notes](#) and [Generative Classifier documentation](#).

QUESTION NO: 3

Which activity enables the identification of the document type by using any classifier?

- A. Digitize Document activity.
- B. Train Classifiers Scope activity.
- C. Present Classification Station activity.
- D. Classify Document Scope activity.

ANSWER: D

Explanation:

The **Classify Document Scope activity** is correct because it is the Document Understanding activity designed to determine a document's type by orchestrating one or more configured classifiers. Within its scope, you add the applicable classifier activities, such as an Intelligent Form Classifier, Keyword Based Classifier, or Machine Learning Classifier. The scope sends the digitized document and its data to those classifiers, then produces classification results identifying the document type and associated confidence information. This classification result can be routed to later stages of a Document Understanding workflow, including data extraction and validation. The activity supports combining classifiers so that a workflow can use the classifier best suited to the documents it receives, while providing a single scope for configuring and executing the classification step. UiPath's Document Understanding process places classification after digitization and before extraction, making the Classify Document Scope activity the appropriate activity for identifying document types. See UiPath's [Classify Document Scope documentation](#) and its overview of the [Document Understanding framework](#).

QUESTION NO: 4

Which UiPath Communications Mining model performance factor relates to the proportion of messages in the dataset that have informative label predictions?

- A. Coverage.
- B. Balance.
- C. Average label performance.
- D. Underperforming labels.

ANSWER: A

Explanation:

Coverage. is the model performance factor that measures the proportion of messages in a dataset for which the model produces informative label predictions. In UiPath Communications Mining, coverage indicates how broadly the trained model can meaningfully classify the available communications, rather than leaving messages without useful predicted labels. High coverage means that the taxonomy and training data enable the model to identify relevant labels across a large share of the dataset, helping users turn more incoming messages into actionable, structured information. It is therefore particularly useful when assessing whether a model has been trained sufficiently to handle the variety of language, intents, and message types present in production data. Coverage should be considered alongside prediction quality so that the model is both applicable to a substantial portion of messages and provides dependable classifications where it is applicable. UiPath's Communications Mining documentation describes the model evaluation metrics and how they help assess a model's readiness and improvement opportunities. See the [UiPath model performance documentation](#) and the [UiPath training documentation](#).

QUESTION NO: 5

Which of the following data structures in a UiPath workflow allow dynamic resizing, making it suitable for scenarios where the number of elements is not predetermined?

- A. Integer
- B. Tuple
- C. List
- D. Array

ANSWER: C

Explanation:

List is the appropriate data structure when a UiPath workflow needs a collection whose number of elements is not known in advance. UiPath uses .NET collection types, and a generic `List<T>` can grow or shrink during execution as items are added or removed. This makes it practical for automation scenarios such as accumulating extracted records, storing files found in a folder, or collecting transaction items returned from an application.

In UiPath Studio, a List variable is created by selecting a type such as `System.Collections.Generic.List<String>`, `List<Int32>`, or another required item type. Activities such as *Add To Collection* can append values to the List while the workflow runs. A List also preserves item order and provides indexed access, which supports subsequent iteration with activities such as *For Each*. UiPath's documentation describes collection variables and their use in workflows, while Microsoft's .NET documentation confirms that `List<T>` represents a strongly typed list of objects accessible by index and supports dynamically adding and removing elements.

References: [UiPath Studio — Managing Variables](#); [Microsoft Learn — List<T> Class](#).

QUESTION NO: 6

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. The UiPath Document Understanding Process project template uses an Orchestrator queue to hold the input documents for processing, typically with one input file represented by one queue item. This queue-based approach supports scalable, reliable processing because robots can retrieve document work independently and process multiple items according to the configured execution model.

Auto Retry must be disabled for the queue used by this template. Document Understanding can pause a transaction while a user completes a validation task in Action Center. Retrying the queue item automatically during that human-in-the-loop stage can start duplicate processing or create conflicts with the template's own transaction-state and exception-handling flow. Instead, the project manages processing errors and retries through its designed workflow logic, preserving the intended coordination between document extraction, validation, and final results. UiPath documents the template and its queue-based input processing in the [Document Understanding Process documentation](#). For the queue setting itself, see UiPath's [Managing Queues documentation](#).

QUESTION NO: 7

For an analytics use case, what are the recommended minimum model performance requirements in UiPath Communications Mining?

- A. Model Ratings of "Good" or better and individual performance factors rated as "Good" or better.
- B. Model Ratings of "Good" and individual performance factors rated as "Excellent".
- C. Model Ratings of "Excellent" and individual performance factors rated as "Good" or better.
- D. Model Ratings of "Excellent" and individual performance factors rated as "Excellent".

ANSWER: A

Explanation:

Model Ratings of "Good" or better and individual performance factors rated as "Good" or better is the recommended minimum for an analytics use case. Communications Mining evaluates a trained model through an overall Model Rating and through individual performance factors that indicate how reliably the model identifies and classifies the concepts needed for the use case. For analytics, the model must be sufficiently dependable both overall and in each relevant factor so that trends, volumes, and other reported insights are based on consistently meaningful classifications. A rating of Good establishes the recommended baseline quality threshold; a higher rating is also acceptable, hence the phrase "or better." Applying the same Good-or-better threshold to the individual factors is important because a strong overall result alone does not ensure that every dimension used in analysis has adequate quality. This recommendation supports using model outputs to understand communication themes and operational patterns with an appropriate level of confidence. UiPath describes model performance ratings and the factors used to assess model quality in its [Model performance documentation](#). Additional guidance on evaluating and improving model quality is available in the [Improving model performance documentation](#).

QUESTION NO: 8

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: D

Explanation:

Tools for message analytics and monitoring is correct because the Reports section in UiPath Communications Mining is designed to help users understand the communications contained in a dataset and observe how those communications change over time. Reports provide analytical views of message data, including volumes and trends, and allow users to segment and filter the dataset to investigate operational patterns. This makes the section useful for monitoring incoming communication, identifying changes in demand or topic distribution, and sharing data-driven insight with stakeholders. Rather than being limited to the model-building workflow, Reports focuses on the content and characteristics of the messages that the model is processing. It therefore supports ongoing analysis after data has been ingested and classified, helping teams track the communications landscape and make informed process decisions. UiPath describes reporting as part of the Communications Mining dataset experience, with dedicated reports and analytics capabilities for exploring message data. See the [UiPath Communications Mining Reports documentation](#) and the [UiPath Communications Mining datasets documentation](#).

QUESTION NO: 9

What are the options available in the Export Now tab of the Export Files dialog box in Document Manager?

- A. Download to Excel, Download, and Export to AI Center.

B. Download to Excel, Export to AI Center, and Export All.

C. Download to Excel and Export All.

D. Export to AI Center, Export All. and Download.

ANSWER: A

Explanation:

Download to Excel, Download, and Export to AI Center is correct because these are the three export destinations/actions provided through the *Export Now* workflow in Document Manager. Download creates a local export package of the dataset files and associated labeling data, which can be retained, moved, or used outside the UiPath platform. Download to Excel provides the exported dataset information in an Excel-friendly format, supporting inspection and reporting. Export to AI Center sends the dataset directly to AI Center, where it can be used in the Document Understanding machine-learning workflow, such as creating a dataset for training or evaluating a model. These choices support the main immediate-export scenarios: local file export, spreadsheet-oriented export, and direct handoff to UiPath's AI model-management service. The relevant Document Understanding guidance is available in UiPath's [Exporting documents documentation](#). For the role of datasets in the wider AI workflow, see the UiPath [AI Center datasets documentation](#).

QUESTION NO: 10 - (SIMULATION)

What is the correct order of migrating a dataset from Document Manager to a Modern Project? Instructions: Drag the Description found on the left and drop on the correct Step found on the right.

ANSWER: See the explanation for the answer

Explanation:

Correct order:

1. Go to the document type you want to export and select **Open document type**.
2. From the **Filter documents** drop-down list, select **Training and validation set**, then select **Export**.
3. Leave **Current search results** selected, enter a name for the export job, and select **Download**.
4. Navigate to and open the Modern project into which the dataset will be imported. Select an existing custom document type or create a new one.
5. Select **Upload**, choose the ZIP file exported from the classic project, and wait for the upload to complete.

The dataset must first be exported and downloaded from the Document Manager/classic project, then uploaded into the selected document type in the Modern project.

QUESTION NO: 11

A developer configured the UI Automation Project Settings and the Properties of a Click activity as shown in the following exhibits:

If the target element is not found during execution in Run mode, how long will it take until an error is thrown (based on default project settings)?

A. 0.15

B. 0.2

C. 15

D. 30

ANSWER: C

Explanation:

15 is correct because the Click activity's configured target-search timeout is 15,000 milliseconds, which equals 15 seconds. When the activity runs and cannot identify its target, UiPath continues attempting to find that target until the applicable timeout elapses. At that point, the activity fails and the workflow raises an error unless the exception is handled elsewhere in the automation.

For UI Automation activities, timeout values are expressed in milliseconds. The activity-level timeout is the relevant value when it has been explicitly configured, rather than relying on the project's default UI Automation timeout. Run mode does not change a 15,000-millisecond activity timeout into a delay setting; delay-before and delay-after values control when the action is performed around the interaction, not how long UiPath searches for a missing UI element.

UiPath documents the Click activity and its target-related configuration at [UiPath Click activity documentation](#). The relationship between project-level UI Automation defaults and activity behavior is described in [UiPath UI Automation project settings](#).

QUESTION NO: 12

What are the three types of classifier trainers available in packages UiPath.IntelligentOCR.Activities and UiPath.DocumentUnderstanding.ML.Activities?

- A. Intelligent Keyword Classifier Trainer, Language Based Classifier Trainer, and Image Based Classifier Trainer.
- B. Image Based Classifier Trainer, Format Based Classifier Trainer, and Machine Learning Classifier Trainer.
- C. Machine Learning Classifier Trainer, Language Based Classifier, and Keyword Based Classifier Trainer.
- D. Keyword Based Classifier Trainer, Intelligent Keyword Classifier Trainer, and Machine Learning Classifier Trainer.

ANSWER: D**Explanation:**

Keyword Based Classifier Trainer, Intelligent Keyword Classifier Trainer, and Machine Learning Classifier Trainer is correct because these are the classifier-training activities provided for training and improving document classification in UiPath Document Understanding. The Keyword Based Classifier Trainer is used to configure and train classification based on relevant document keywords. The Intelligent Keyword Classifier Trainer supports a more intelligent keyword-based approach, using labeled documents to learn the keywords and their significance for each document type. The Machine Learning Classifier Trainer is supplied through the Document Understanding ML activities and trains an ML classifier using labeled classification data, supporting scenarios in which classification requires learned patterns beyond a manually maintained keyword list. Together, these trainers correspond to UiPath's supported approaches for creating classification training data and improving classifier performance within a Document Understanding workflow. Their outputs can be persisted and subsequently supplied to the corresponding classifier activities when processing new documents. UiPath's activity documentation describes the keyword-based trainer activity and its role in training a keyword classifier, while the ML activity documentation covers the machine-learning-based classification workflow. See [UiPath Keyword Based Classifier Trainer documentation](#) and [UiPath ML Classifier Trainer documentation](#).

QUESTION NO: 13

Which UiPath Communications Mining model performance factor assesses the proportion of the entire dataset that has informative label predictions?

- A. Average label performance.
- B. Coverage.
- C. Balance.
- D. Underperforming labels.

ANSWER: B

Explanation:

Coverage. is the Communications Mining model-performance factor that measures the proportion of the overall dataset for which the model produces informative label predictions. In practical terms, it indicates how much of the communication data is meaningfully represented by the model's taxonomy, rather than being left without useful predicted labels. Strong coverage means the model can identify relevant intents, topics, or concepts across a broad share of the communications being analyzed.

Coverage is especially useful when assessing whether a model's labels adequately reflect the real range of requests and information within a dataset. A model can have well-performing individual labels but still offer limited operational value if many communications do not receive informative predictions. Reviewing coverage therefore helps teams determine whether additional training data, improved label definitions, or new labels are needed to ensure the model captures the major patterns in the data. UiPath includes coverage among the measures used to evaluate the overall health and usefulness of a Communications Mining model. See UiPath's [Model rating documentation](#) and the [guidance for improving a model](#).

QUESTION NO: 14

In a Document Understanding project, the user needs to extract information from PDF documents with the following requirements:

The documents can contain scanned or digitally typed text.

The documents can contain checkboxes, and these must be extracted.

The automation must use the logical processors in the most efficient way to obtain the maximum degree of parallelism. What are the properties provided to the Digitize Document activity in the Digitize phase?

- A. ApplyOcrOnPdf -> Auto
- B. ApplyOcrOnPdf -> Auto DegreeOfParallelism -> -1 DetectCheckboxes -> True
- C. ApplyOcrOnPdf -> No DegreeOfParallelism -> True DetectCheckboxes -> True
- D. ApplyOcrOnPdf -> Auto DegreeOfParallelism -> Max DetectCheckboxes -> True
- E. ApplyOcrOnPdf -> No DegreeOfParallelism -> -1 DetectCheckboxes -> False

ANSWER: B

Explanation:

ApplyOcrOnPdf -> Auto DegreeOfParallelism -> -1 DetectCheckboxes -> True is the appropriate configuration for these requirements. Setting *ApplyOcrOnPdf* to *Auto* lets Digitize Document determine whether OCR is needed for each PDF. This supports mixed inputs efficiently: digitally generated PDFs can use their existing text layer, while scanned PDFs, which consist of image content, can be processed through OCR to produce usable document text and structure.

Enabling *DetectCheckboxes* ensures that checkbox controls are identified during digitization. The resulting checkbox information is then available to later Document Understanding activities and extraction workflows, allowing the automation to capture whether a box is selected or unselected.

A *DegreeOfParallelism* value of *-1* directs the activity to use all available logical processors. This is the setting intended to maximize parallel processing capacity without manually setting a fixed processor count, making it suitable when the goal is the highest practical degree of parallelism on the executing machine. UiPath documents these settings in the [Digitize Document activity reference](#); broader digitization workflow context is available in the [Document Understanding digitization documentation](#).

QUESTION NO: 15

Which of the following is a type of communication that is typically interpreted by UiPath Communications Mining?

- A. Scanned letters
- B. Shared email inboxes
- C. Call data in real-time
- D. Real-time chat data

ANSWER: B

Explanation:

UiPath Communications Mining is designed to interpret high volumes of unstructured, text-based business communications, making **Shared email inboxes** the correct answer. Organizations commonly use shared mailboxes for customer support, operations, claims, complaints, and service requests. Communications Mining ingests these messages and applies machine-learning models to identify intent, classify messages into categories, and extract relevant information such as customer details, reference numbers, or requested actions. This turns inbox content that would otherwise require substantial manual review into structured, actionable data that can support automation and operational decision-making.

The product is particularly suited to communications that are already available as written records and can be connected through supported data sources. After messages are imported, users can train and refine models using annotations, then use the resulting classifications and extracted fields in downstream UiPath automations or analytics workflows. UiPath describes Communications Mining as a platform for discovering and automating processes from communications data, including email-based customer interactions. See the [UiPath Communications Mining overview](#) and the [UiPath documentation on data sources](#) for details.

QUESTION NO: 16

In UiPath Communications Mining, which phase is the starting point of the model training process, where similar intents and conversation themes are grouped?

- A. Refine
- B. Explore
- C. Discover
- D. Setup

ANSWER: C

Explanation:

Discover is the starting phase for building a Communications Mining model. In this phase, Communications Mining analyzes the uploaded message data and identifies recurring language, themes, and potential intents. Similar communications are grouped into clusters, allowing users to inspect representative examples and recognize the business processes, requests, or issues that occur in the dataset.

This discovery work provides the foundation for the taxonomy used in model training. Users can turn the patterns found in clusters into meaningful labels and then apply those labels to relevant messages. The labeled examples teach the model how to recognize the same intents and themes across previously unseen communications. Beginning with discovery is therefore important because it helps teams create labels grounded in the actual content and terminology present in their data, rather than relying only on assumptions about what customers or employees may be discussing.

UiPath describes Communications Mining as a platform for discovering and extracting insights from unstructured communications and for training models to classify those communications. For further guidance, see the [UiPath Communications Mining getting started documentation](#) and the [UiPath model creation and training documentation](#).

QUESTION NO: 17

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

Explanation:

During the Classify stage, Document Understanding determines the type of each digitized document by matching it to a document type defined in the project taxonomy. The taxonomy supplies the structure used by the workflow, including the available document types and the fields relevant to those types. Using configured classification methods through activities such as Classify Document Scope, the workflow evaluates document content and layout and returns classification results in a unified format. A document can consequently be assigned to one of the taxonomy document types when its classification meets the configured criteria, or it can be skipped when it cannot be reliably classified or is not eligible for a selected document type. This decision is essential because the assigned document type determines which extraction configuration and target fields are used in the subsequent data-extraction stage. Classification can also use confidence thresholds and routing choices so that uncertain cases can be handled according to the automation's design. UiPath describes the processing flow and the role of classification in its [Document Processing Framework documentation](#). Details on configuring classifiers are available in the [Classify Document Scope activity documentation](#).

QUESTION NO: 18

What type of variable is used to store information about a duration in UiPath?

- A. String
- B. System.DateTime
- C. Integer
- D. System.TimeSpan

ANSWER: D

Explanation:

System.TimeSpan is the appropriate variable type for storing a duration, meaning an elapsed interval of time rather than a particular calendar date and clock time. A TimeSpan value can represent intervals such as five minutes, two hours and thirty minutes, or several days. In UiPath Studio, variables use .NET types, so System.TimeSpan can be selected as a variable type and used in expressions and activities that calculate, compare, add, or subtract durations. For example, subtracting one DateTime value from another produces a TimeSpan representing the time elapsed between those moments. A TimeSpan can also be created with expressions such as `TimeSpan.FromMinutes(30)`, which is useful when a workflow needs to define timeout periods, measure processing durations, or perform date/time arithmetic. UiPath's documentation explains that Studio supports .NET data types for variables, including time-related types, while Microsoft's .NET documentation defines TimeSpan specifically as a time interval. See [UiPath documentation on managing variables](#) and [Microsoft's System.TimeSpan reference](#).

QUESTION NO: 19

What additional property does the ML Extractor have compared to the other types of extractors?

- A. Timeout
- B. Endpoint
- C. ApiKey

ANSWER: D

Explanation:

ML Skill is correct because the ML Extractor activity uses a deployed Machine Learning skill to process document data and return extracted fields. Unlike rule-based or prebuilt extractor configurations, an ML Extractor must be linked to the specific ML skill that hosts the trained model to be used for extraction. The ML Skill property identifies that model endpoint within UiPath AI Center or the configured ML Skills environment, allowing Document Understanding to send the document and relevant input to the appropriate trained skill.

This configuration is essential because an ML model's behavior, supported document types, and extracted fields depend on the particular skill selected. For example, an organization can deploy separate skills for invoices, purchase orders, or specialized custom document classes, then select the required skill in the ML Extractor configuration. UiPath's Document Understanding framework combines this machine-learning extraction capability with validation and downstream processing, while the selected ML skill supplies the model intelligence used to identify and extract data. See UiPath's [ML Extractor documentation](#) and [ML Skills documentation](#).

QUESTION NO: 20

Which environment variable is relevant for Evaluation pipelines?

- A. eval.enable_ocr
- B. eval.redo_ocr
- C. eval.enable_qpu
- D. eval.use_cuda

ANSWER: B

Explanation:

`eval.redo_ocr` is the environment variable relevant to Evaluation pipelines. It is used when evaluating a Document Understanding extraction model in an AI Center pipeline and instructs the evaluation process to run OCR again on the input documents. This makes it possible to measure extraction performance with newly generated OCR results, rather than relying only on OCR information already associated with the documents. The setting is particularly useful when assessing how OCR quality affects downstream field extraction accuracy, such as after changing or improving the OCR configuration. For this parameter to have practical value, an OCR engine must have been configured for the ML package or pipeline scenario so that the pipeline has an OCR capability to invoke. Evaluation pipelines are intended to produce metrics that help assess an ML package's performance against labeled data; controlling whether OCR is redone helps ensure that the evaluation setup reflects the intended document-processing conditions. UiPath documents evaluation-pipeline configuration and supported parameters in its [Evaluation pipelines documentation](#), with broader pipeline configuration guidance available in the [ML pipelines documentation](#).

QUESTION NO: 21

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:

To create and edit a Taxonomy file specific to the current automation project is correct because Taxonomy Manager is the Document Understanding design tool used to define the structure of documents that an automation processes. In Taxonomy Manager, developers build and maintain the project taxonomy by defining document types, groups, categories, and the fields associated with each document type. This structure provides a consistent business model for the document-processing workflow: for example, it can define that an invoice document contains fields such as invoice number, invoice date, supplier name, and total amount.

The saved taxonomy is used by Document Understanding activities and processes to identify the relevant document types and data fields to classify and extract. Maintaining it at the project level lets the automation use terminology and field definitions that match its specific business requirements, while also making those definitions reusable throughout the workflow. UiPath describes Taxonomy Manager as the place to create, import, and manage taxonomies for Document Understanding projects. For further details, see the [UiPath Taxonomy Manager documentation](#) and the [UiPath taxonomy overview](#).

QUESTION NO: 22

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

Explanation:

It could increase model bias. In Communications Mining, Search is useful for finding specific communications, terms, or patterns, but it should be used sparingly during training because search-selected examples are not necessarily representative of the overall dataset. Repeatedly searching for familiar keywords or known phrasing can concentrate annotations around a narrow subset of communications. As a result, the model may learn associations that reflect the search terms and the annotator's selection pattern more strongly than the naturally occurring variety in the data.

Training from a broader, more randomly presented set of communications helps the model encounter different language, contexts, and edge cases. This produces labels that better represent the population of messages the model will evaluate in production. Limiting Search therefore supports a less biased training set and improves the model's ability to generalize to unseen communications, rather than performing well only on wording similar to examples deliberately located through searches. UiPath's Communications Mining guidance describes model training practices and the importance of representative data when improving model performance. See the [UiPath Communications Mining training documentation](#) and the [Communications Mining overview](#).

QUESTION NO: 23 - (SIMULATION)

What is the correct order of uploading a package exported from UiPath AI Center?

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:

Correct order:

1. Export the package from AI Center.
2. On the **ML Packages** page, click **Import ML Package**.
3. In the **Upload package** field, select/upload the exported .zip package file.
4. Click **Create**.

This exports the ML package first, then begins an import, supplies the downloaded ZIP archive, and finally creates the imported package.

QUESTION NO: 24

What does the Document Classification step do?

- A. Identifies what type of document the robot is currently processing.
- B. Presents a document processing-specific user interface for validating and correcting automatic classification outputs.
- C. Empowers the closing of the feedback loop to any classification algorithm capable of learning.
- D. Retrieves the text from any PDF or image, using, only if necessary, the OCR engine.

ANSWER: A

Explanation:

Identifies what type of document the robot is currently processing. Document Classification is the Document Understanding stage that determines a document's type, such as an invoice, receipt, purchase order, or another predefined document category. This determination is essential because it lets the automation route the document to the appropriate downstream extraction logic, including the taxonomy and extraction model configured for that document type. Classification can use one or more classifiers and can evaluate either an entire document or pages within a document, depending on the workflow design. The resulting classification information is then available to subsequent Document Understanding activities so that the robot can extract the relevant fields using the correct schema. UiPath describes classification as the process of identifying the document type before data extraction, helping automations handle mixed document sets consistently and at scale. See the UiPath [Document Classification documentation](#) and the [Document Understanding activities overview](#) for the role of classification in a document-processing workflow.

QUESTION NO: 25

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

Explanation:

.doc is the appropriate filter because the asterisk is a wildcard that matches any sequence of characters following the .doc extension prefix. Consequently, this pattern includes traditional Microsoft Word .doc files and modern .docx files, allowing one For Each File in Folder activity to process the Word documents commonly stored in a local folder. It can also match other Word-related extensions that begin with .doc, such as macro-enabled document formats, which is useful when the requirement is to iterate broadly through Microsoft Word document files rather than restrict processing to a single Word file format. The For Each File in Folder activity provides a Filter property for selecting files by a matching pattern before the loop processes them. UiPath's activity documentation describes this activity and its file-filtering capability at [UiPath For Each File in Folder](#). The wildcard behavior used by file-search patterns is also documented by Microsoft at [Directory.GetFiles](#)

[method](#). Using `*.doc*` therefore efficiently targets the relevant Word document extensions without requiring separate iterations.

QUESTION NO: 26

How should the data be managed after extraction in the UiPath Document Understanding process?

- A. Directly pass the extracted data to other processes via Arguments.
- B. Serialize the data and store it in a storage bucket or similar location.
- C. Store the data temporarily in local folders before processing.
- D. Keep the data only in the queue items without serialization.

ANSWER: B

Explanation:

Serialize the data and store it in a storage bucket or similar location. is the appropriate approach when extracted Document Understanding data must be retained and used by later workflow stages, jobs, or processes. Extraction results and the document object model are structured runtime objects; serializing them converts relevant information into a durable format that can be persisted and reconstructed or consumed later. This is especially useful for long-running or multi-stage automations, where the extraction step and downstream processing may not run in the same execution context. A UiPath Storage Bucket provides centrally managed, Orchestrator-connected storage for files and data used by automations, avoiding dependence on a particular robot machine. The workflow can also export the extraction results into formats that are convenient for downstream business processing, such as a DataSet or other structured output, before persisting the needed artifact. UiPath documents the [Export Extraction Results activity](#) for converting Document Understanding extraction output into usable structured data. For persistent robot-accessible storage, see UiPath's [Storage Buckets documentation](#). This pattern supports reliable handoff, traceability, and reuse of extracted information.

QUESTION NO: 27

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

Explanation:

The page number from which the field was extracted and the exact position on the page is correct because Document Understanding extraction output can retain metadata that identifies where a reported field value came from in the source document. This provenance information includes the page containing the extracted value and its bounding-region coordinates, which describe the value's location on that page. Including this metadata makes an exported result more useful for validation, troubleshooting, audit scenarios, and downstream processes that need to link a structured value back to its visual source. For example, a reviewer can use the page and bounding-box information to quickly locate an invoice total or supplier name in the original document, rather than searching through the document manually. UiPath describes extraction results as including field-level information and document references that support validation and traceability. The exact output structure can vary with the extraction and export workflow, but location metadata is the relevant additional information described here. See UiPath's [Export extraction results documentation](#) and its [Document Understanding processes documentation](#) for details on handling extracted document data and associated metadata.