



AI Studio
User Manual
Version 1

Copyright © 2026 Computer Talk Technology, Inc. All rights reserved.

No part of this publication may be reproduced, transmitted, or translated in any form or by any means, electronic, mechanical, manual, optical or otherwise, including photocopying, recording, or any information storage and retrieval system, without the prior permission in writing from Computer Talk Technology, Inc.

ComputerTalk Trademarks

ice, iceAdministrator, iceAlert, iceBar, iceBar for web, iceBar for Teams© Mobile, iceCampaign, iceChat, iceJournal, iceManager, iceMobile Connect, iceMonitor, icePay, icePhone, iceReporting, iceSurvey, iceWorkflow Designer, Say My Name, and AI Studio are trademarks of ComputerTalk Technology, Inc.

Microsoft, Excel, and Windows are either registered trademarks or trademarks of Microsoft Corporation in the United States and/or other countries.

Adobe, Acrobat, and Reader are either registered trademarks or trademarks of Adobe Systems Incorporated in the United States and/or other countries.

All other company and product names used herein may be the trademarks or registered trademarks of their companies.

Microsoft product screen shot(s) reprinted with permission from Microsoft Corporation.

Part Number: UM_AISM_20260804

AI Studio Version 1.0

Table of Contents

Chapter 1: Getting Started	3
Single Sign-On.....	7
Signing On with Single Sign-On.....	7
AI Settings	9
Common Error and Warning Messages.....	10
Authentication error	10
Loss of connection	11
API Keys.....	11
Chapter 2: Knowledge Management	15
Knowledge Base.....	16
Sources	16
Import	18
Searching for a source	25
Verified Answers	27
Right-Click Menu	28
Detail View	29
Edit an answer group response.....	30
Add an alternate question.....	31
Edit an alternate question	32
Add New Question.....	33
Delete a question-answer pair.....	34
Searching for a question.....	35
Import	36
Export	37
Refresh.....	38
Simulation Lab	39
Compare Mode	39
Singular Testing.....	41
Chapter 3: Natural Language Routing.....	45
Project Management.....	46
Project.....	47
Authoring.....	48
Intents	49
Entities	54
Utterance	60
Training Jobs	69
Model Performance	72
Overview	73

Model Performance	76
Test Set Details	77
Dataset Distribution.....	78
Confusion Matrix	78
Deployments	80
Playground.....	82
Single Query	82
Multi-turn Sandbox.....	84
Batch Testing.....	88
Chapter 4: AI Agents	90
AI Agents	91
Configuration	91
Categories	93
Category Breakdown.....	96
Knowledge Gaps.....	98
Creating an AI Agent.....	102
Chapter 5: Dashboards.....	104
Platform View.....	105
Bot Performance	106
Trends	107
Interactions Over Time	107
Knowledge Agent Answer Mode Distribution.....	108
Confidence and Success Rate	109
Escalation Rate Chart	110
Categories and Sentiment.....	111
Category Breakdown.....	111
Sentiment Breakdown.....	112
Interactions	113
Single Agent View	115
FAQ Agent Metrics.....	116
Overview	116
Answer Quality.....	117
Categories	120
Interactions	120
Natural Language Routing Metrics.....	120
Overview	120
Intent Breakdown	121
Routing.....	122
Interactions	123
Automation Agent.....	123
Overview	123
Journey	124

Outcomes	126
Interactions	127



Welcome to iceAI Studio

As AI bots become more prevalent in today's business world, many contact centers are increasingly adopting the use of voicebots and chatbots to answer common questions and enable customers to perform self-service tasks.

iceAI Studio is a tool that allows you to manage your knowledge base, your conversational language understanding model or your verified answers, all from one convenient location.

This manual is designed for administrators who contribute to the configuration and management of AI Agents in your contact center. Depending on the type of bot you have purchased, an administrator can be responsible for updating the knowledge base and verified answers, or managing the intents, entities and training data for Natural Language Routing.

This manual covers several topics:

- **Chapter 1: Getting Started** provides important information regarding AI Studio, common error and warning messages, and AI settings.
- **Chapter 2: Knowledge Base** describes how to manage the knowledge base used for genAI bots. Administrators can use the simulation lab to test the bot, and compare genAI answers to verified question-answer pairs.
- **Chapter 3: Natural Language Routing** describes how to manage the intents, entities and training data used to develop the natural language understanding model used for the NLR bot. From this page, administrators can also test their model.
- **Chapter 4: AI Agents** describes how to create and configure AI agents.
- **Chapter 5: Dashboards** provides contact center performance data across all agents.

In discussing how this application works, this manual assumes that you have the following:

- A basic understanding of fundamental contact center terms such as inbound, outbound, user, IVR, and contact.
- Basic navigation skills for standard web-based Windows graphical user interfaces (GUIs). This includes the ability to right-click and left-click, select options from a right-click menu, resize and minimize windows, and navigate and scroll with a mouse pointer.

The following conventions are used in this manual:

- **Notes** highlight important information.
- **Cautions** are used to bring attention to functions and features that can impact the information viewed.
- Words displayed in **bold** font are defined within the paragraph.
- *Italics* are used to indicate buttons found on the software interface.

Where you can activate a feature in more than one way (button, menu, icon, and so forth), the manual describes only one method.

This version of the manual is specific to AI Studio and requires ice server version 15.2.



Chapter 1: Getting Started

ComputerTalk provides a range of AI agents, commonly referred to as bots:

1. Custom Question Answering (FAQ)

The FAQ bot searches verified answers for a match that meets the configured confidence level, then relays that response to the customer.

All question-answer pairs in verified answers are configured in Knowledge Management under the Verified Answers tab.

If you are configuring an FAQ bot, please refer to Chapter 2: Knowledge Management.

2. FAQ GenAI

The FAQ genAI bot uses a hybrid search model. It first checks verified answers for a match that meets the configured confidence level. If there is no confident match, the bot falls back on genAI to provide an answer.

If you are configuring a FAQ genAI bot, please refer to Chapter 2: Knowledge Management.

3. Natural Language Routing

The NLR bot is a custom workflow that utilizes AI to understand conversational language. It enables users to build custom natural language understanding models to predict the overall intention of an incoming utterance and extract important information from it.

If you are configuring an NLR bot, please refer to Chapter 3: Natural Language Routing.

AI Studio provides a unified platform that allows administrators to create, manage, configure, test, and report on their bots across multiple performance metrics.

To fully utilize AI Studio and the various configurations available, you must have the following:

- iceBot configured in workflow
- An ice user account
- ice server version 15

Note: If you do not have an ice account created, contact your ice administrator.

This chapter includes information about login procedures.

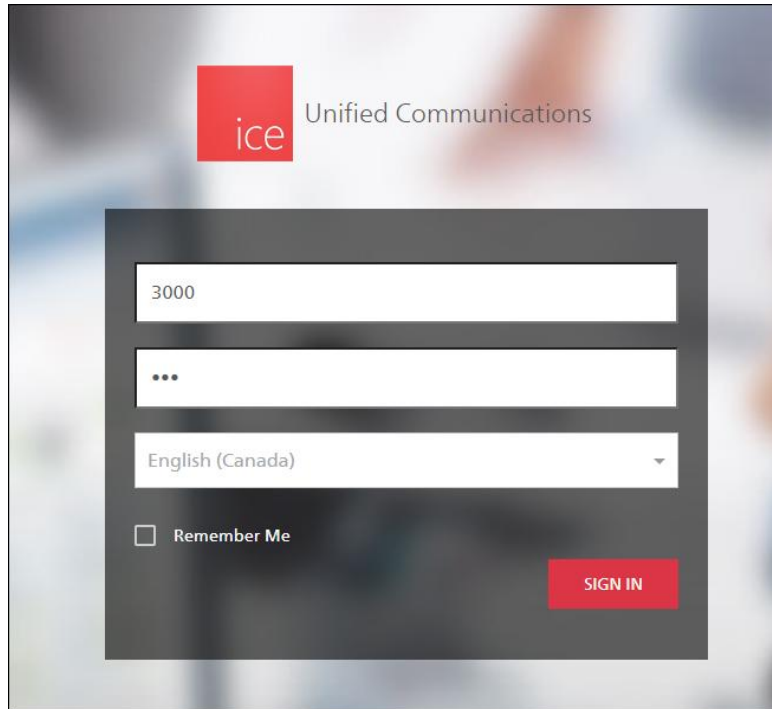
Once familiar with the AI Studio interface, you may proceed to subsequent chapters for detailed information on how to manage your knowledge base.

To access AI Studio, you must first log onto iceManager.

iceManager is a web-based application and can be used on any computer running a web browser. To sign in, you must provide a user ID and password, or use your single sign-on credentials. Contact your ice administrator if you do not have this information.

To sign in to iceManager:

1. Open your web browser and go to your iceManager site.

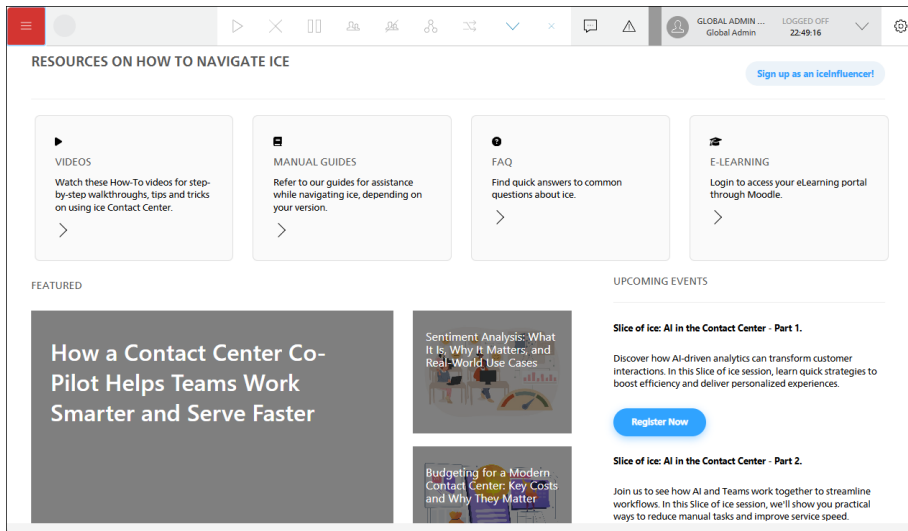


2. In the 'Username' field, enter your four-digit user ID.
3. In the 'Password' field, enter your password.
4. If you wish to view iceManager in a language other than English, click the dropdown and select the language of choice. AI Studio will be displayed in the selected iceManager language.
5. Select the 'Remember Me' check box if you wish to have your username pre-populated the next time you access the Sign In page.

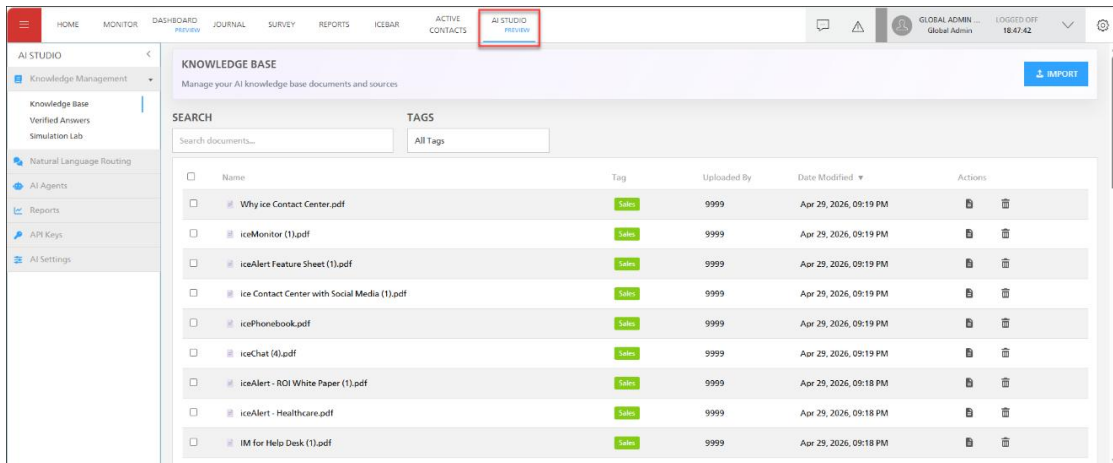
Note: This option is not recommended for shared computers.

6. Click Sign In.
7. Once you have signed in, you will see the home page.

Below is a screenshot of a sample home screen.



- Click on the *AI Studio* button in the Navigation Panel.



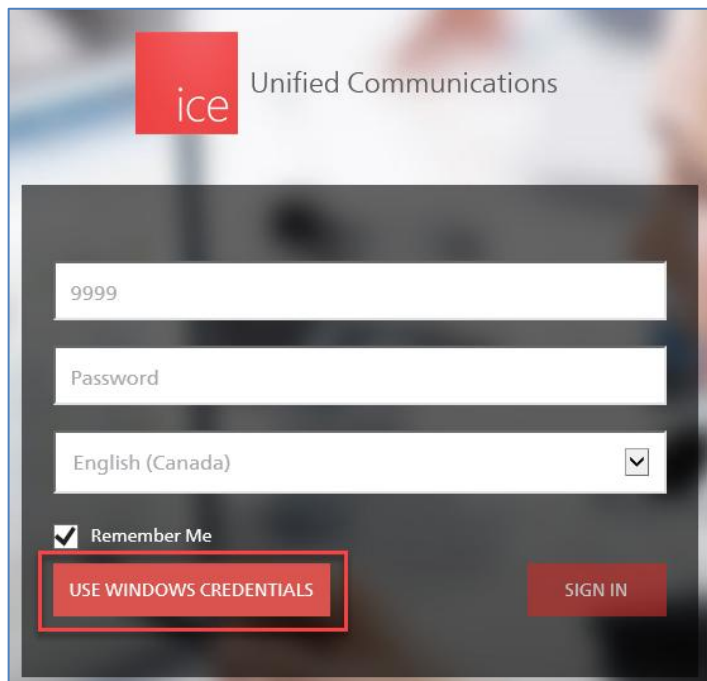
Single Sign-On

If your organization has enabled Single Sign-On for iceManager, you will be able to sign on using your Windows credentials.

Note: To enable Single Sign-On, it will need to be configured using Active Directory in iceAdministrator. For further information on how to enable Single Sign-On, please review the *iceAdministrator User Manual*.

Signing On with Single Sign-On

Once Single Sign-On is properly configured, launch the iceManager website, and click the Use Windows Credentials button rather than entering your username and password.

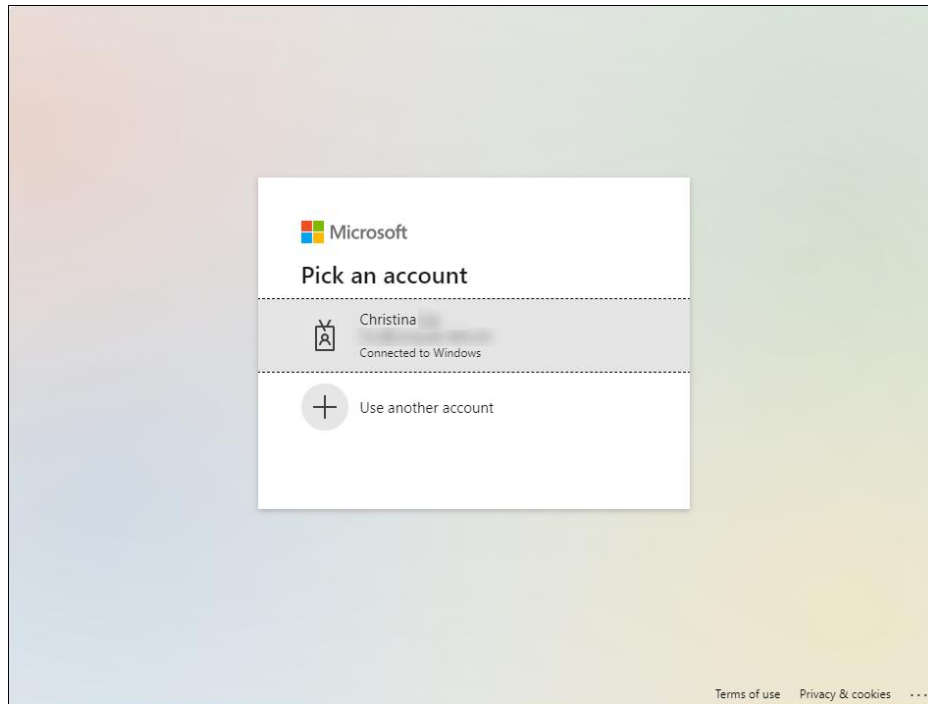
The image shows a web browser window displaying the iceManager login page. At the top left is the 'ice' logo (a red square with the word 'ice' in white) and the text 'Unified Communications'. Below this is a dark grey login form. The form contains three input fields: the first has '9999' entered, the second is labeled 'Password', and the third is a dropdown menu showing 'English (Canada)'. Below these fields is a checkbox labeled 'Remember Me' which is checked. At the bottom of the form are two buttons: 'USE WINDOWS CREDENTIALS' (highlighted with a red rectangular box) and 'SIGN IN'.

1. When you click OK, the Single Sign-On dialog box will appear.

Note: If you wish to skip this step for future logins, check the box for *Remember Me*. This way, you will not have to enter your user ID each time you sign in.

2. To sign in with this method, you will use the same credentials that you use to log on to your computer or your company email. Enter your username and password in the boxes.

Note: This dialog box may look different, depending on the way your administrator has configured the system.



AI Settings

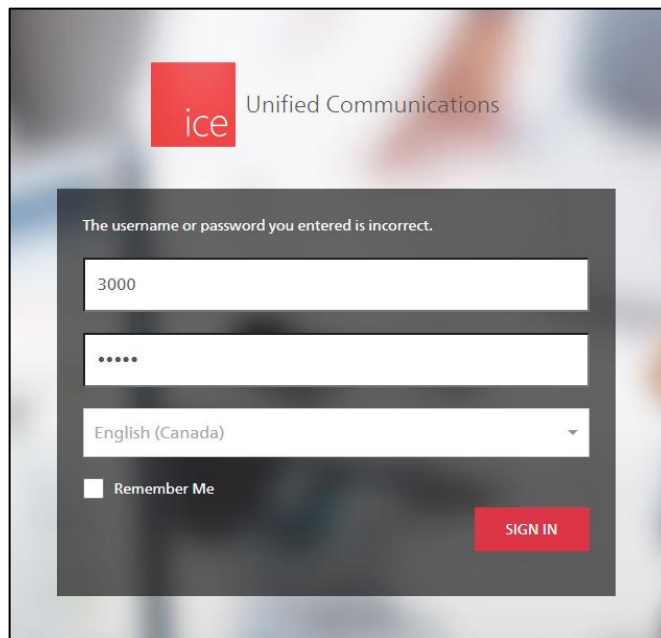
The screenshot displays the 'AI SETTINGS' page in the AI Studio interface. The left sidebar contains a navigation menu with the following items: AI STUDIO, Knowledge Management (with sub-items: Knowledge Base, Verified Answers, Simulation Lab), Natural Language Routing, AI Agents, Reports, API Keys, and AI Settings. The main content area is titled 'AI SETTINGS' and includes the subtitle 'Configure tenant-wide defaults for AI features.' Below this, there is a section for 'TEXT TO SPEECH' with a 'Default Voice Name' input field. The input field has a placeholder text 'e.g. en-US-AriaNeural'. A blue 'SAVE' button is located below the input field. A note below the input field states: 'The Azure TTS voice name used by all AI bots unless overridden on a per-bot basis.' The top of the interface features a header bar with a settings gear icon, a user profile icon, the text 'DEFAULT NAME ... Administrator', 'LOGGED OFF 18:13:31', and a dropdown arrow.

On this page, administrators can configure the default voice model used by all AI bots unless overridden on a per-bot basis.

Common Error and Warning Messages

Authentication error

If a user types the wrong User ID or the wrong password, the following message appears.

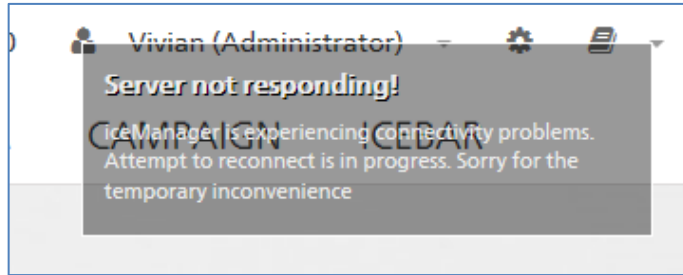


If you cannot remember your password or User ID, an ice administrator can reset it in iceAdministrator. For more information refer to the *iceAdministrator User Manual*.

Loss of connection

A few seconds after connection is lost, an error message appears in the top right corner of the screen:

"Server not responding! iceManager is experiencing connectivity problems. Attempt to reconnect is in progress. Sorry for the temporary inconvenience".



The message fades away after a few minutes. iceManager will keep attempting to reconnect until it is successful.

Verify that you are connected to the Internet. If you are connected, but still receive the "Server not responding" message, contact your ice administrator.

API Keys

This section allows administrators to create API Keys which are used to authenticate with the AI Studio web service. They are created by users with administrator permissions and have access to all API methods.

API KEYS

API KEYS ↺

ID	Name	Created By	Created Date	Enabled
api...	iceWorkflow	Matthew (9999)	2026-04-06 15:24:43	✓

Check [API Documentation](#) for more information.

CREATE API KEY ⓘ

Disable and Enable the API Key

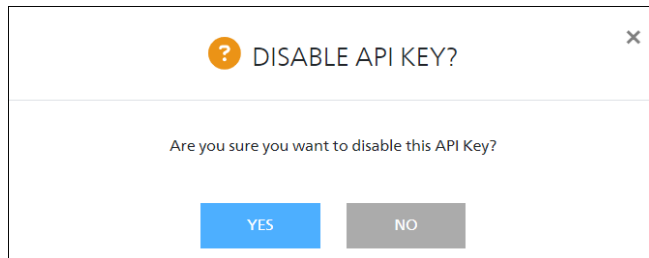
To disable the API Key, right-click on the key in the grid and select “Disable”.

API KEYS ↺

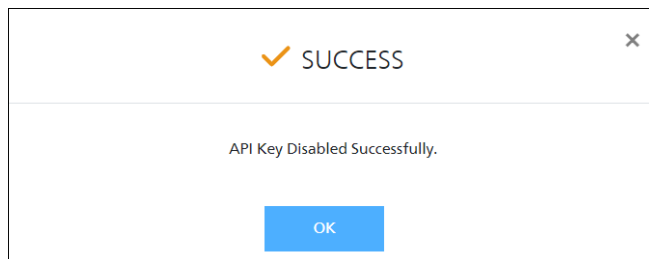
ID	Name	Created By	Created Date
1	Andrei	Global Admin (9999)	2024-05-10 14:32:38

Disable

Select “Yes” in the following window to disable the API Key.



The following window will display upon successfully disabling the key.



To enable, right click on the key and select *Enable*.

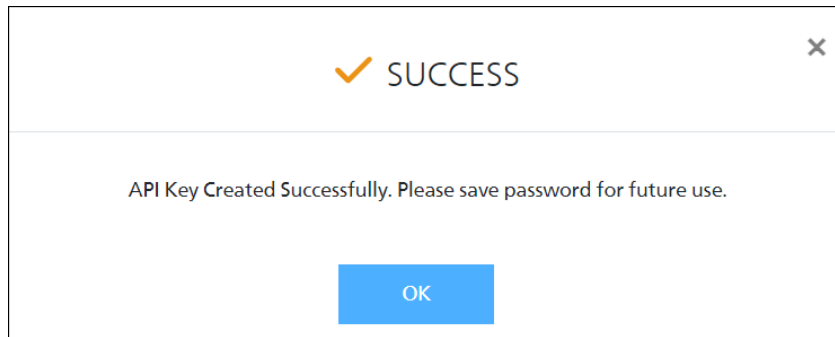
Create API Key

Administrators can create API Keys.



To create a key, enter a name for the key in the Key Name field and click the *Submit* button.

Upon successful creation, the following window will display.



A key will be created and can be copied using the copy icon next to the key field.

A form titled "CREATE API KEY" with an information icon. It contains two input fields: "Key Name" with the value "June" and "Key" with a masked value of dots. To the right of the "Key" field is a copy icon. At the bottom are two buttons: "CLEAR" (red) and "SUBMIT" (blue). There is an expand/collapse arrow in the top right corner.

Note: The Key is only available once on submit. As it cannot be displayed again, it is important to save the Key securely if it will be required at a later time.



Chapter 2: Knowledge Management

Knowledge management consists of the knowledge base, verified answers, and the simulation lab. These three sections allow you to configure the resources needed for Knowledge Bots, also known as FAQ bots.

The knowledge base contains the sources that are used to generate AI answers. Sources are manually assigned tags, which allows bots to only reference a specific subset of sources, while also allowing you to manage all knowledge sources from one central knowledge base. You can modify your sources to enhance answers and improve the accuracy of the genAI bot.

Users with FAQ bots will create and manage all verified question-answer pairs using the Verified Answers page. You can also configure answer groups, create alternate questions and export your set of verified answers.

This section closes with a Simulation Lab that allows you to test your knowledge base and verified answers to ensure that you have appropriate coverage for the questions posed to your bot. You can test your bot's performance and compare your verified answers against AI generated responses.

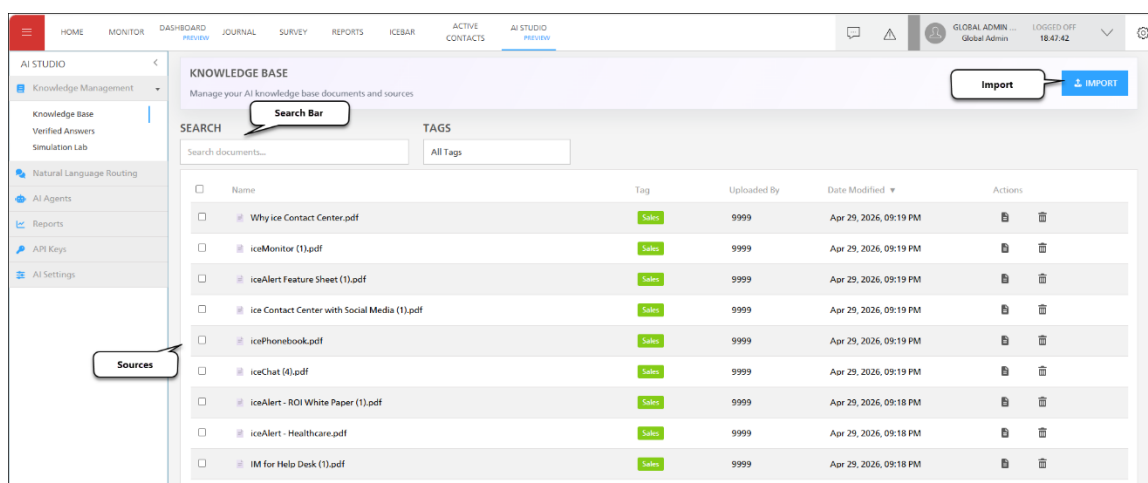
Knowledge Base

The knowledge base contains the sources that are used to generate AI answers. You can modify your sources to enhance answers and improve the accuracy of the genAI bot.

If your bot uses a hybrid search where it first checks for a match in the verified answers and falls back on genAI to provide an answer if there is no confident match, you will use the knowledge base to manage those sources.

Administrators can upload documents (PDF, DOCX, TXT, CSV or XLSX) or import from a website through the Knowledge Base.

The Knowledge Base consists of the sources, the search bar and the import button.



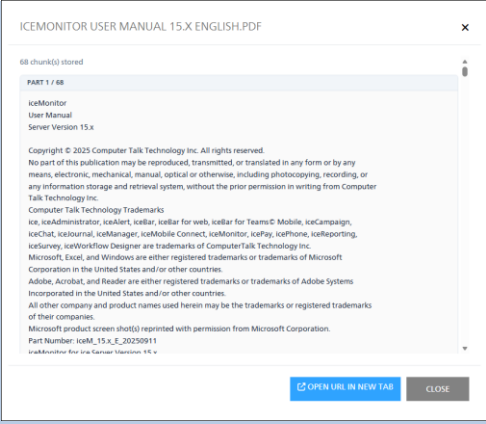
Sources

A list of all sources that have been uploaded to the project is displayed in the sources grid.

The table describes the fields in the sources grid.

Column	Description
Name	The name of the source.
Tag	The tag assigned to the file. Tags are used to organize and filter the documents and filter the sources that a bot references.

Column	Description												
	A single implementation may have multiple different bots sharing the same knowledge base. In these cases, administrators can specify the tags for each specific bot, and limit the sources they reference. To change a tag, hover over the tag. An edit icon will appear. <table><tr><td>☐</td><td> IceSurvey (1).pdf</td><td> Sales  </td><td>9999</td></tr><tr><td>☐</td><td> Why ice Contact Center.pdf</td><td> Sales </td><td>9999</td></tr><tr><td>☐</td><td> IceMonitor (1).pdf</td><td> Sales </td><td>9999</td></tr></table> Click on the icon. The following field appears: Tag Sales ✓ × Sales Enter in your new tag and click the checkmark icon to save. Click the “x” to cancel.	☐	IceSurvey (1).pdf	Sales 	9999	☐	Why ice Contact Center.pdf	Sales	9999	☐	IceMonitor (1).pdf	Sales	9999
	☐	IceSurvey (1).pdf	Sales 	9999									
	☐	Why ice Contact Center.pdf	Sales	9999									
	☐	IceMonitor (1).pdf	Sales	9999									
	Uploaded By	Displays the user ID of the ice user who uploaded the file.											
Date Modified	The date that the file was last modified.												
Actions	For webpages, this column contains a refresh button. This will refresh the content from the webpages. For other documents, this column contains a view and a delete button. The view button () allows administrators to view the stored content.												

Column	Description
	 The delete button () allows administrators to delete the source.

Import

Administrators can upload documents or crawl websites through the Import button. The supported file types include: PDF, DOCX, TXT, CSV or XLSX.

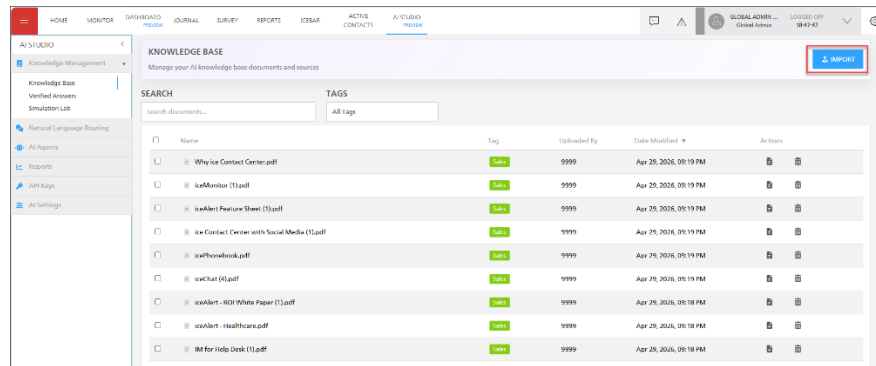
Upload Documents

Each file is parsed to extract clean text, split into token-bounded chunks with overlap, embedded and stored in the database.

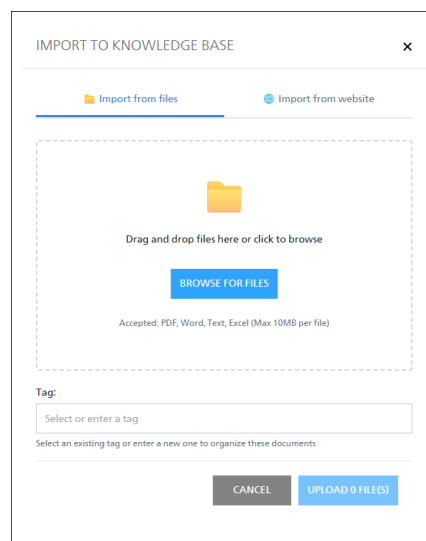
Note: The maximum size per file is 10 MB.

Follow the steps below to import a new file.

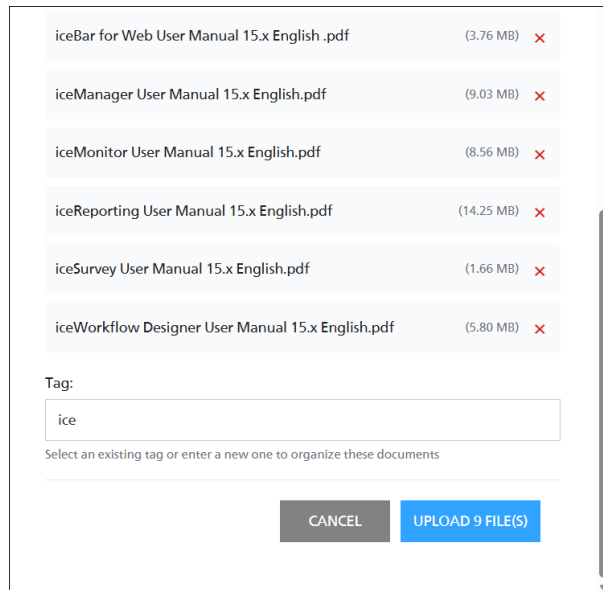
1. Click the *Import* button.



2. Select the *Import from files* tab at the top of the window.



3. Select one or more files to import through the drag and drop option, or through the *Browse for Files* button.
4. Enter a tag in the tag field. This tag will be used to filter your search results. If your bot is configured with hybrid search, the bot will first search through your verified question-answer pairs. If there is no high confidence match, ice will then search the knowledge base for any related information to generate an AI answer.



5. Click Upload Files. You will see the confirmation message.

Note: If you upload a duplicate file, the new upload will replace the original file.

Crawl a Website

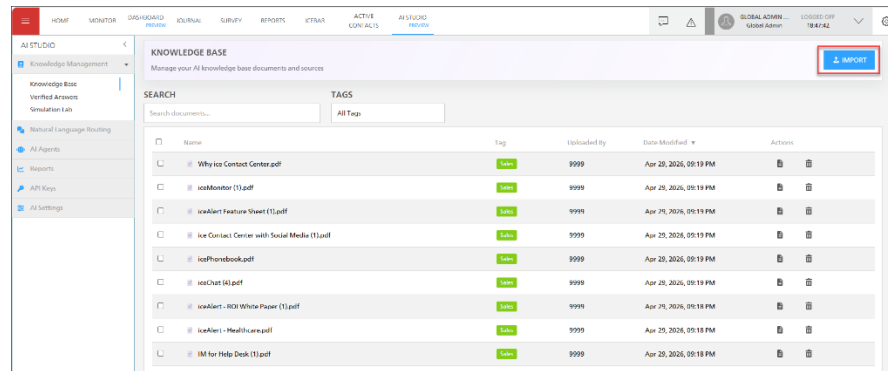
Administrators can also import information from a website. The web crawler is a built-in background service that fetches, cleans, and segments web pages directly into the knowledge base. Administrators can choose between basic and advanced mode.

Basic mode allows administrators to input a website. It starts at the base URL, extracts all links on each page and follows links within the same domain up to the configured depth and page limit.

Note: Not all websites are able to be crawled. There may be limitations due to website configurations or security permissions.

To use basic mode, follow the steps below:

1. Click the Import button.



2. Select the *Import from website* tab at the top of the window.

The 'IMPORT TO KNOWLEDGE BASE' dialog box is shown. It has two main tabs: 'Import from files' and 'Import from website'. The 'Import from website' tab is selected. Under this tab, there are two sub-tabs: 'Basic Mode' (selected) and 'Advanced Mode (Sitemap)'. The 'Basic Mode' section includes a 'Website URL' field with the value 'https://example.com', a 'Maximum Pages (optional)' field with the value '5000', and a 'Crawl Depth (optional)' field with the value '1'. A note states: 'Note: Website crawling will extract text content from web pages'. There is also a 'Tag' field with the placeholder 'Select or enter a tag'. At the bottom, there are 'CANCEL' and 'START CRAWL' buttons.

3. Select basic mode and enter the website you wish to crawl.
4. Enter a tag in the tag field. This tag will be used to filter your search results and organize your documents.
5. Click Start crawl. You will see the confirmation message when this crawl is complete.

Advanced mode allows administrators to input a sitemap URL which will fetch the sitemap and load the URL list into a grid as shown in the image below.

IMPORT TO KNOWLEDGE BASE

Import from files

Import from website

Basic Mode

Advanced Mode (Sitemap)

SELECT URLS FROM SITEMAP

Advanced Mode (Sitemap): Parse your sitemap.xml and manually select which pages to index. This gives you precise control over what content is added to your knowledge base.

Tag:

Select or enter a tag

Select an existing tag or enter a new one to organize these documents

CANCEL

START CRAWL (0 URLs)

ADVANCED CRAWL: SITEMAP SELECTION

Sitemap URL

https://www.samsung.com/sitemap.xml

FETCH SITEMAP

Enter the URL to your sitemap.xml file

Advanced Filters

Filter by URL Pattern (Regex)

for example /blog/.*

Use regular expressions to filter URLs

Filter by Last Modified Date

From:

yyyy-mm-dd

To:

yyyy-mm-dd

CLEAR FILTERS

SELECT VISIBLE

DESELECT ALL

☐	URL	Last Modified	Priority	Change Frequ...
☐	https://www.samsung.com/ae/sitemap.xml	2018-06-14	N/A	N/A
☐	https://www.samsung.com/ae_ar/sitemap.xml	2018-06-14	N/A	N/A
☐	https://www.samsung.com/africa_en/sitemap.xml	2018-06-14	N/A	N/A
☐	https://www.samsung.com/africa_fr/sitemap.xml	2018-06-14	N/A	N/A

Total URLs: 97

Filtered: 97

Selected: 0

CANCEL

START CRAWL

From this grid, administrators can manually select the URLs to scrape.

To use advanced mode, follow the steps below:

1. Click the Import button.

2. Select the Import from website tab at the top of the window.
3. Select advanced mode and click "Select Urls from sitemap".
4. In the window that opens, enter your sitemap URL and click "Fetch Sitemap".

ADVANCED CRAWL: SITEMAP SELECTION

Sitemap URL

Enter the URL to your sitemap.xml file

Enter a sitemap URL and click "Fetch Sitemap" to begin.

5. The URL list will be displayed in the grid below. You can further filter the list using *Advanced Filters*, or select the URLs you would like to crawl.

ADVANCED CRAWL: SITEMAP SELECTION

Sitemap URL

Enter the URL to your sitemap.xml file

Advanced Filters

Filter by URL Pattern (Regex)

Use regular expressions to filter URLs

Filter by Last Modified Date

From: To:

☐	URL	Last Modified	Priority	Change Frequ...
☐	https://www.samsung.com/ae/sitemap.xml	2018-06-14	N/A	N/A
☐	https://www.samsung.com/ae_ar/sitemap.xml	2018-06-14	N/A	N/A
☐	https://www.samsung.com/africa_en/sitemap.xml	2018-06-14	N/A	N/A
☐	https://www.samsung.com/africa_fr/sitemap.xml	2018-06-14	N/A	N/A

Total URLs: 97 Filtered: 97 Selected: 0

6. Once you have selected your URLs, click *Start Crawl*.
7. Enter a tag in the tag field. This tag will be used to filter your search results and organize your documents.

IMPORT TO KNOWLEDGE BASE

Import from files Import from website

Basic Mode Advanced Mode (Sitemap)

SELECT URLS FROM SITEMAP

Advanced Mode (Sitemap): Parse your sitemap.xml and manually select which pages to index. This gives you precise control over what content is added to your knowledge base.

Tag:
Select or enter a tag

Select an existing tag or enter a new one to organize these documents

CANCEL START CRAWL (0 URLS)

- Click Start crawl. You will see the confirmation message when this crawl is complete.

The following table describes the available fields when crawling a website.

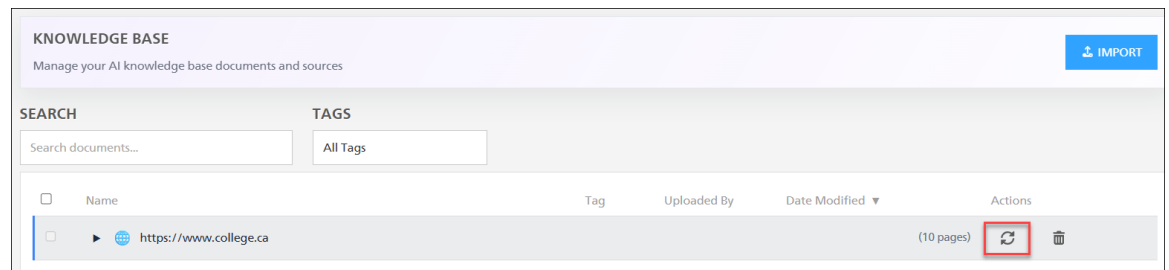
Field	Description
Basic Mode	
Website URL	Enter the URL of the website to crawl. The system will extract and index the content.
Maximum Pages (optional)	Limit the number of pages to crawl. Default: 10, Maximum: 100
Crawl Depth (optional)	How many levels deep to follow links. Default: 1, Maximum: 3
Advanced Mode (Sitemap)	
Select URLs from Sitemap	Enter the URL to your sitemap.xml file and click Fetch Sitemap to begin.

Updating a crawled website

If changes have been made to a website that has already been crawled, you will need to re-crawl the site. The changes will not automatically appear in the knowledge base. However, this is a very simple process, and requires just one click of a button.

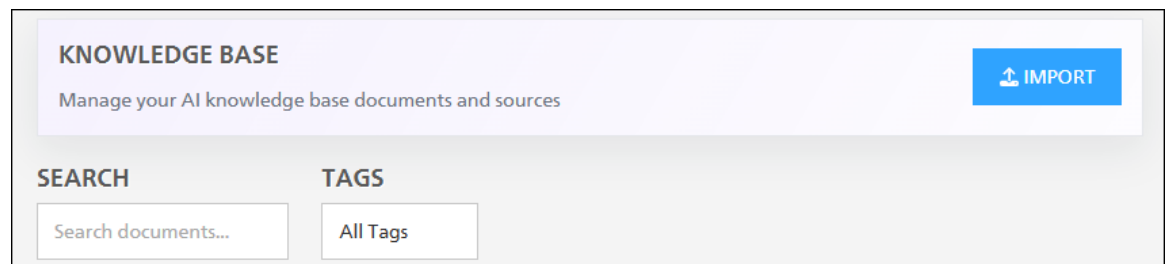
To re-crawl a website, follow the steps below:

1. Locate the website in the knowledge base.
2. In the actions menu, click on the re-crawl website button.



When the process is complete, you will see a confirmation message.

Searching for a source



To search for a source, complete the following steps:

1. If you know the name or part of the name of the source, enter it in the search box. The results will filter accordingly in the sources grid.
2. If you are looking for a source by the tag, click on the tags field to open the tags dropdown list, and select a tag. The sources grid will update to display the sources tagged with the selected tag.

The knowledge base provides the sources for GenAI responses. Administrators can ensure that the knowledge base is up to date with relevant information by adding, deleting and replacing documents.

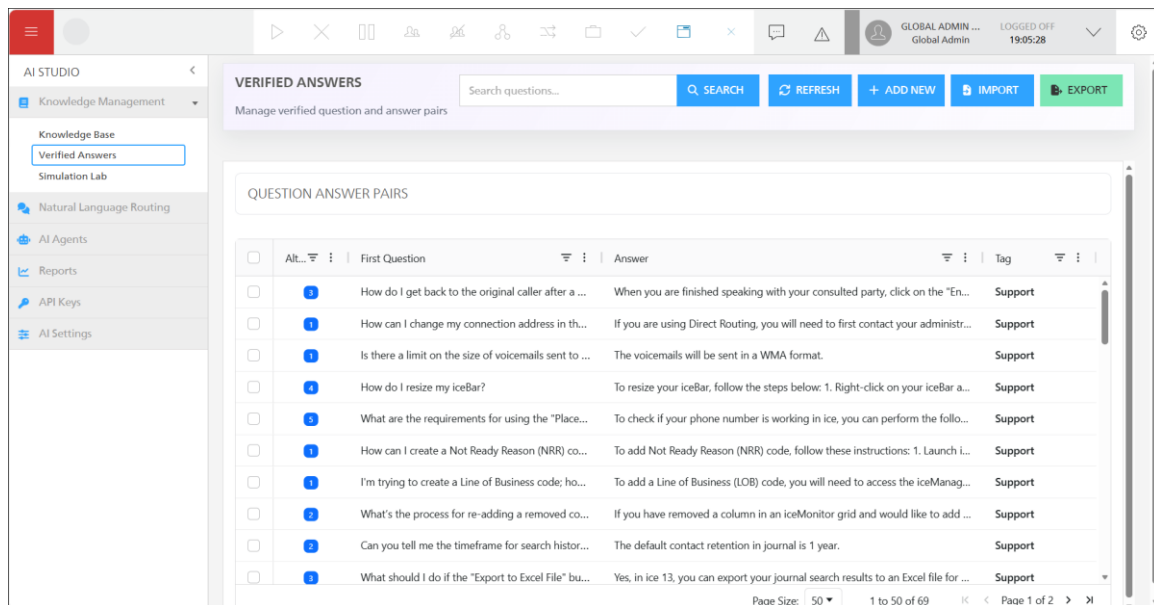
Verified Answers

Verified answers are human-authored Q&A pairs that represent the gold standard of the knowledge base. They are returned verbatim and are not modified by any language model. When a caller's question semantically matches a verified answer above the confidence threshold, the answer is returned exactly as written by the administrator.

The Verified Answer page allows administrators to import, manually author, manage and export question-answer pairs.

The page consists of a search bar, buttons to manage the question-answer pairs, and the question-answer pairs grid.

The Question-Answer Pairs grid displays a list of all the verified question-answer pairs in the project.



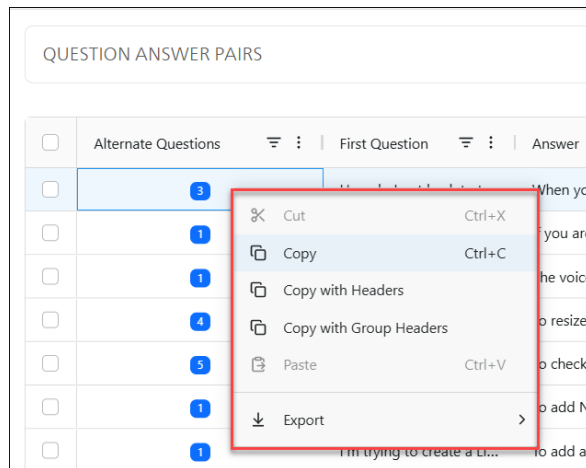
The table below describes each column.

Column	Description
Alternate Questions	You may have questions that are semantically similar to existing questions in the project. In this case, administrators can create alternate

	questions that share one answer via Answer Groups. The <i>Questions</i> column indicates the total number of questions associated with this Answer Group.
First Question	The first question created in the Answer group.
Answer	The answer associated with the question and any configured alternate questions.
Tag	The tag assigned to the question-answer pair. Tags are used to filter the verified question-answer pairs that are referenced by the bot. A single implementation may have multiple different bots sharing the same verified answers page, and knowledge base. In these cases, administrators can specify the tags for each specific bot, and limit the sources they reference.

Right-Click Menu

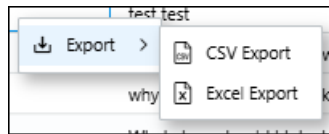
Right-click on a row in the grid to view the right-click menu options.



The *Copy* option allows you to copy the value of the selected field.

- If you wish to copy the field with headers, click *Copy with Headers*.
- If you wish to copy the field with the group headers, click *Copy with Group Headers*.

The *Export* option allows you to export all rows in the grid.

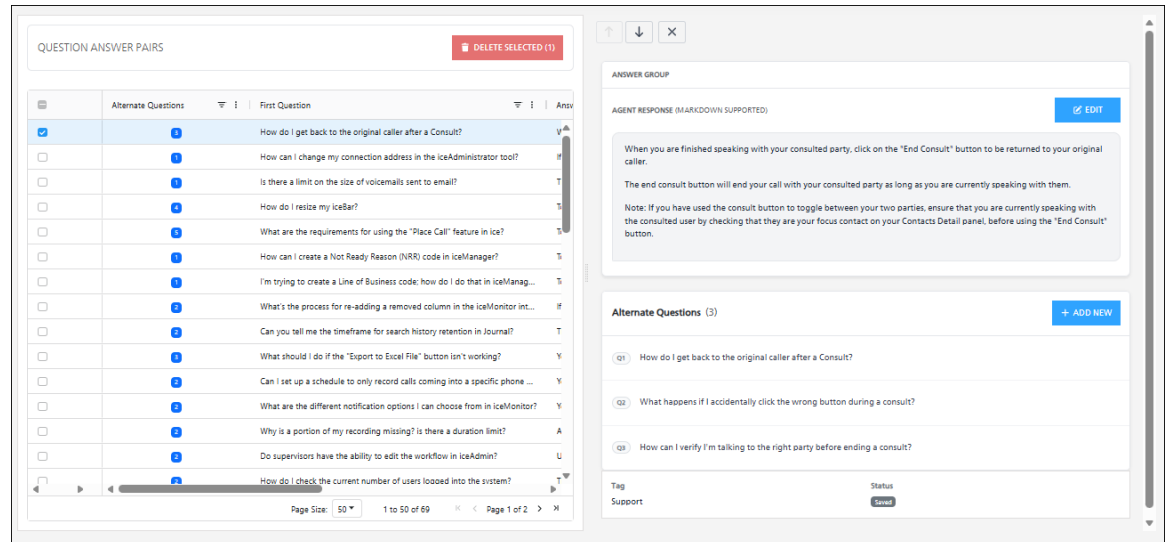


Select from CSV and Excel file formats.

Detail View

To see more information for a specific question-answer pair:

1. Click on the row in the question-answer pairs grid.
2. The Details view for that question-answer pair will display.



The Details View consists of the following sections:

- Answer Group: Agent Response
 - The response associated with the selected question. This response will be returned for all questions in this answer group.
- Alternate Questions
 - All questions associated with this answer group will be displayed. This includes the first question and any alternative questions added. These questions will return the same response.
- Navigation Buttons

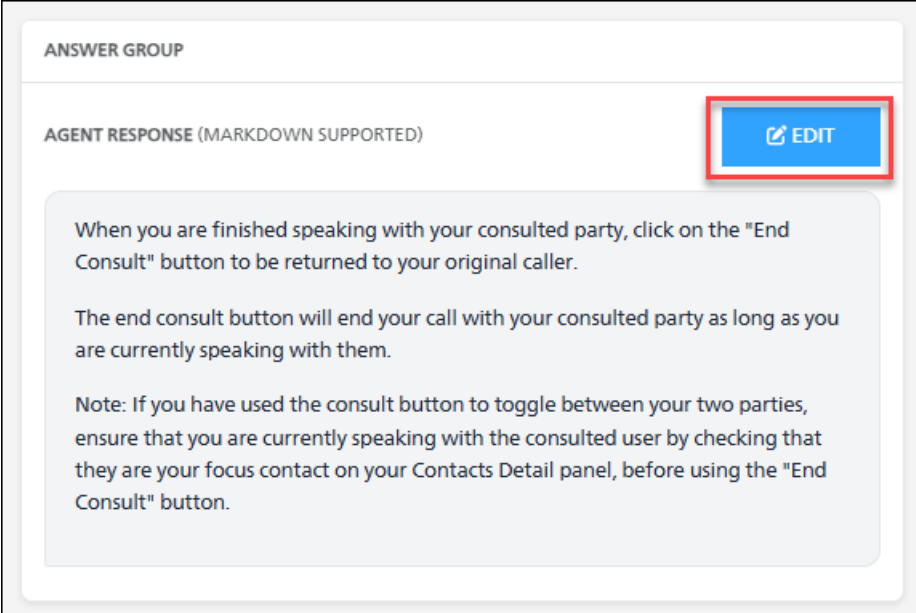
- The navigation buttons above the Answer Group section allow you to view the Details View for the next or previous question-answer pair or close the Details View.

The Details View opens on the right, allowing you to view the complete list of questions in the project on the left. This allows for easy comparison between the general and detailed views. You can view the Details view for other questions by clicking on the rows in the grid on the left.

Edit an answer group response

All questions within an answer group will receive the answer configured in the Agent Response field.

1. If you wish to edit the answer, click the Edit button.



The screenshot shows a user interface for editing an answer group. At the top, there is a header 'ANSWER GROUP'. Below it, the label 'AGENT RESPONSE (MARKDOWN SUPPORTED)' is displayed. To the right of this label is a blue button with a pencil icon and the text 'EDIT', which is highlighted with a red rectangular box. Below the label, there is a text area containing the following content:

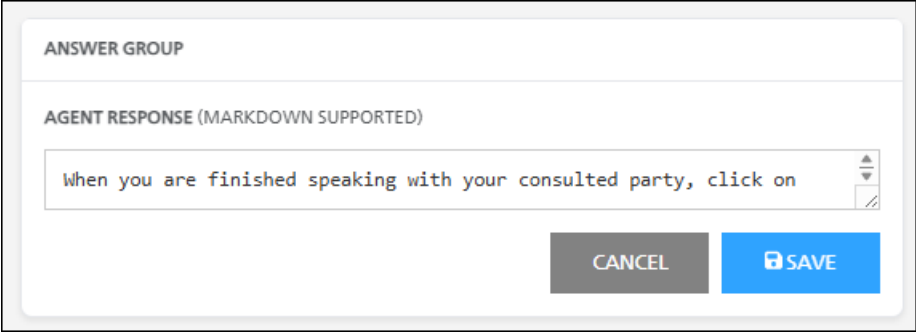
When you are finished speaking with your consulted party, click on the "End Consult" button to be returned to your original caller.

The end consult button will end your call with your consulted party as long as you are currently speaking with them.

Note: If you have used the consult button to toggle between your two parties, ensure that you are currently speaking with the consulted user by checking that they are your focus contact on your Contacts Detail panel, before using the "End Consult" button.

2. You may make your changes in the field and click Save.

Note: This field supports markdown.



ANSWER GROUP

AGENT RESPONSE (MARKDOWN SUPPORTED)

When you are finished speaking with your consulted party, click on

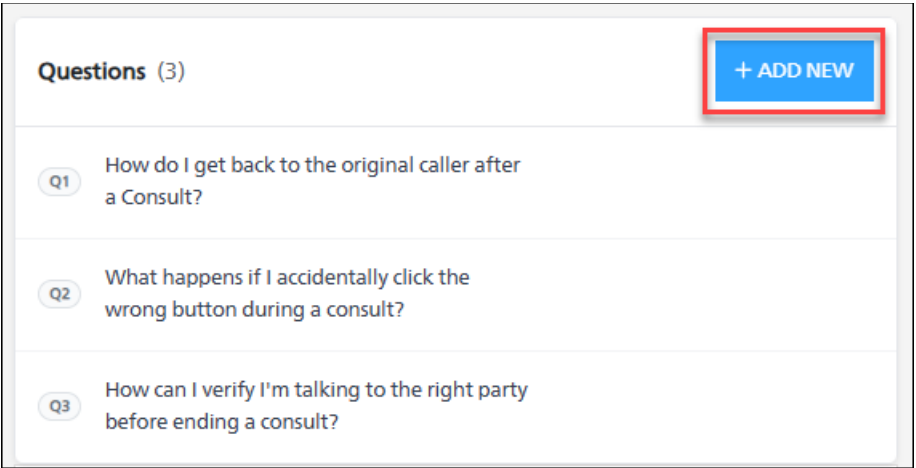
CANCEL SAVE

Add an alternate question

The questions section allows administrators to add alternate questions. These are additional queries that are semantically similar to questions in the project.

To add a new question, follow the steps below.

1. Click the Add New button.



Questions (3)

+ ADD NEW

Q1 How do I get back to the original caller after a Consult?

Q2 What happens if I accidentally click the wrong button during a consult?

Q3 How can I verify I'm talking to the right party before ending a consult?

2. Enter the new question in the Question field. The answer and tag are prepopulated based on the data provided for the first question created in the answer group.

3. Click Save.
4. All alternate questions will be listed in this section.

Note: If you modify the answer or tag, the question-answer pair will be saved as a new and separate question-answer pair in the grid. It will no longer be associated with the original answer group.

Edit an alternate question

To edit an alternate question, follow the steps below.

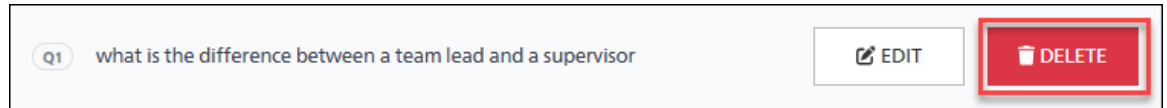
1. Hover over the question you wish to edit. The edit button will appear.

2. Click Edit. The text field is now editable.

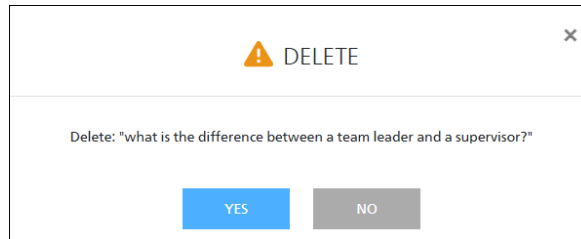
3. Make your edits and click save.

To delete an alternate question, follow the steps below:

1. Hover over the question you wish to delete. The delete button will appear.



2. Click delete.



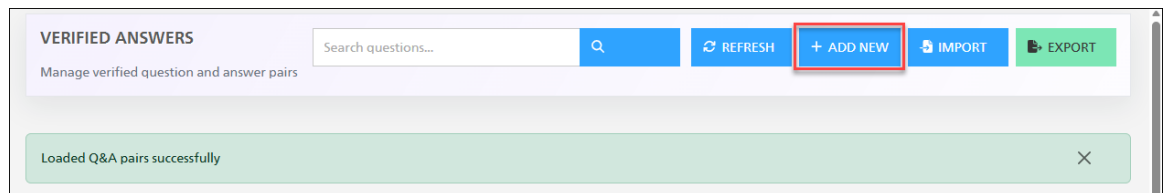
3. Your question will be deleted.

Note: Deleted questions cannot be recovered.

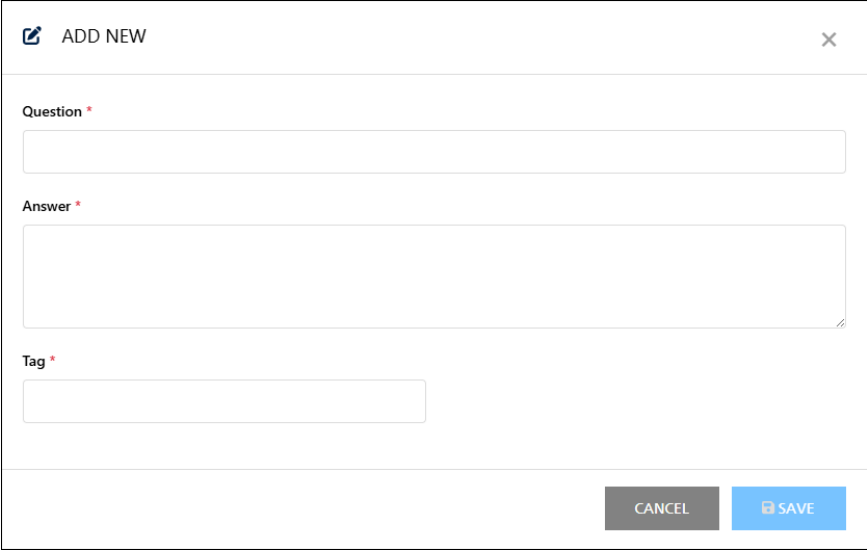
Add New Question

There may be times when you need to manually add one or more questions. You may do this from the *Add New* button.

1. To add a new question to the verified answers, click *Add New*.



2. The following window opens.



ADD NEW

Question *

Answer *

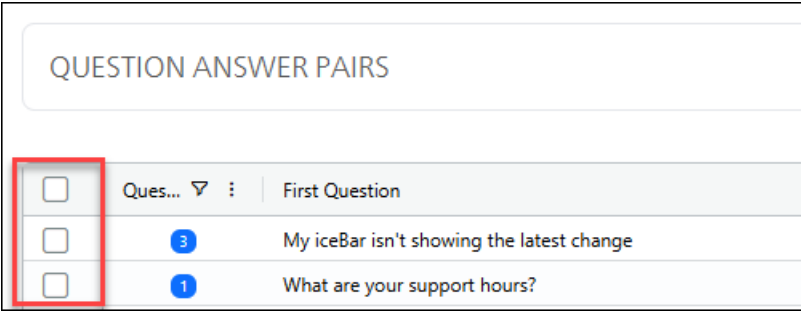
Tag *

CANCEL SAVE

3. Enter your question in the *Question* field.
4. Enter your answer in the *Answer* field.
Note: If you click the *Add New* button while the Detail View is open, the answer field will automatically populate with the Answer Group response.
5. Enter your *Tag*.
6. Click *Save* to save your question-answer pair. This will be saved in your Question-Answer Pairs grid.

Delete a question-answer pair

1. To delete a question-answer pair, select the checkbox next to the question in the grid.

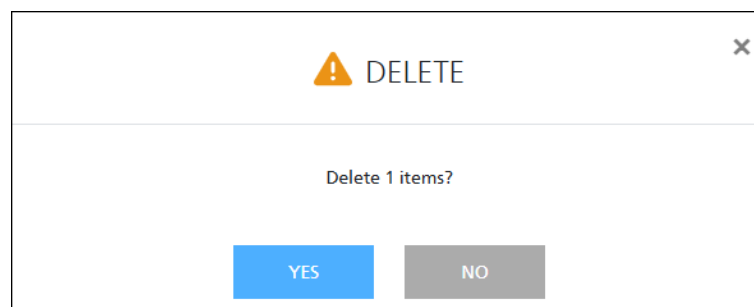


QUESTION ANSWER PAIRS		
☐	Ques... ▾	First Question
☐	3	My iceBar isn't showing the latest change
☐	1	What are your support hours?

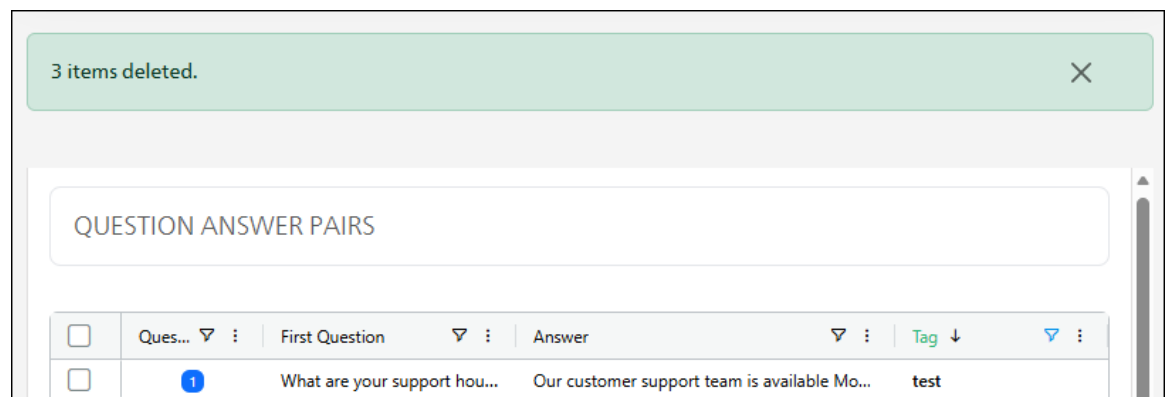
2. The delete button will appear in the top right corner of the screen.

QUESTION ANSWER PAIRS					DELETE SELECTED (1)
	Alter...	First Question	Answer	Tag	
☒	3	How do I get back to the original caller after a Consult?	When you are finished speaking with your consulted party, click on the "E..."	Support	
☐	1	How can I change my connection address in the iceAdministrator tool?	If you are using Direct Routing, you will need to first contact your adminis...	Support	
☐	1	Is there a limit on the size of voicemails sent to email?	The voicemails will be sent in a WMA format.	Support	

3. Select all questions that you wish to delete.
4. Click Delete selected.



5. Your question-answer pair will be deleted.



Searching for a question

You may wish to see a list of questions containing a certain term, or validate if a question already exists in your verified answers.

To search for a question, complete the following steps:

1. If you know the term or contents of the question, enter it in the search box.
2. Click the search button. The results will filter accordingly in the sources grid.

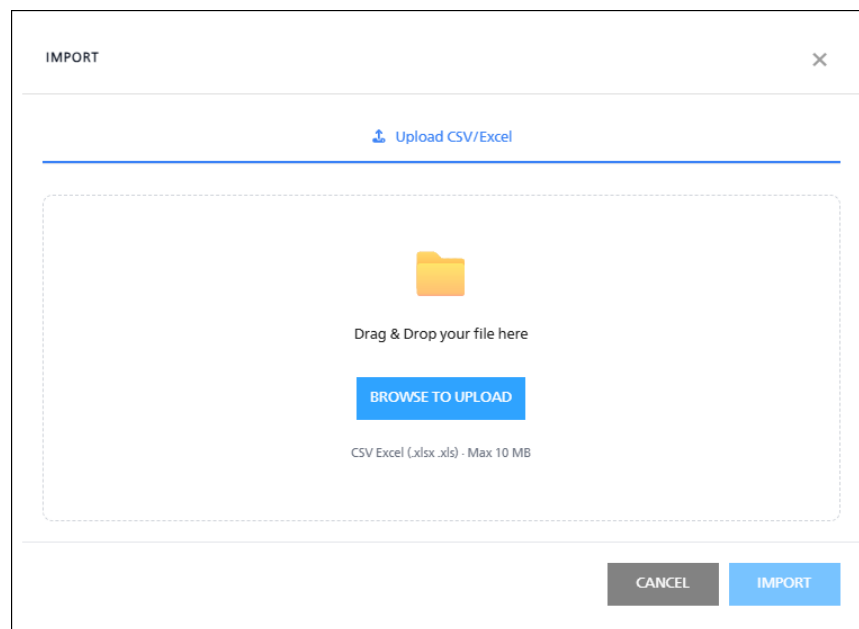
Import

You may wish to import a configured list of question-answer pairs from a CSV file. The mandatory columns are: Question, Answer, Tag. Two additional optional columns can also be included: Category and Subcategory.

Note: The maximum size per file is 10 MB.

Follow the steps below to import a CSV file.

1. Click the Import button.
2. The following window will open.



3. You may drag and drop one or multiple files from your workstation into the field, or click the Browse File button to select a file.

qna-export-2026-06-10.csv
41.4 KB

DEFAULT TAG ⓘ
Select or type a tag

☒ First row is header DELIMITER ' , '

PREVIEW & COLUMN MAPPING Question Answer Tag / Source

Question	Answer	Tag / Source
QUESTION	ANSWER	TAG
	When you are finished speaking with your consulted party, click on the "End Consult" button to be returned to your original caller. The end consult button will end your call with your consulted party as long as you are currently speaking with them. Note: If you have used the consult button to toggle between your two parties, ensure that you are currently speaking with the consulted user by checking that they are your focus contact on your Contacts Detail panel, before using the "End Consult" button.	Support

- If your file does not follow the recommended format of Question, Answer, Tag, you may use the column mapping drop-down menus to match the column to the appropriate category.
- Enter a default tag. This will be used if any rows in the file do not have a tag.

qna-export-2026-06-10.csv
41.4 KB

DEFAULT TAG ⓘ
Support

☒ First row is header DELIMITER ' , '

- Click Import. You will see a confirmation message when the import is complete.

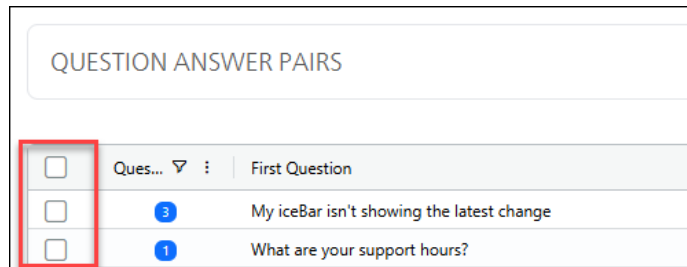
CANCEL IMPORT

Export

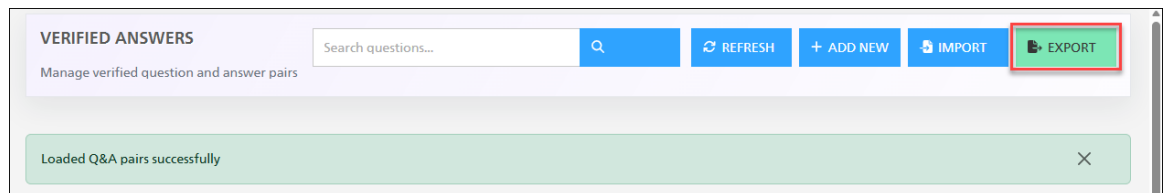
You may need to create a backup of your verified questions. You may export selected questions or all verified answers.

To export your question-answer pairs, follow the steps below.

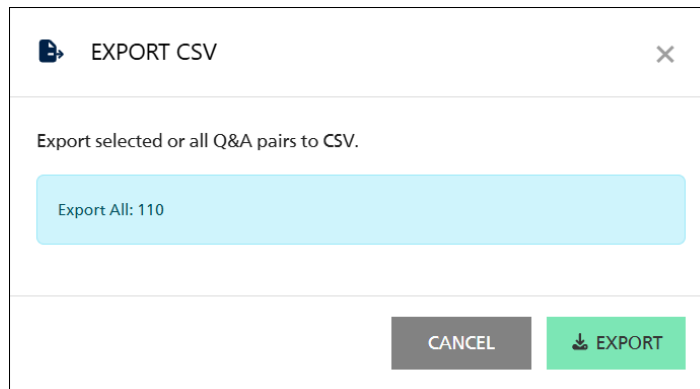
1. Select the question-answer pairs you wish to export by selecting the checkboxes in the grid. If you wish to export all questions, do not select any pairs.



2. Click Export.



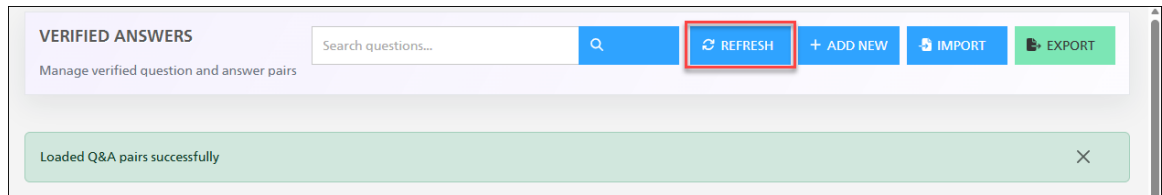
3. The following window will open.



4. The number of selected questions will be displayed at the bottom of the window. When you are ready to export the questions, click Export.
5. Your file will be available in your browser downloads.

Refresh

If you wish to refresh the Question-Answer Pairs grid after making edits, click the Refresh button.

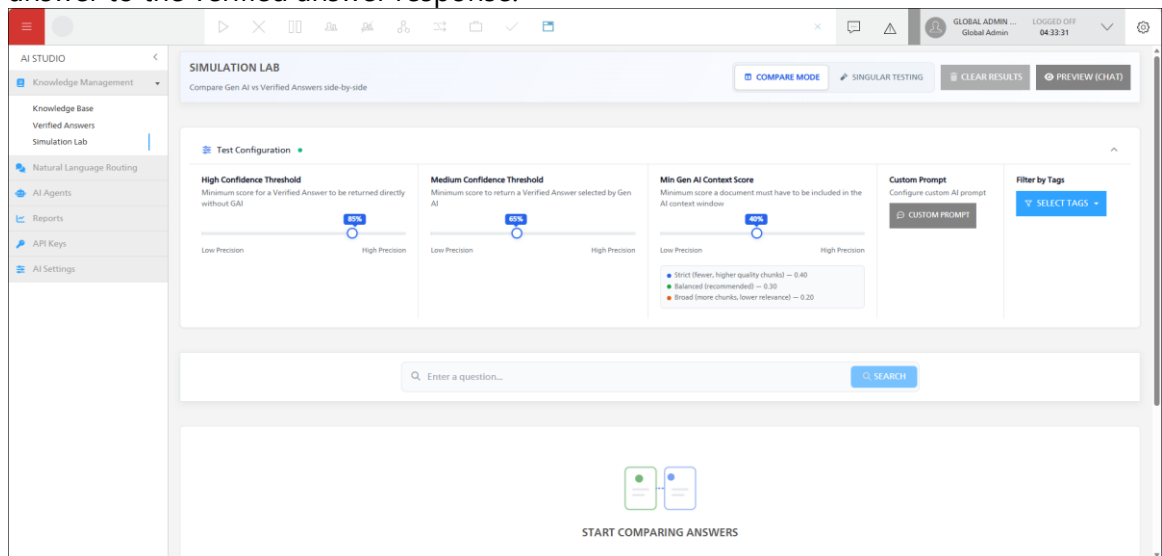


Simulation Lab

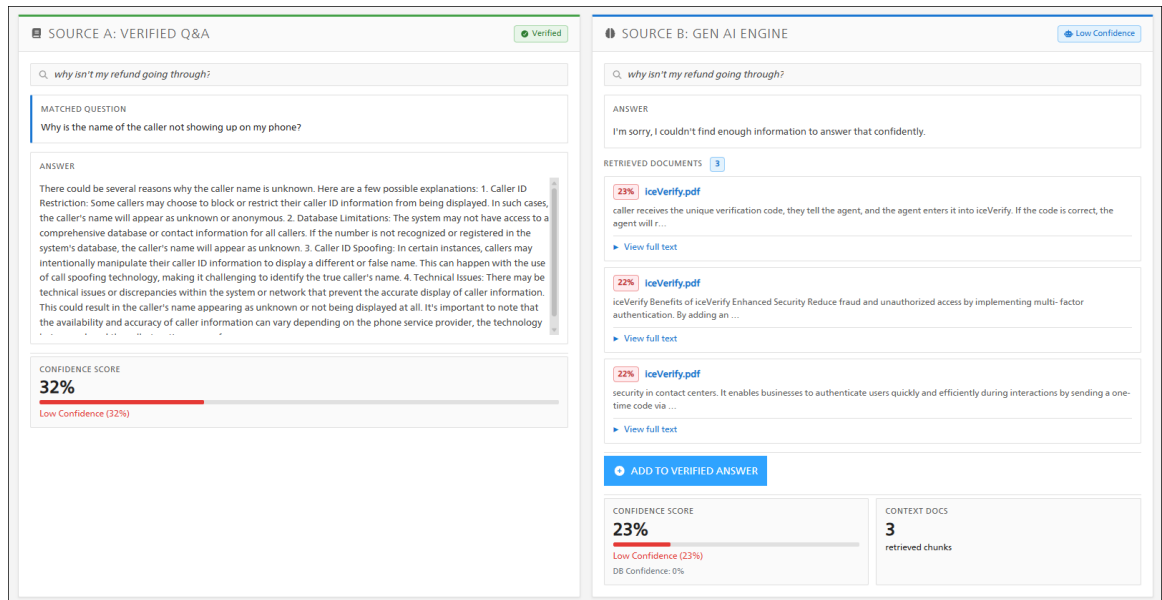
The simulation lab allows administrators to test the knowledge base and verified answers. Administrators can choose from two modes: compare and singular testing.

Compare Mode

In compare mode, administrators can test their queries and compare the genAI answer to the verified answer response.



This provides a platform to compare the human-authored verified answer to the AI-generated response, and evaluate any modifications that need to be made. Additionally, administrators can adjust the threshold values to see the suggested genAI answers at different confidence levels.

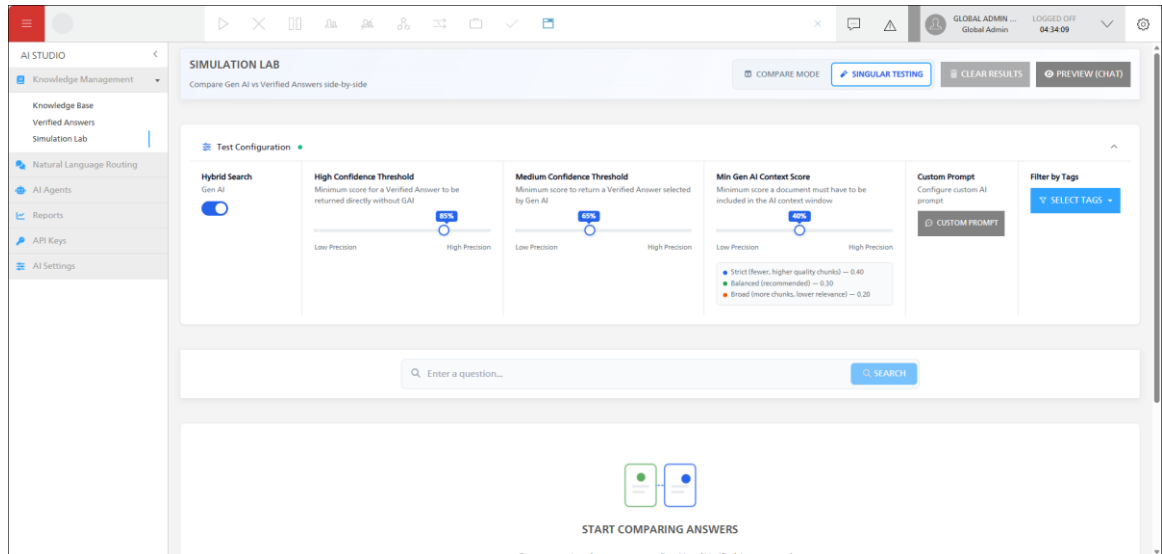


The following table describes the test configuration settings:

Setting	Description
High Confidence Threshold	Set the minimum score for a verified answer to be returned directly without generative AI (GAI).
Medium Confidence Threshold	Set the minimum score to return a verified answer selected by genAI.
Min GenAI Context Score	Set the minimum score a document must have to be included in the AI context window.
Custom Prompt	A custom prompt overrides the default GAI prompt for this Agent. If left blank the default help-desk prompt is used.
Filter by Tags	Select the tags associated with the sources in the knowledge base to be used in the AI response. The bot will only refer to the sources with corresponding tags when generating a response.
Preview	Preview (Chat) – Click to see the response in markdown. Chat (Voice) – Click to see the response in plain text.

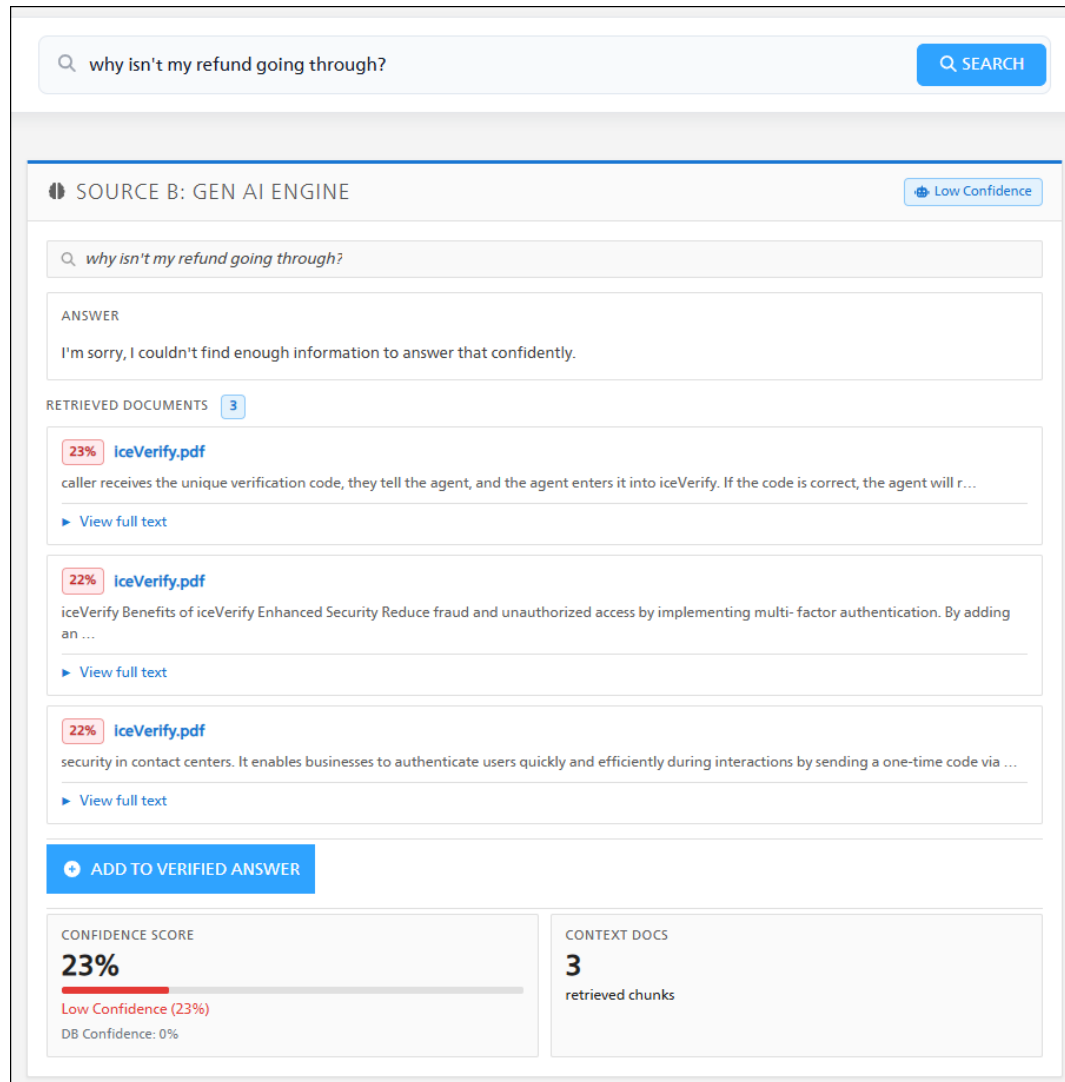
Singular Testing

Singular testing allows administrators to call the hybrid search endpoint to test the knowledge base and verified answers.



When the hybrid search toggle is enabled, it allows genAI responses to be presented if there is no verified answer that meets the high confidence threshold, and there is a genAI answer that meets the medium confidence threshold.

In the following screenshot, the medium confidence threshold is configured at 65%. However, the top 3 responses show 23%, 22% and 22% respectively. As these options do not meet the 65% threshold, and there is no verified answer matching the question, the response received is "I'm sorry, I couldn't find enough information to answer that confidently."



At the bottom of the window, there is an option to add to the verified answers to fill this gap.

Click "Add to verified answer." The following window opens:

Add this Gen AI response as a Verified Answer pair

Review and edit the question and answer before saving to Verified Answer.

Question *

why isn't my refund going through?

Answer *

I'm sorry, I couldn't find enough information to answer that confidently.

Tag

SALES

SUPPORT

BILLING

ICE

MAPLE

MATT

TEST

Or type a custom tag...

CANCEL

✓ SAVE TO VERIFIED ANSWER

Administrators can edit the question, answer and tag associated with the new question-answer pair.

When the hybrid search is disabled, the bot will not return a genAI response. This allows administrators to compare the bot experience when AI is enabled with the experience when it is not.

SOURCE A: VERIFIED Q&A Verified

MATCHED QUESTION

Why is the name of the caller not showing up on my phone?

ANSWER

There could be several reasons why the caller name is unknown. Here are a few possible explanations: 1. Caller ID Restriction: Some callers may choose to block or restrict their caller ID information from being displayed. In such cases, the caller's name will appear as unknown or anonymous. 2. Database Limitations: The system may not have access to a comprehensive database or contact information for all callers. If the number is not recognized or registered in the system's database, the caller's name will appear as unknown. 3. Caller ID Spoofing: In certain instances, callers may intentionally manipulate their caller ID information to display a different or false name. This can happen with the use of call spoofing technology, making it challenging to identify the true caller's name. 4. Technical Issues: There may be technical issues or discrepancies within the system or network that prevent the accurate display of caller information. This could result in the caller's name appearing as unknown or not being displayed at all. It's important to note that the availability and accuracy of caller information can vary depending on the phone service provider, the technology being used, and the caller's actions or preferences.

CONFIDENCE SCORE

32%

Low Confidence (32%)

Users can configure the settings to match the bot design to test the user experience that live customers will have.



Chapter 3: Natural Language Routing

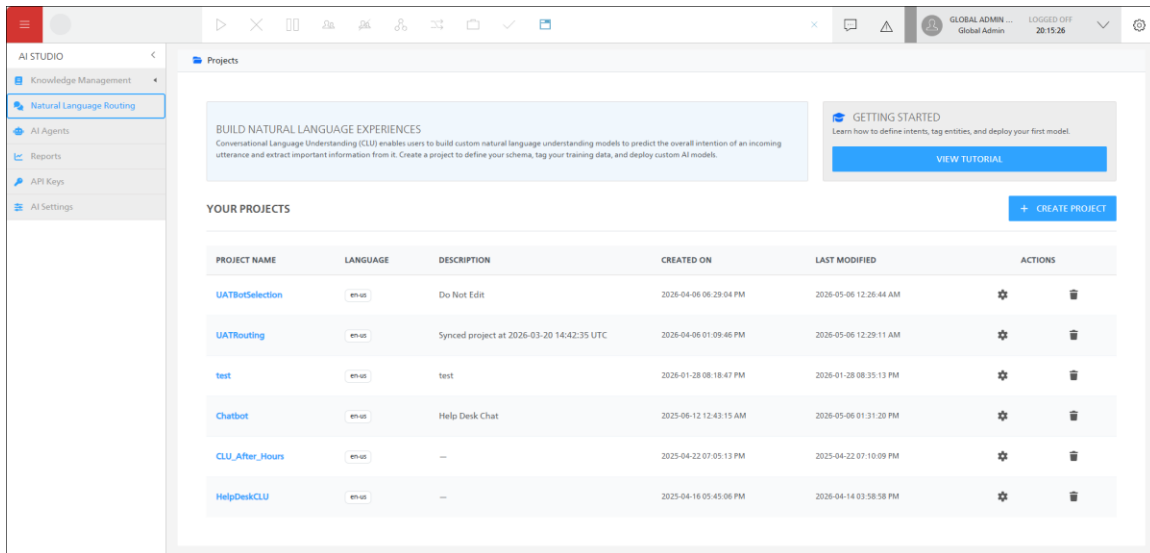
Natural Language Routing bots are custom workflows that utilize AI to understand conversational language. For example, you can enhance or even replace a DTMF menu with a Natural Language Routing (NLR) model to create a more natural flow and to support a greater number of intents per menu.

NLR allows users to build NLR models to predict the overall intention of an incoming utterance and extract important information from it.

The NLR interface within AI Studio provides a cohesive workspace for administrators to define conversational schemas, label training data, and manage the end-to-end lifecycle of Natural Language Processing models.

Project Management

The Natural Language Routing page acts as a centralized dashboard for all of your NLR projects.



The table below explains each project field.

Field	Description
Project Name	The name of the NLR project.
Language	The language(s) supported for the project.
Description	A short description of the project.
Created On	The date the project was created.
Last Modified	The date the project was last modified.
Actions	<u>Settings</u> The settings button allows administrators to configure the description and the confidence threshold. Note: The confidence threshold is the minimum confidence level required from an answer to be returned to a customer. <u>Delete</u> The delete button allows administrators to delete a project.

Project

From the Project Management page, select your project to access the project schema and model lifecycle.

The project page consists of the menu bar on the left, and the Intents, Entities and Utterances tabs on the right.

The menu bar includes the following pages:

- Authoring
- Training Jobs
- Model Performance
- Deployments
- Playground

The screenshot shows the AI Studio Project Management interface for a project named 'HelpDeskCLU'. The left sidebar contains a menu with 'Active Project' (HelpDeskCLU (en-us)), 'DATA LABELING' (Authoring), and 'MODEL LIFECYCLE' (Training jobs, Model performance, Deployments, Playground). The main area has tabs for 'Intents', 'Entities', and 'Utterances'. A green notification bar states 'Project data loaded successfully.' Below this is a search bar 'Search intents...' and an '+ ADD INTENT' button. A table lists intents with columns: NAME, UTTERANCES, DESCRIPTION, and ACTIONS.

NAME	UTTERANCES	DESCRIPTION	ACTIONS
None	12	Used when the user's input doesn't match any known intent or is out of scope.	[Icon]
Sales	20	—	[Icon]
Agent	30	—	[Icon]
Lookup	28	—	[Icon]
New	18	—	[Icon]

Showing 1-5 of 5

PREVIOUS 1 NEXT

Authoring

When AI processes human language, it breaks down the interaction into three main elements: utterances, intents and entities.

Utterances are the spoken or written statements made by a user in natural language when interacting with a bot. These can be simple statements, questions, or commands, or any expression of human language.

For example:

- What is the status of my ticket?
- How can I add a holiday?
- Speak to an agent
- Sales

Intents are the underlying purpose that a user intends to achieve through their utterances. For example, if a customer says “I need to fly to Paris next Friday”, the primary intent might be categorized as Book_Flight. The interface allows users to add, edit, and delete intents in real time.

Other examples:

- Ticket Inquiry
- How To
- Agent
- Sales

Entities are variables within an utterance that provide additional context and details relevant to the user’s request. Entities help the bot understand specific details used to fulfill the user’s request.

Using the previous example of booking a flight, the actionable entities that the system needs to extract are the destination (Paris) and the timeline (next Friday).

Entities answer questions such as “what, where, when, and how much.”

For example:

- Ticket number
- Location
- Date

Entities can be extracted through both rigid rules and contextual machine learning.

The first step in defining your model schema is to map out the specific conversational boundaries the AI needs to comprehend by defining Intents and Entities.

Intents

Remember, an intent is the underlying purpose that a user intends to achieve through their utterance. It represents the user's core goal, action or underlying motivation within a given message.

The Intents tab allows users to add, edit, and delete intents in real time.

NAME	UTTERANCES	DESCRIPTION	ACTIONS
Demo_Banking	3	Online Banking Demo	[Trash Icon]
None	0	Used when the user's input doesn't match a...	[Trash Icon]
Demo_ServiceAgent	22	Testing Style update	[Trash Icon]
Demo_SalesAgent	11	—	[Trash Icon]
Demo_AgentAssist	10	—	[Trash Icon]

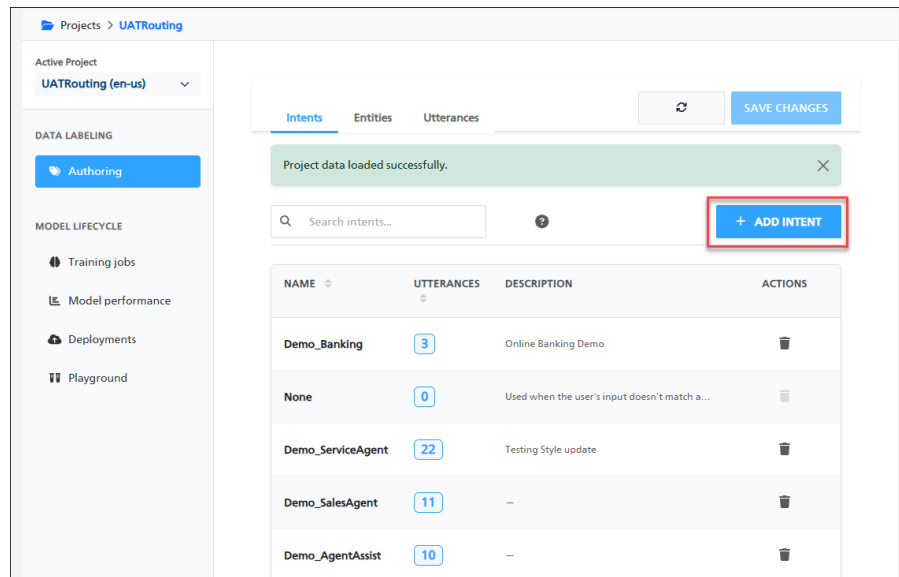
The following table describes the fields in the Intents grid.

Column	Description
Name	The name of the intent.
Utterances	The number of utterances tagged to the intent. Click on the number to see the utterances tagged to this intent.
Description	A description of the intent's purpose.
Actions	Administrators can delete intents using the delete icon in the actions column. Note: Deleted intents cannot be recovered.

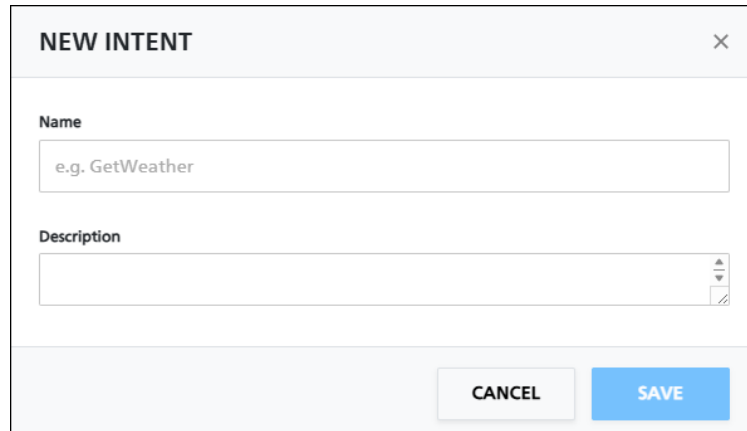
Add an intent

To add a new intent, follow the steps below:

1. Click the *Add Intent* button.

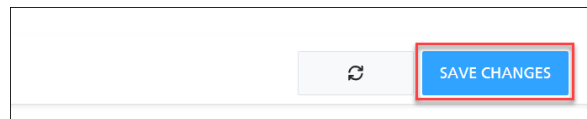


2. Enter the name of the new intent.



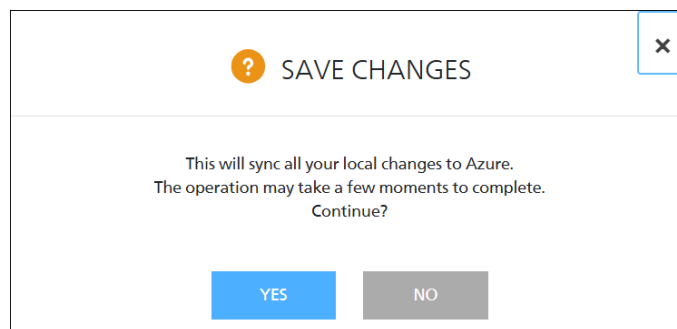
A dialog box titled "NEW INTENT" with a close button (X) in the top right corner. It contains two input fields: "Name" with the placeholder text "e.g. GetWeather" and "Description" which is empty. At the bottom right, there are two buttons: "CANCEL" and "SAVE".

3. Add a description for the intent.
4. Click *Save*. This saves your changes locally. You will also need to sync these changes to Azure.
5. Click *Save changes* in the top right corner of the screen.



A button labeled "SAVE CHANGES" with a blue background and white text, highlighted with a red rectangular border. To its left is a grey button with a circular refresh icon.

6. The following window will open. Click *Yes* to sync your changes.

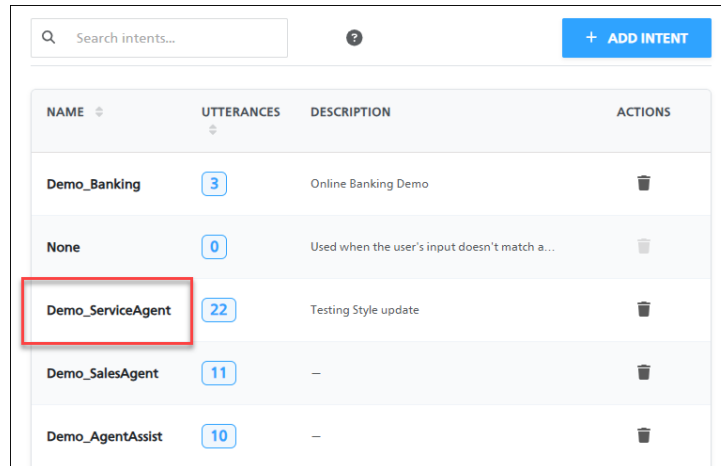


A confirmation dialog box titled "SAVE CHANGES" with a question mark icon and a close button (X) in the top right corner. The text inside reads: "This will sync all your local changes to Azure. The operation may take a few moments to complete. Continue?". At the bottom, there are two buttons: "YES" (blue) and "NO" (grey).

Edit an intent

To edit an intent, follow the steps below:

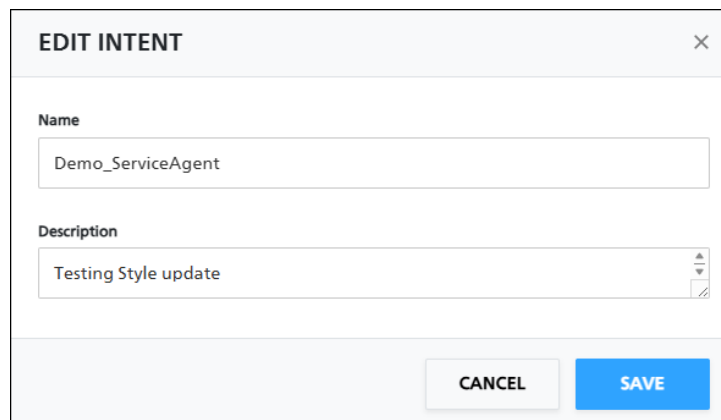
1. Click on the intent name in the grid.



Search intents... ? + ADD INTENT

NAME	UTTERANCES	DESCRIPTION	ACTIONS
Demo_Banking	3	Online Banking Demo	
None	0	Used when the user's input doesn't match a...	
Demo_ServiceAgent	22	Testing Style update	
Demo_SalesAgent	11	—	
Demo_AgentAssist	10	—	

2. Edit the intent name and description fields as needed.



EDIT INTENT

Name

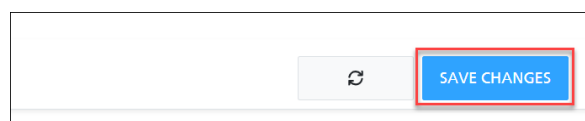
Demo_ServiceAgent

Description

Testing Style update

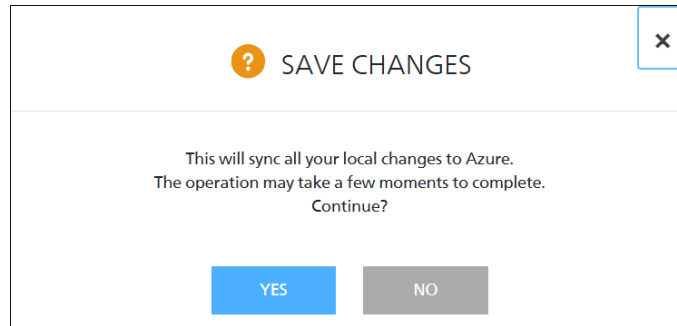
CANCEL SAVE

3. Click *Save*. This saves your changes locally. You will also need to sync these changes to Azure.
4. Click *Save changes* in the top right corner of the screen.



↻ SAVE CHANGES

5. The following window will open. Click *Yes* to sync your changes.

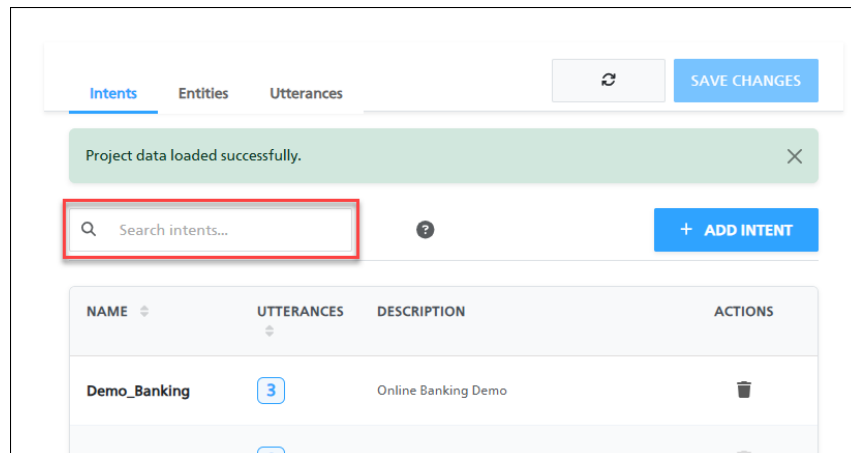


Search for an intent

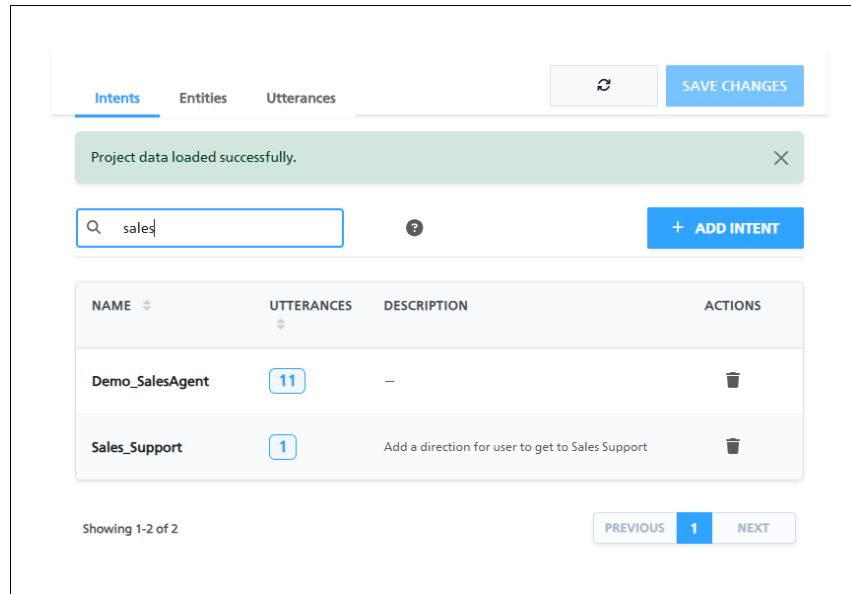
You may wish to see a list of intents containing a certain term, or validate if an intent already exists.

To search for an intent, follow the steps below:

1. If you know the term or contents of the intent, enter it in the search box.

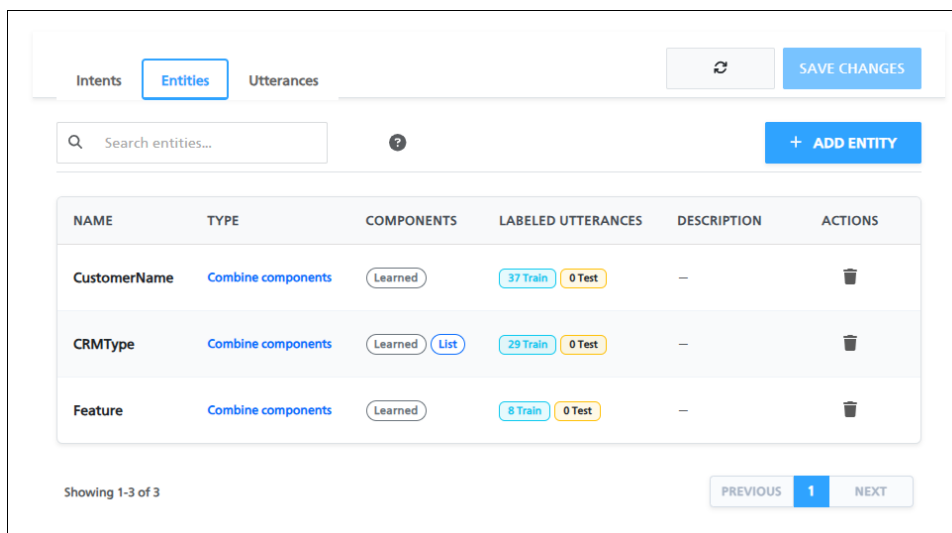


2. The results will filter accordingly in the intents grid.



Entities

Remember, entities represent the specific data points required to fulfill the intent.



The following table describes the fields in the Entities grid.

Column	Description
Name	The name of the entity.

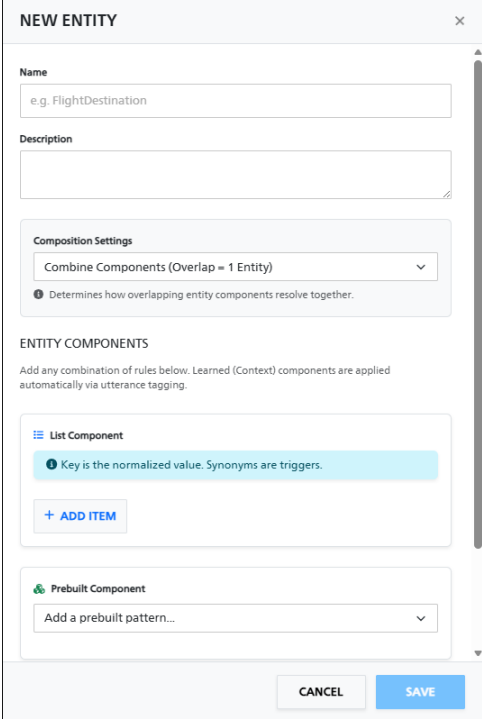
Column	Description
Type	Because a single entity can be defined by multiple components at the same time, their predictions will inevitably overlap. For example, a "Ticket Quantity" entity might utilize both a <i>Prebuilt</i> number component (to catch all numbers) and a <i>Learned component</i> (to understand that the number relates specifically to tickets). When a user says, "Book two tickets," both components will trigger on the word "two." The interface allows you to resolve these collisions through two distinct composition settings: • Combined Entities: This setting merges the overlapping components into a single, unified entity extraction. The model takes the union of the components, meaning you get the exact text boundaries predicted by the context-aware learned component, enriched with the normalized keys or standardized formatting from the list or prebuilt components. • Separate Entities: This option treats the overlapping predictions entirely independently. Instead of merging them, the model will return each overlapping component as its own distinct entity instance, allowing your downstream application logic to decide which specific extraction to prioritize or process.
Components	Rather than relying on a single method for extraction, Azure entities are highly modular. They can be composed of multiple distinct components, allowing the AI model to capture data through both rigid rules and contextual ML. • Learned Components: This is the most dynamic type of extraction. An entity does not inherently start with a learned component; it automatically becomes "learned" the moment you begin tagging it within the Utterance table. By highlighting text spans in your training utterances, you teach the underlying machine learning model to recognize the entity based on its surrounding context and sentence structure, allowing it to predict entities even if the exact words have never been seen before. • List Components: These rely on exact text matching against a fixed, closed set of related words and synonyms. For example, a "destination" entity might include a list component where "NYC", "The Big Apple," and "New York"

Column	Description
	all map to a normalized key. If any of these exact synonyms appear in the user's input, the entity is instantly recognized regardless of the surrounding context. • Prebuilt Components: These allow you to leverage out-of-the-box recognition for common, universally standardized data types. Without needing to supply any of your own training data, you can attach prebuilt components to instantly extract complex variables like Dates, Times, Numbers, Locations, etc. • Regex Components: For highly structured data that follows a strict, predictable pattern, such as flight number, an invoice code, or a product ID. These components use regular expressions to capture and extract the exact text match.
Labeled Utterances	Shows the number of utterances that have been labeled with each entity. The data is divided into training data and test data. For more information, see Utterance.
Description	A description of the entity.
Actions	Administrators can delete entities using the delete icon in the actions column. Note: Deleted entities cannot be recovered.

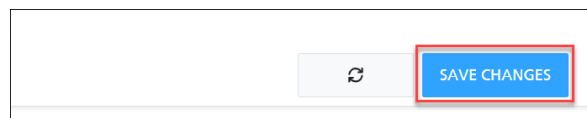
Add an entity

To add a new entity, follow the steps below:

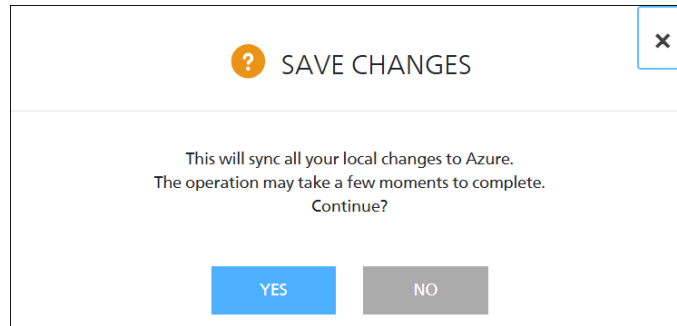
1. Click the *Add Entity* button.
2. The following window will open.



3. Enter the entity name.
4. Enter a description for the entity.
5. Select the composition setting. This setting determines how overlapping entity components resolve together. For more information, please refer to Entities.
6. Add any combination of list components, prebuilt components and regex components. Learned components are applied automatically via utterance tagging. For more information on entity components, please refer to Entities.
7. Click *Save*. This saves your changes locally. You will also need to sync these changes to Azure.
8. Click *Save changes* in the top right corner of the screen.



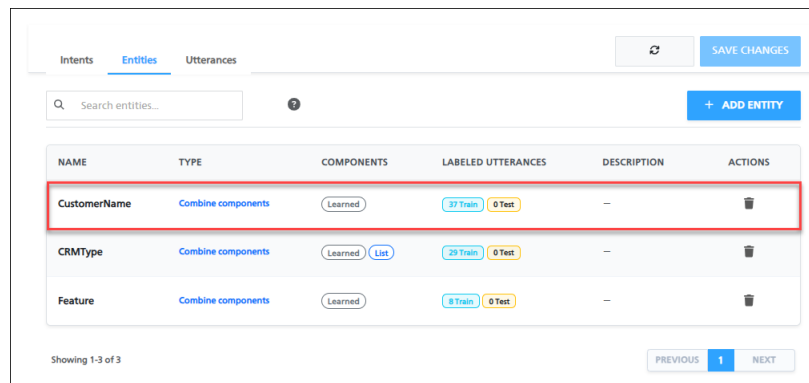
9. The following window will open. Click *Yes* to sync your changes.



Edit an entity

To edit an entity, follow the steps below:

1. Click on the entity in the grid.



2. The following window will open.

EDIT ENTITY [X]

Name

CustomerName

Description

Composition Settings

Combine Components (Overlap = 1 Entity) [v]

ⓘ Determines how overlapping entity components resolve together.

ENTITY COMPONENTS

Add any combination of rules below. Learned (Context) components are applied automatically via utterance tagging.

List Component

ⓘ Key is the normalized value. Synonyms are triggers.

+ ADD ITEM

Prebuilt Component

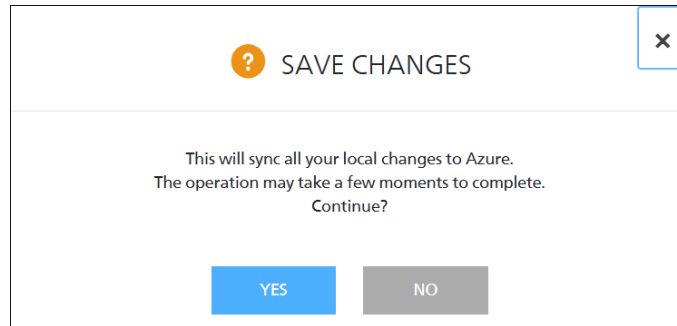
Add a prebuilt pattern... [v]

CANCEL SAVE

3. Edit the Name, Description, Composition Settings, and Entity Components as needed.
4. Click *Save*. This saves your changes locally. You will also need to sync these changes to Azure.
5. Click *Save changes* in the top right corner of the screen.

SAVE CHANGES

6. The following window will open. Click *Yes* to sync your changes.

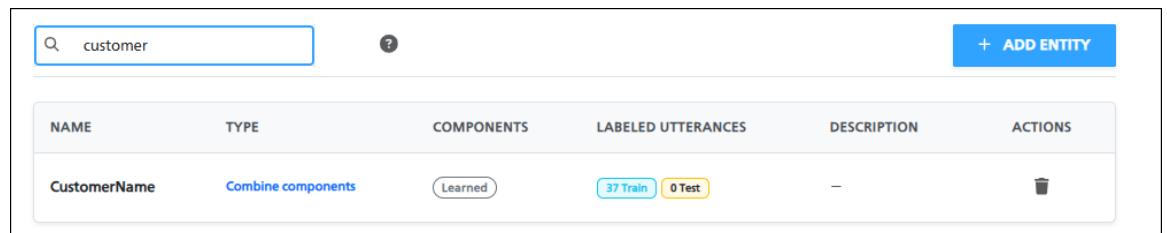


Search for an entity

You may wish to see a list of entities containing a certain term, or validate if an entity already exists.

To search for an entity, follow the steps below:

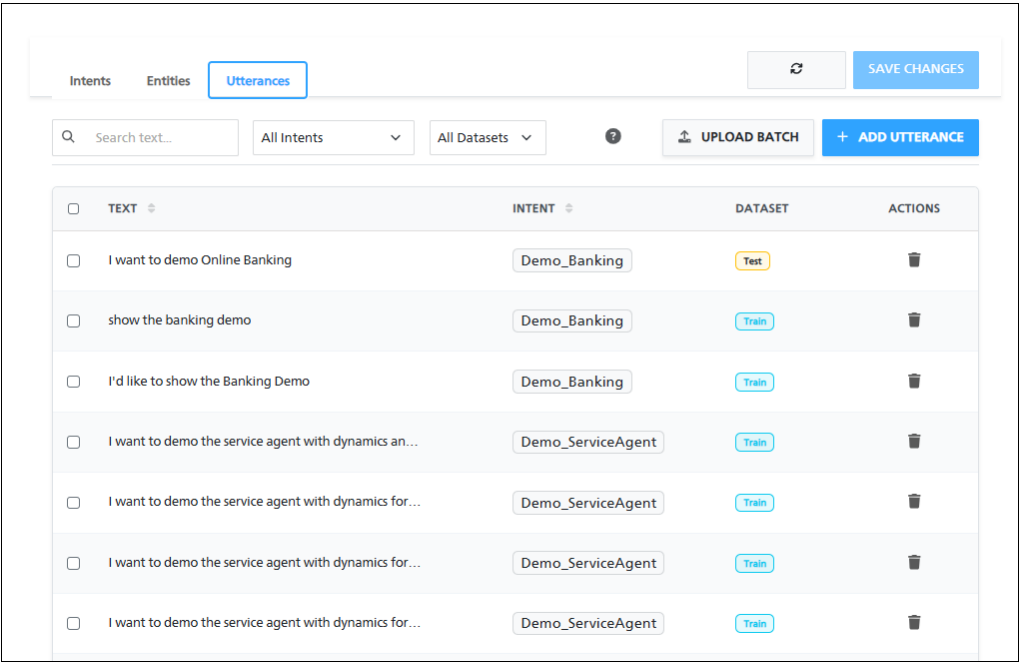
1. If you know the term or contents of the entity, enter it in the search box.
2. The results will filter accordingly in the entities panel.



Utterance

The core of the model building process occurs in the data labeling section. This process transforms raw training utterances into structured, machine-readable data.

The Utterances tab allows administrators to manage massive datasets using batch filtering, editing, creation and deletion of utterances. To prevent model ambiguity and to ensure high-quality training data, AI Studio does not allow for duplicate utterances.



The following table describes the fields in the Utterances grid.

Column	Description
Text	The utterance text.
Intent	The intent labeled in the utterance.
Dataset	The dataset to which the utterance is assigned. The dataset is divided into a training set and a testing set and can be assigned when creating or editing an utterance. The training set is used to train the model. This is the set which the model uses to learn the labeled utterances. The testing set is a blind set of data that is not used in the training process, but only used during model evaluation. After the model is trained, it uses the utterances in the testing set to make predictions.
Actions	Administrators can delete utterances using the delete icon. Note: Deleted utterances cannot be recovered.

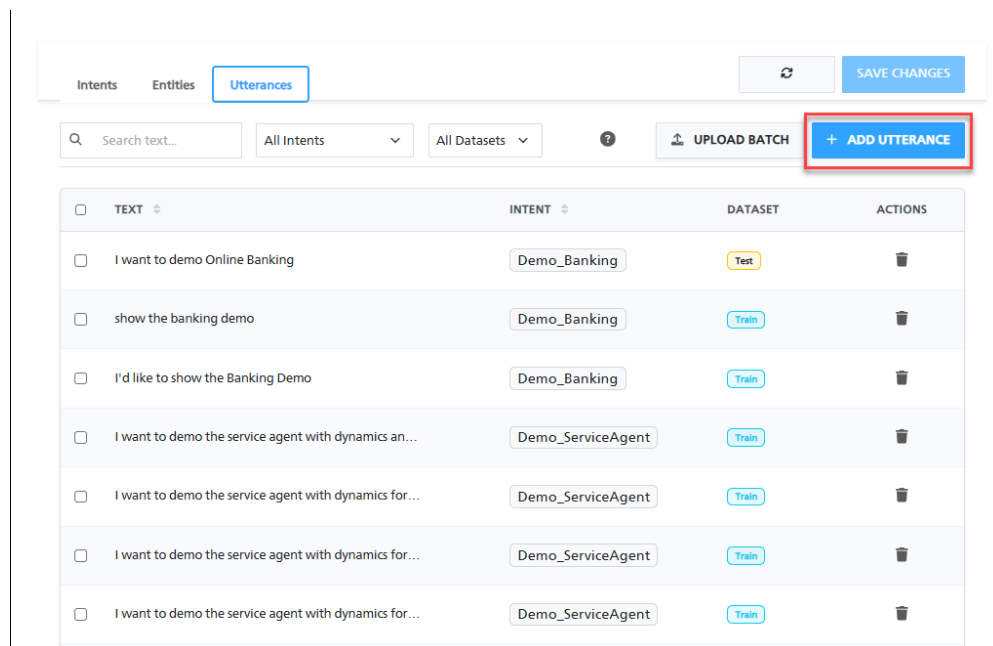
Add an utterance

There are two ways to add an utterance.

- Manually add an utterance using the *Add Utterance* button.
- Upload Batch

To manually add an utterance:

1. Click the *Add Utterance* button.



2. The following window will open.

The screenshot shows the 'NEW UTTERANCE' dialog box. It has a title bar with a close button. Inside, there is a text input field labeled 'Text (Highlight text to label entities)'. Below the text input, there are three dropdown menus: 'Language' (set to 'English (US)'), 'Intent' (set to 'Select intent...'), and 'Dataset' (set to 'Train'). At the bottom left, there is a small icon and the text 'Locked to default'. At the bottom right, there are two buttons: 'CANCEL' and 'SAVE'.

3. Enter your utterance text.
4. To label entities within the utterance, highlight the text.

The screenshot shows the 'NEW UTTERANCE' dialog box. At the top, there's a title bar with a close button. Below it, a text input field contains 'I need to demo banking for Computer Talk', with 'Computer Talk' highlighted. An 'Entity' dropdown menu is open, showing a list of entities: 'CustomerName', 'CRMTYPE', and 'Feature'. Below the text field, there are three dropdown menus: 'Language' (set to 'English (US)'), 'Select intent...' (with a dropdown arrow), and 'Dataset' (set to 'Train'). At the bottom, there are 'CANCEL' and 'SAVE' buttons.

5. A list of entities will open. Select the appropriate entity to label the text.
6. You will see a preview of the labeled utterance below the text field.

The screenshot shows the 'NEW UTTERANCE' dialog box after the entity has been assigned. The text input field still contains 'I need to demo banking for Computer Talk'. Below it, the 'Entity Preview' section shows the text 'I need to demo banking for Computer Talk CUSTOMERNAME', where 'CUSTOMERNAME' is the label assigned to 'Computer Talk'. The 'Language' and 'Dataset' dropdowns remain the same, and the 'Select intent...' dropdown is still open.

7. Select the *Intent* for this utterance from the drop-down list.

NEW UTTERANCE

Text (Highlight text to label entities)

I need to demo banking for Computer Talk

Entity Preview

I need to demo banking for Computer Talk CUSTOMERNAME

Language: English (US) Locked to default

Intent: Select intent... (Dropdown menu open showing: Demo_Banking, None, Demo_ServiceAgent, Demo_SalesAgent, Demo_AgentAssist, Identify_Feature, Demo_Municipality, Sales_Support)

Dataset: Train

CANCEL SAVE

8. Select the *Dataset*.

Language: English (US) Locked to default

Intent: Select intent...

Dataset: Train (Dropdown menu open showing: Train, Test)

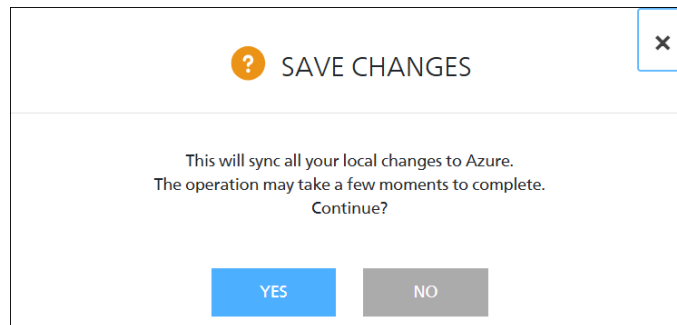
CANCEL SAVE

Note: If you have configured a project with more than one language, select the language of the utterance in this step.

7. Click *Save*. This saves your changes locally. You will also need to sync these changes to Azure.
8. Click *Save changes* in the top right corner of the screen.

SAVE CHANGES

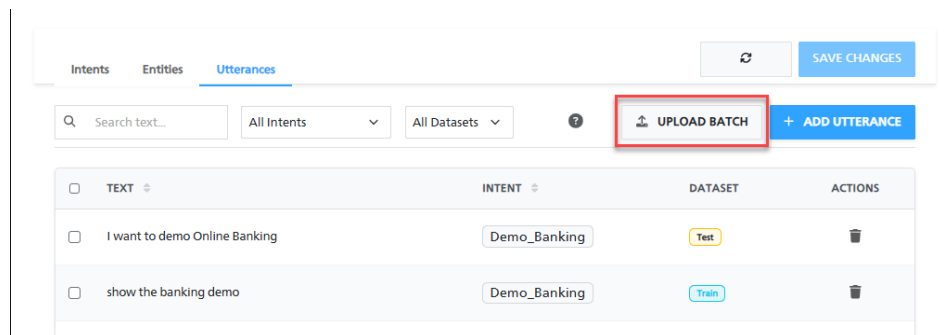
9. The following window will open. Click Yes to sync your changes.



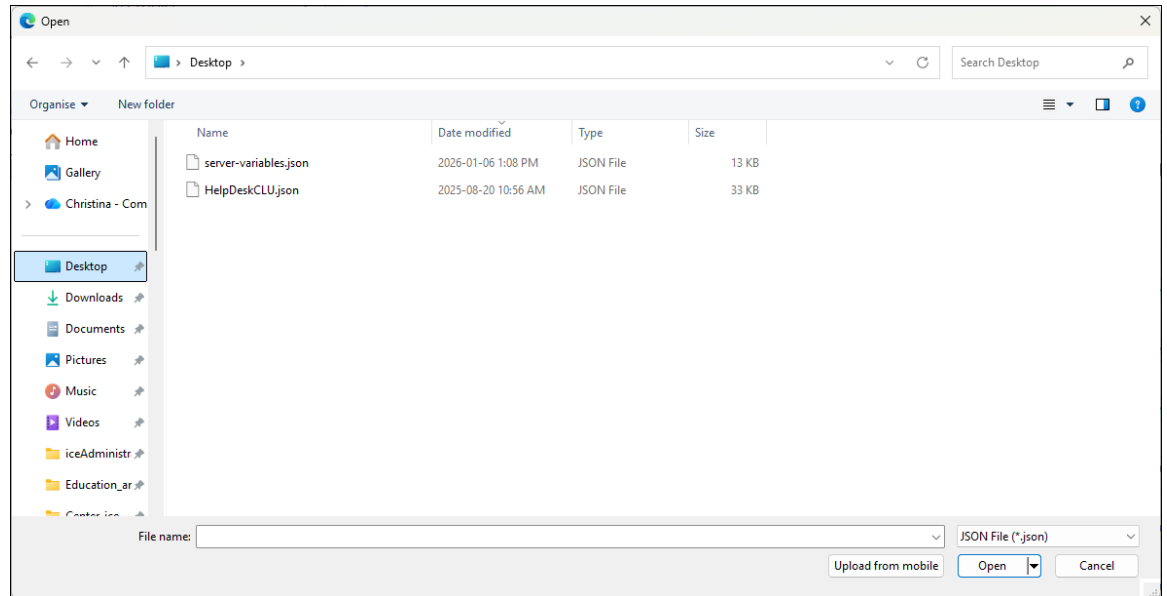
Alternatively, administrators can upload a JSON file containing all their utterances, tagged intents, and entities. If an intent or entity does not exist already in NLR, the system will create one based on the user upload. This allows for mass data uploads and efficiency.

To batch upload your utterances, follow the steps below:

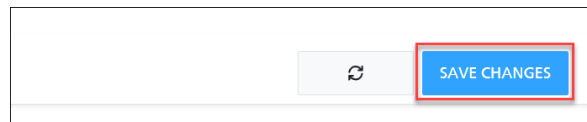
1. Click *Upload Batch*.



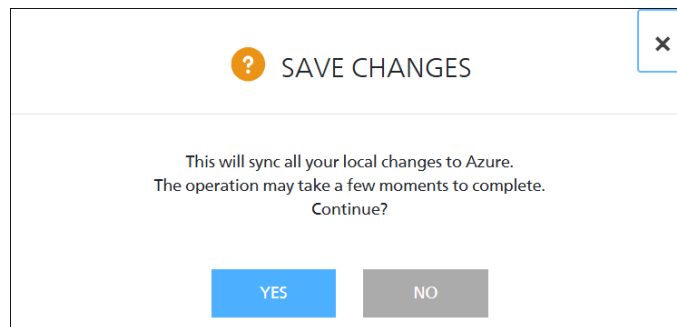
2. The file explorer window will open. Select your JSON file and click *Open*.



3. Click Save changes in the top right corner of the screen.



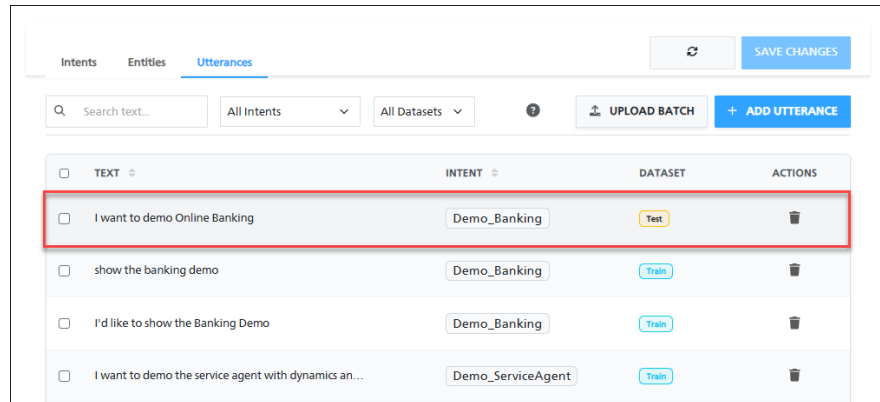
4. The following window will open. Click Yes to sync your changes.



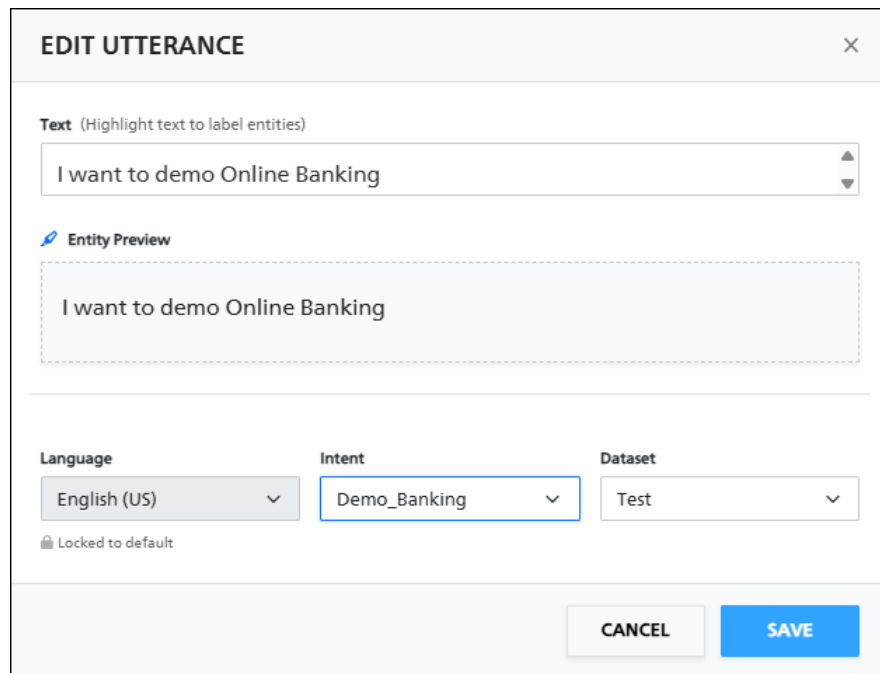
Edit an utterance

To edit an utterance, follow the steps below:

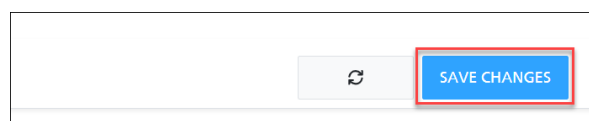
1. Click on the utterance in the grid.



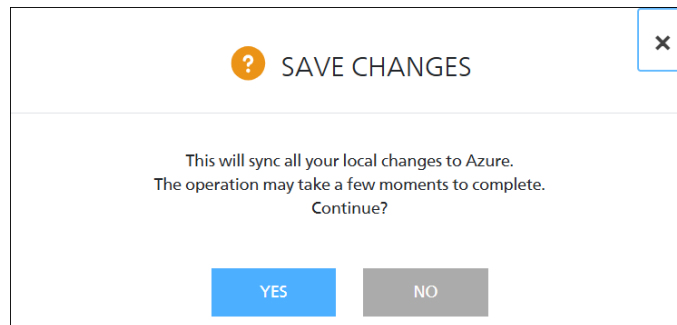
- The following window will open.



- Edit the Text, Entity labeling, Intent and Dataset as needed.
- Click Save. This saves your changes locally. You will also need to sync these changes to Azure.
- Click Save changes in the top right corner of the screen.



- The following window will open. Click Yes to sync your changes.



Search for an utterance

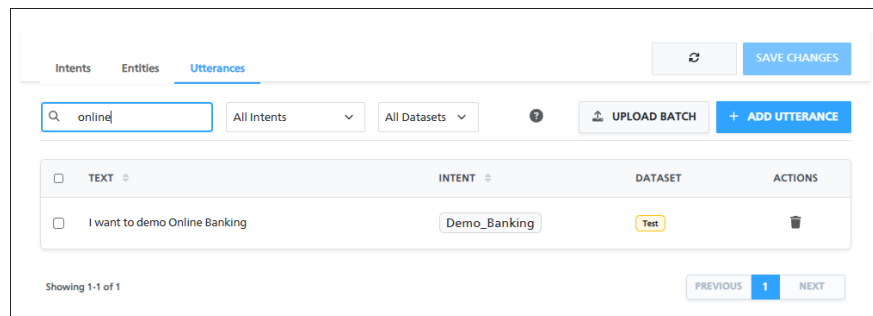
You may wish to see a list of utterances containing a certain term, or validate if an utterance already exists.

To search for an utterance, follow the steps below:

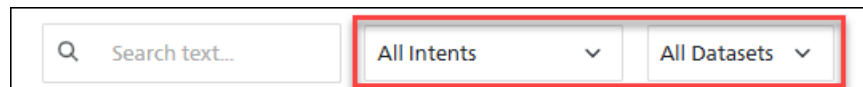
- If you know the term or contents of the utterance, enter it in the filter box.



- The results will filter accordingly in the utterances panel.



- You may further filter your utterances by intent or dataset.



Training Jobs

After your schema has been defined and your utterances have been labeled, the next step is to compile that raw data into a functional machine learning model.

TRAIN YOUR MODEL ?				▶ TRAIN MODEL
TRAINING JOB ID	STATUS	OUTPUT MODEL	CREATED ON	
 845318a0-04b9-48d8-9177-77f7b14...	 SUCCEEDED	 CSR	2026-04-14 05:13:54 PM	
 b728266c-000a-400d-aa69-2f6ed0e6...	 SUCCEEDED	 testing_training_v1	2026-04-14 03:58:26 PM	

To train the model, create a training job.

1. MODEL DETAILS

2. DATA SPLITTING

SELECT A MODE

You can train a new model or overwrite an existing one. Overwriting will replace the old model with the new data.

Create a new training model

Overwrite an existing model name

MODEL NAME *

e.g. v1.0-beta

TRAINING MODE

Standard training (free)

Advanced training

Configuration Summary

MODEL VERSION

Not defined

TRAINING MODE

Standard training (free)

STRATEGY

Not defined

CANCEL

NEXT →

Note: To train the model, the project must have a minimum of 15 utterances.

To create a training job, follow the steps below:

1. Click Train Model.



2. Select your mode. You can choose to create a new training model or overwrite an existing model. Overwriting will replace the old model with new data.

A form titled "SELECT A MODE" with a subtitle explaining that overwriting replaces the old model. It contains two radio button options: "Create a new training model" (selected) and "Overwrite an existing model name". Below these is a text input field labeled "MODEL NAME *" with the placeholder text "e.g. v1.0-beta".

3. Enter a unique model name. Using semantic versioning, such as v2-advanced, allows administrators to safely iterate on their models and maintain a historical record of different schema or dataset experiments.
4. Select the training mode. Standard training is optimized for testing, free and only for English projects. It is ideal for quick iterations and allows you to rapidly test how new intents or entity labels affect the model's baseline understanding.

Advanced training uses deeper and more rigorous computational resources. While the processing time is extended, this mode yields a significantly more accurate and comprehensive model for production.

Note: This will incur additional costs.

A section titled "TRAINING MODE" containing two selectable options. The first option, "Standard training (free)", is highlighted with a blue border and includes a lightning bolt icon and text about faster training times for English projects. The second option, "Advanced training", is in a light gray box and includes a gear icon and text about fine-tuned neural network transformer models.

5. Click Next. This will take you to the data splitting section.
6. Configure your data split.

Automatic Percentage Split allows the system to automatically divide the dataset into the industry standard (80/20) split. This means that 80% of the utterances will be used to train the model, and 20% will be reserved for testing. Administrators can manually adjust the splits by changing the percentage values in the training and testing fields respectively.

CONFIGURE YOUR DATA SPLIT

☒ **Automatic Percentage Split**

TRAINING %

80

%

TESTING %

20

%

☐ **Manual Dataset Allocation**

Manual dataset allocation allows administrators to use the testing and training assignments from the utterances page. When creating or editing an utterance, you can manually assign either the training or testing dataset. Manual allocation uses these numbers to train and test the model.

CONFIGURE YOUR DATA SPLIT

☐ Automatic Percentage Split

☒ **Manual Dataset Allocation**

Uses the 'Dataset' field (Train/Test) assigned to each utterance in your project.

TRAINING SET
49 96%

TESTING SET
2 4%

- Click *Start Training*. Once the model has initiated training, you can view its progress in the dashboard.

Model Performance

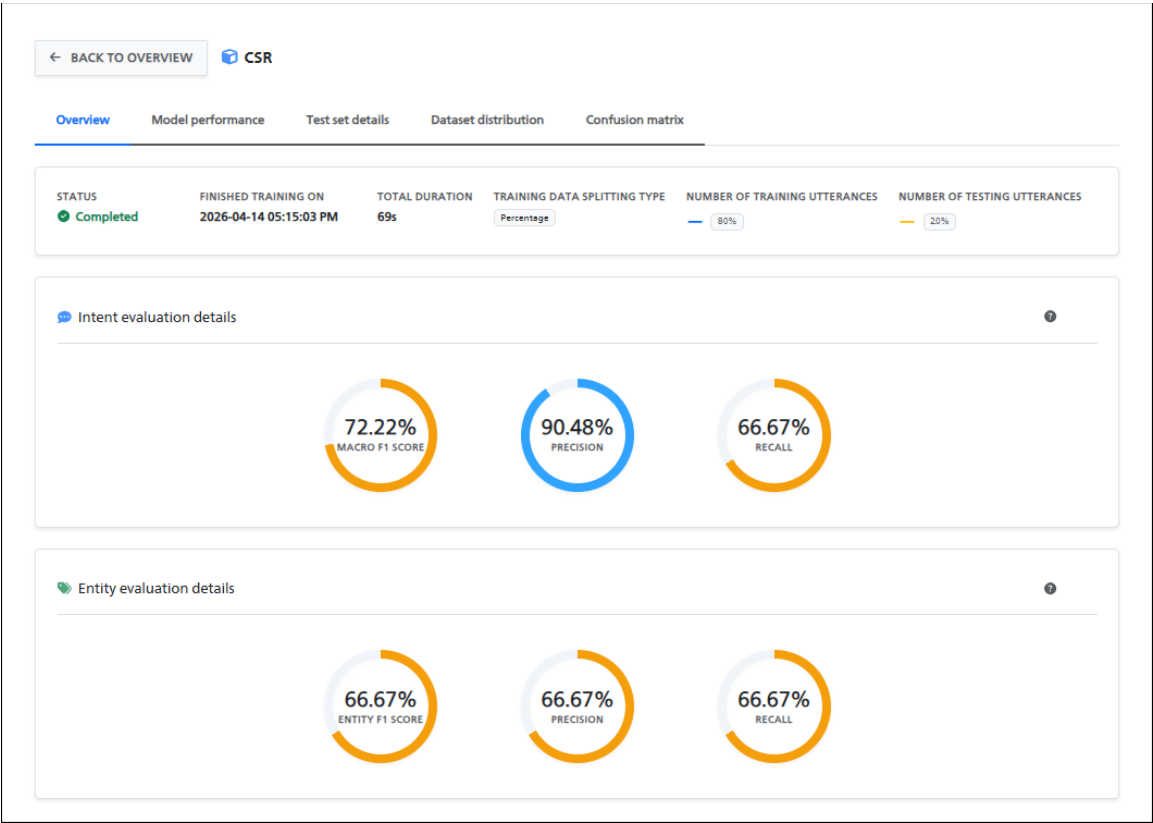
Once a training job is complete, you can view the model performance in the dashboard. The information in this section allows administrators to move beyond simply knowing if a model works, to understanding if there are any weaknesses, or areas for improvement.

EVALUATE YOUR MODEL
?

MODEL NAME	CONFIG VERSION	STATUS	CREATED ON	EXPIRATION DATE	RERUN EVALUATION	ACTIONS
CSR	v2022-09-01	✓ SUCCEEDED	2026-04-14 05:15:03 PM	2026-08-31		
Test	v2022-09-01	✓ SUCCEEDED	2026-04-07 11:14:45 AM	2026-08-31		
testing_training_v1	v2022-09-01	✓ SUCCEEDED	2026-04-14 03:59:34 PM	2026-08-31		

The model performance page allows you to view all the training jobs that have run.

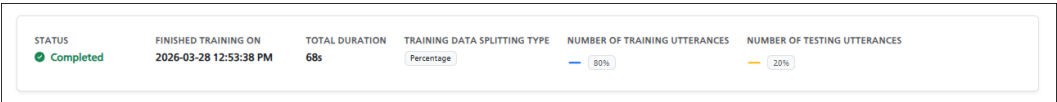
Click on a model in the grid to view the performance details.



Overview

The overview panel displays core macro metrics including overall model statistics, intent evaluation details, entity evaluation details, and guidance and recommendations.

Model Overview

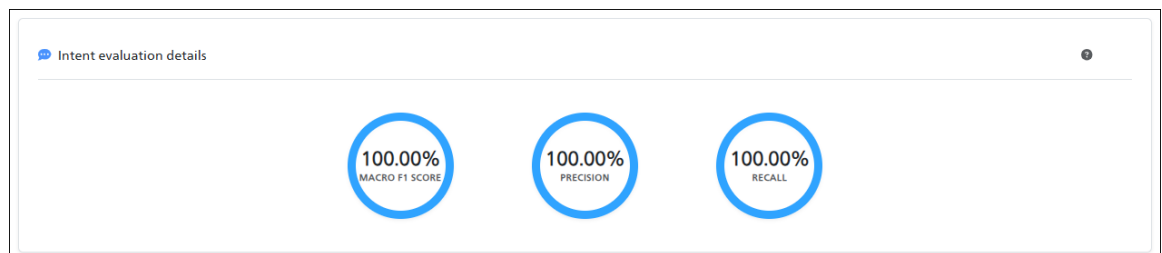


The following table describes each column in the overview.

Column	Description
Status	The status of the training job.
Finished Training On	The date and time that the training was completed.
Total Duration	The total amount of time taken to train the model.
Training Data Splitting Type	The method in which the model data is split.
Number of Training Utterances	The number or percentage of utterances in the training set.
Number of Testing Utterances	The number or percentage of utterances in the testing set.

Intent Evaluation Details

Intent metrics are calculated by comparing the model's predictions against the actual labels in your Testing set.



Precision: How often the model is correct when it predicts a specific intent. This is also a measure of quality. A high precision score indicates a cautious model that rarely makes false claims.

Recall: Out of all the actual instances of an intent in your test set, how many did the model successfully find? This is also a measure of quantity. Recall measures the model's completeness. A high recall score means the model casts a wide net and rarely misses a target.

F1 Score: A balanced average between Precision and Recall. This can be used as your primary indicator of health. A score of 80% or higher is healthy. 50 – 79% is considered a warning. Less than 50% is considered poor health.

Entity Evaluation Details

Entity metrics are calculated by comparing the model's predictions against the actual labels in your Testing set.

Precision: How often the model is correct when it predicts a specific entity. This is also a measure of quality.

Recall: Out of all the actual instances of an entity in your test set, how many did the model successfully extract? This is also a measure of quantity.

F1 Score: A balanced average between Precision and Recall. This can be used as your primary indicator of health. A score of 80% or higher is healthy. 50 – 79% is considered a warning. Less than 50% is considered poor health.

Guidance and Recommendations

The guidance and recommendations displayed in this section are provided by Azure to improve the model. It will typically cover three points:

- Does the training set have enough data
- Are all intents or entities present in the test set
- Is there an unclear distinction between intents or entities

For example, in the following screenshot, to improve the accuracy of the model, each intent should have a minimum of 15 utterances in your training data. Azure identifies the intents that do not meet the recommended minimum.

GUIDANCE & RECOMMENDATIONS

Review these specific issues to improve your model's overall accuracy.

1 Issue Detected

> **Intents without enough labels in training set**

We recommend a minimum of 15 utterances per intent in your training data.

- > None
- > check_balance
- > pay_bill
- > recent_transactions

Model Performance

The model performance is divided into two views, intent and entity. Click between the two buttons to see the different views.

[← BACK TO OVERVIEW](#)

Banking

Overview

Model performance

Test set details

Dataset distribution

Confusion matrix

INTENT

ENTITY

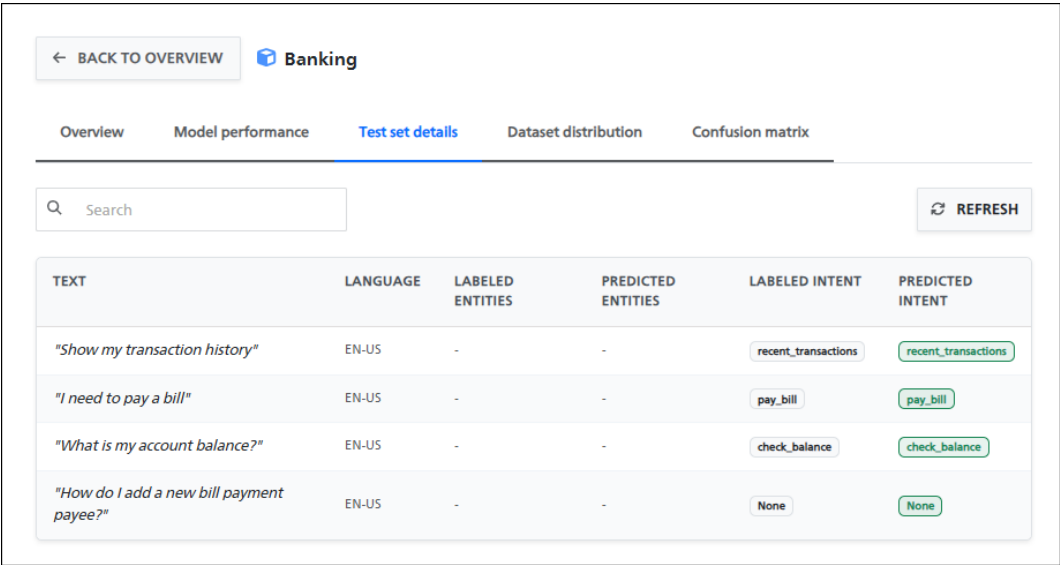
REFRESH

INTENT	PRECISION %	RECALL %	INTENT F1 SCORE	TRAINING LABELS	TESTING LABELS
recent_transactions	100%	100%	100%	5	1
pay_bill	100%	100%	100%	5	1
check_balance	100%	100%	100%	5	1
None	100%	100%	100%	5	1

On each page, you can see the precision, recall, and F1 scores for each intent or entity, as well as the number of labeled utterances in the training set, and the testing set.

Test Set Details

While the macro metrics provide a bird’s eye view of model health, diagnosing the exact root cause of an error is done by looking at the raw data. The test set panel serves as the ground-level inspection tool, allowing users to review the specific performance of every individual utterance evaluated during testing.

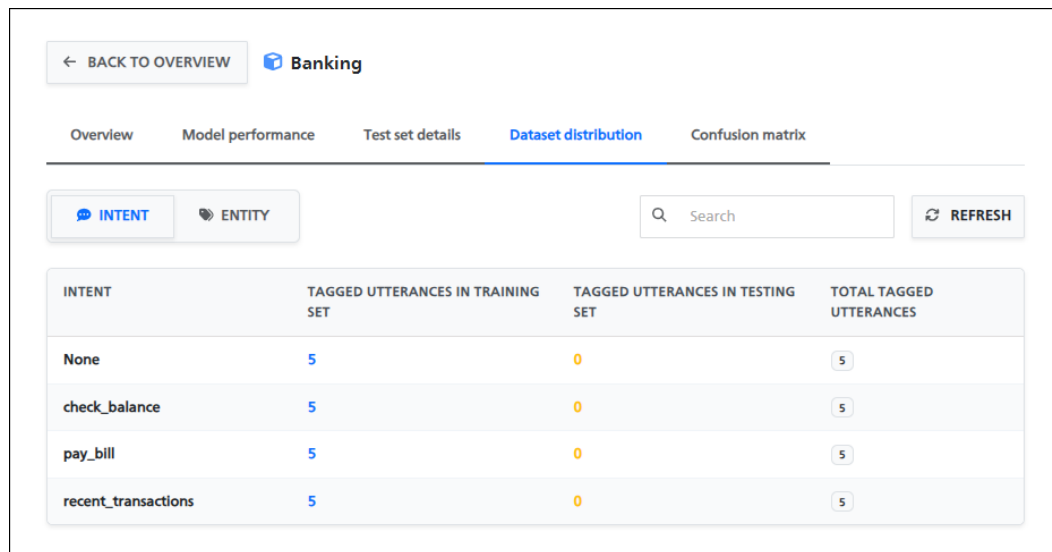


The following table describes each of the columns in the panel.

Column	Description
Text	The utterance text.
Language	The language of the utterance.
Labeled Entities	The user labeled entity.
Predicted Entities	The entity the model predicted.
Labeled Intent	The user labeled intent.
Predicted Intent	The intent the model predicted.

Dataset Distribution

The dataset distribution page is split into two views: intent and entity. Click on the buttons to switch the view.



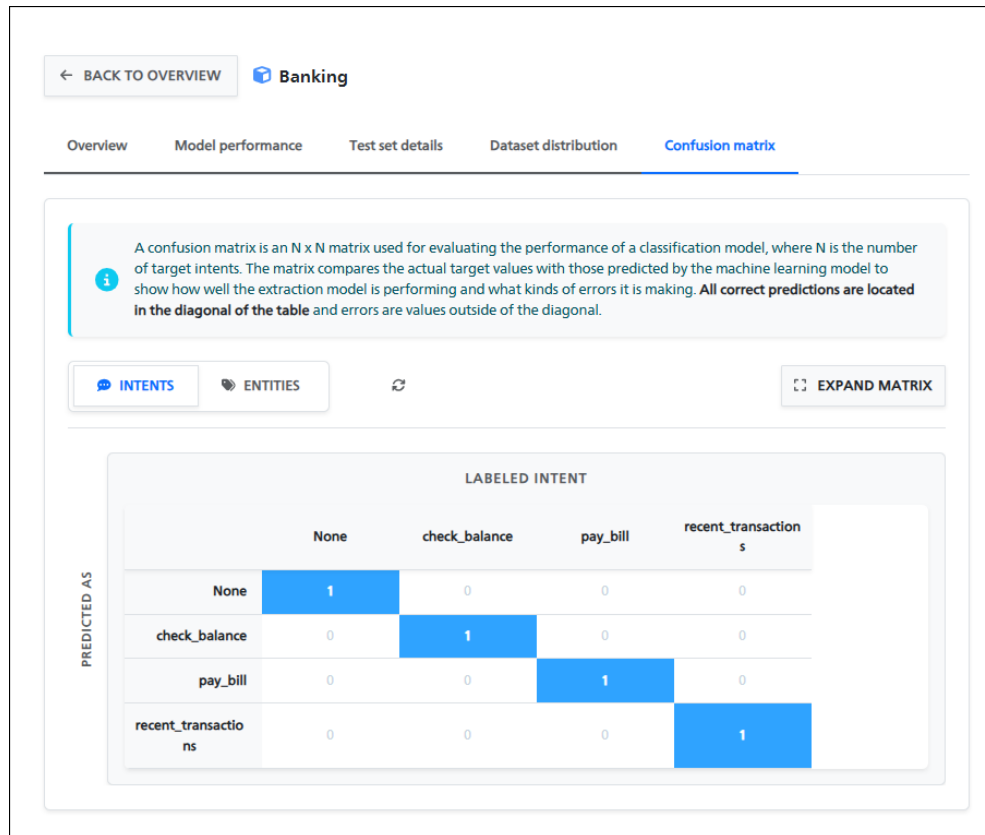
← BACK TO OVERVIEW Banking			
Overview	Model performance	Test set details	Dataset distribution
INTENT	ENTITY	Search	REFRESH
INTENT	TAGGED UTTERANCES IN TRAINING SET	TAGGED UTTERANCES IN TESTING SET	TOTAL TAGGED UTTERANCES
None	5	0	5
check_balance	5	0	5
pay_bill	5	0	5
recent_transactions	5	0	5

Each view displays the intent or entity, the number of tagged utterances in the training set, the number of tagged utterances in the testing set and the total tagged utterances.

This panel shows the raw distribution between the training and testing dataset per instance and entity. The splits can be assigned manually or automatically using percentages.

Confusion Matrix

A confusion matrix is an $N \times N$ matrix used for evaluating the performance of a classification model, where N is the number of target intents. The matrix compares the actual target values with those predicted by the machine learning model to show how well the extraction model is performing and identify what types of errors it is making. All correct predictions are found in the diagonal of the table, and errors are values outside of the diagonal.

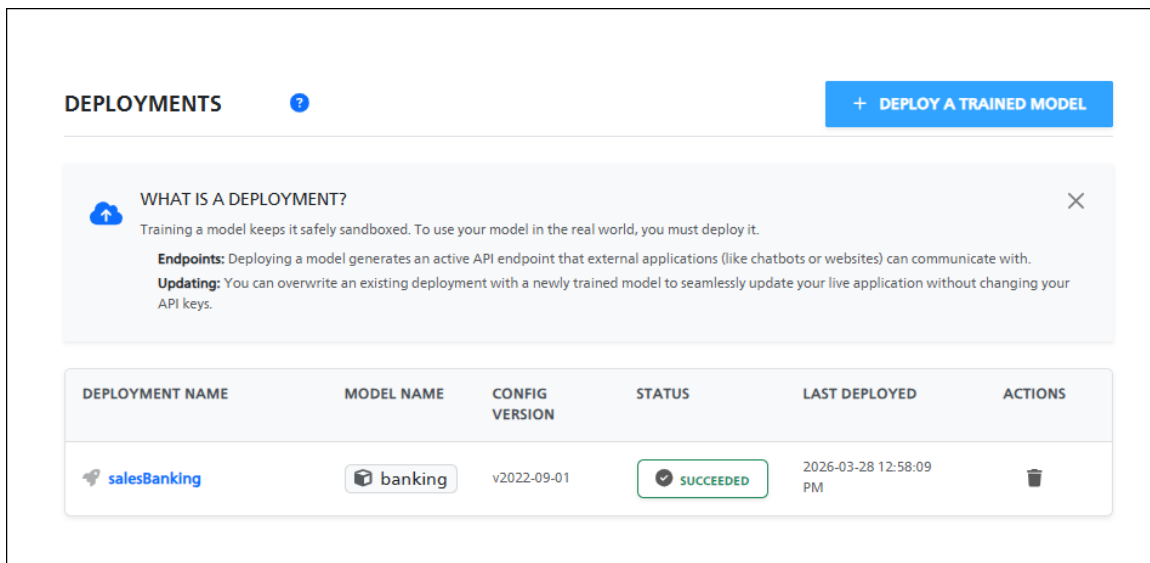


The confusion matrix serves as the primary diagnostic tool for evaluating the model's predictive accuracy, functioning as an interactive heatmap. By cross-referencing the model's predicted labels (rows) to the ground truth labels (columns), the matrix highlights structural weaknesses within the training data. A healthy model will display a high concentration of true positives along the main diagonal axis, indicating that the predicted and actual values align. Any red tinted cells outside the diagonal exposes exact points of confusion.

Deployments

Once a training job has successfully completed and a model has been created from the data, it is ready for deployment. You cannot use your model in ice until it has been deployed.

Note: If you have made changes to your model, you will need to re-train and redeploy the model to see the changes in ice.



DEPLOYMENTS ? [+ DEPLOY A TRAINED MODEL](#)

WHAT IS A DEPLOYMENT? ×

Training a model keeps it safely sandboxed. To use your model in the real world, you must deploy it.

Endpoints: Deploying a model generates an active API endpoint that external applications (like chatbots or websites) can communicate with.

Updating: You can overwrite an existing deployment with a newly trained model to seamlessly update your live application without changing your API keys.

DEPLOYMENT NAME	MODEL NAME	CONFIG VERSION	STATUS	LAST DEPLOYED	ACTIONS
 salesBanking	 banking	v2022-09-01	 SUCCEEDED	2026-03-28 12:58:09 PM	

To deploy a model, follow the steps below:

1. Click Deploy a Trained Model.



DEPLOYMENTS ? [+ DEPLOY A TRAINED MODEL](#)

2. Select to create a new deployment or to overwrite an existing deployment.

 **DEPLOY A TRAINED MODEL**

CREATE OR UPDATE A DEPLOYMENT

You can create a new deployment name or overwrite an existing deployment name to add a trained model to them.

☒ **Create a new deployment**

☐ **Overwrite an existing deployment**

You can overwrite an existing deployment with a newly trained model to update your live application without changing your API keys.

3. If you are creating a new deployment, enter a name. If you are overwriting an existing deployment, select the deployment from the list.

DEPLOYMENT NAME *

Create a name for the deployment

DEPLOYMENT NAME *

Select an existing deployment... ▼

4. Assign a model to your deployment.

ASSIGN A MODEL

Assign a trained model to your deployment name.

MODEL NAME *

Select a model... ▼

5. Click Deploy Model.

Playground

Once your model has been deployed, the integrated testing playground provides a sandboxed environment to interact with your bot. This module features three distinct testing modalities to ensure the model behaves predictably across different conversational scenarios.

The screenshot displays the AI Studio Playground interface. At the top, the 'Evaluation Mode' is set to 'SINGLE QUERY', with options for 'MULTI-TURN SANDBOX' and 'BATCH TESTING'. The 'Configuration' section on the left includes a 'PROJECT NAME' field with 'SalesBanking' and a 'DEPLOYMENT NAME *' dropdown menu. The 'Query Sandbox' section features a text input area with the placeholder 'Enter test query or upload a file...', an 'UPLOAD FILE' button, and a 'RUN PREDICTION' button. A 'CLEAR' button is also present. The 'Results' section at the bottom shows a placeholder for analysis results with the text 'Run a test to see analysis results.'

Single Query

The Single Query mode is designed for the rapid, isolated validation of your deployed model. Administrators can input a standalone text utterance and

instantly review the raw JSON response returned by the API. This interface explicitly breaks down the model's logic, displaying the predicted top intent, the associated confidence score, and a clear mapping of any extracted entities. It is the ideal tool for quickly testing how the model handles a specific, complex phrase without the influence of any prior conversational context.

To test the model, first select your deployment in the deployment name dropdown.

Enter your query in the query field and click Run Prediction.

Your result will be displayed in the results field.

The screenshot displays the AI Studio interface for testing a model. At the top, the 'Evaluation Mode' is set to 'SINGLE QUERY'. Below this, the 'Configuration' section on the left includes a 'PROJECT NAME' field with 'SalesBanking' and a 'DEPLOYMENT NAME *' dropdown menu with 'salesBanking' selected. The 'Query Sandbox' section on the right contains a text input field with the query 'What is my account balance?', an 'UPLOAD FILE' button, and a 'RUN PREDICTION' button. The 'Results' section at the bottom shows the 'TOP INTENT' as 'check_balance' with a confidence score of '90.68%' and 'ENTITIES FOUND' as 'No entities found.'. A 'SHOW ANALYSIS DETAILS' button is located at the bottom of the results section.

Evaluation Mode: ?	
SINGLE QUERY	MULTI-TURN SANDBOX
BATCH TESTING	

Configuration

PROJECT NAME

SalesBanking

DEPLOYMENT NAME *

salesBanking

Query Sandbox

What is my account balance?

UPLOAD FILE

RUN PREDICTION

Results

TOP INTENT check_balance 90.68%	ENTITIES FOUND No entities found.
---	---

SHOW ANALYSIS DETAILS

To see more detailed results, expand the Show analysis details section.

^ HIDE ANALYSIS DETAILS

All Intent Scores

INTENT	SCORE
check_balance	90.68%
recent_transactions	84.89%
None	83%
pay_bill	56.01%

Extracted Entities Details

ENTITY	TEXT	CONFIDENCE
--------	------	------------

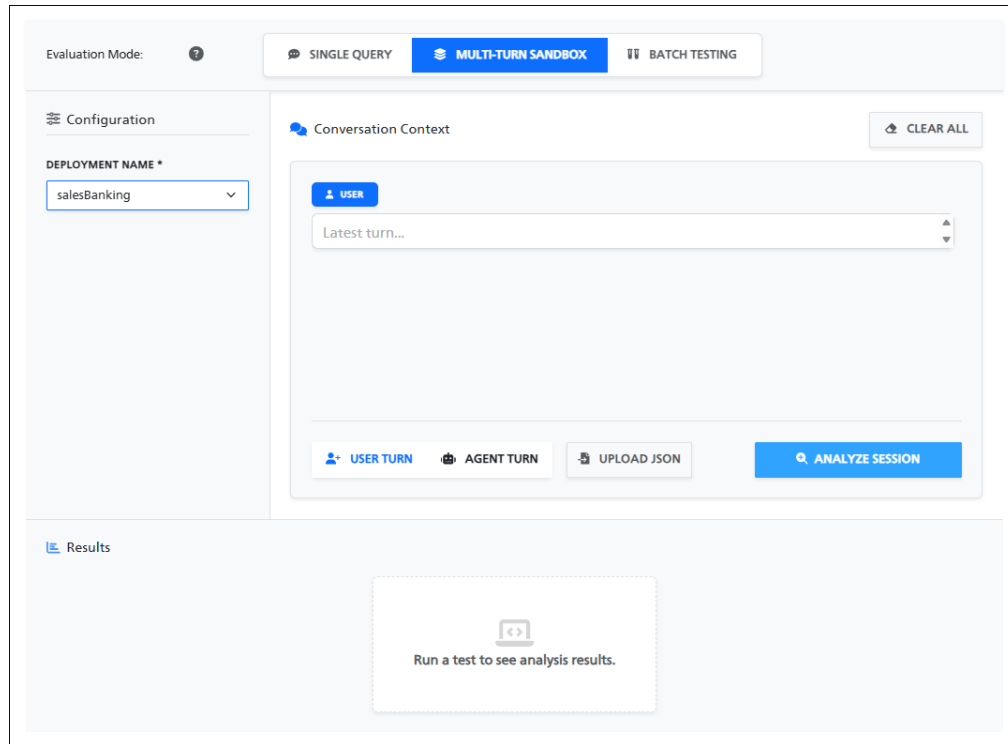
</> Raw JSON Response

```
{
  "success": true,
  "projectName": "SalesBanking",
  "deploymentName": "salesBanking",
  "query": "What is my account balance?",
  "language": "en",
  "topIntent": "check_balance",
  "topIntentConfidence": 0.9068339,
  "allIntents": [
    {
      "intent": "check_balance",
      "confidence": 0.9068339
    },
    {
      "intent": "recent_transactions",
      "confidence": 0.8489068
    },
    {
      "intent": "None",
      "confidence": 0.83
    },
    {
      "intent": "pay_bill",
      "confidence": 0.5601
    }
  ]
}
```

In the analysis details, administrators can view the top intents and their confidence scores, as well any entities. You can also view the raw JSON response returned by the API.

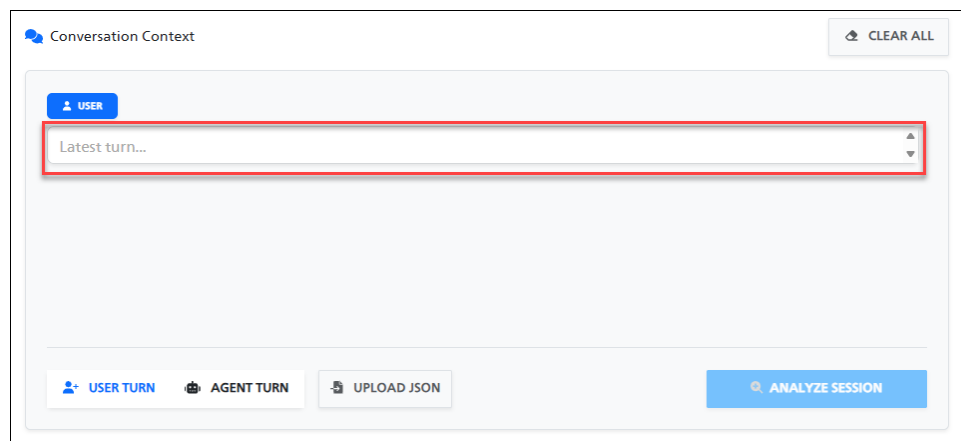
Multi-turn Sandbox

To test the model's ability to predict the intent and entities of the overall conversation, administrators can use the multi-turn sandbox. This evaluates the overall prediction and entities throughout the conversation.

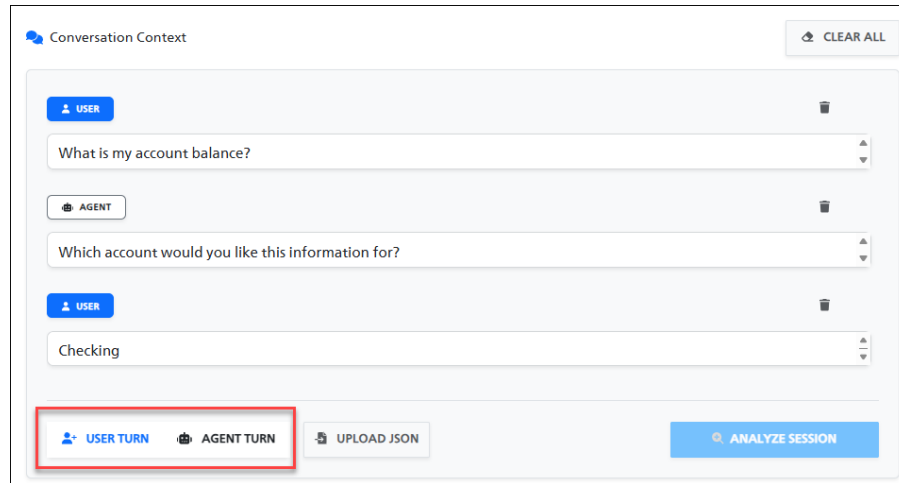


To test the model using the multi-turn sandbox, follow the steps below:

1. Begin by entering your first user utterance.

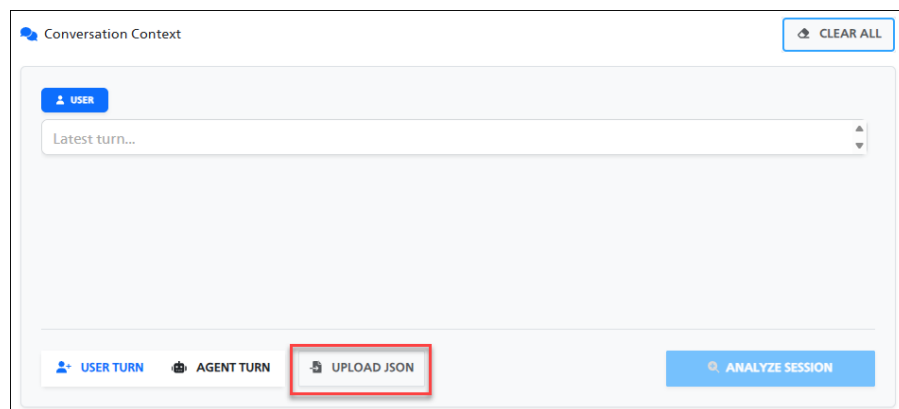


2. To add second utterance, click the related button below. For another user utterance, click User Turn. To add an agent utterance, click Agent Turn.



The screenshot shows the 'Conversation Context' window. At the top right is a 'CLEAR ALL' button. The main area contains three conversation turns. Each turn has a role label (USER or AGENT) in a blue box, a text input field, and a trash icon. The first turn is from a USER with the text 'What is my account balance?'. The second turn is from an AGENT with the text 'Which account would you like this information for?'. The third turn is from a USER with the text 'Checking'. At the bottom, there are three buttons: 'USER TURN' (highlighted with a red box), 'AGENT TURN', and 'UPLOAD JSON'. To the right of these is an 'ANALYZE SESSION' button.

3. You may also upload a JSON doc with the conversation context.



This screenshot shows the 'Conversation Context' window with only one turn visible, labeled 'USER' with the text 'Latest turn...'. The 'UPLOAD JSON' button at the bottom is highlighted with a red box. The 'ANALYZE SESSION' button is also visible.

4. When you have finished entering your conversation utterances, click Analyze Session.
Note: You must have a deployment name selected to be able to analyze the session.
5. The results will be displayed in the Results section.

Results

Notice:

The result may not be accurate since this ConversationalAI task is using neither multi-turn trained model nor LLM.

PREDICTED SESSION INTENT

check_balance

ENTITIES FOUND

No entities found.

SHOW ANALYSIS DETAILS

For more detailed information, expand the analysis details.

Results

Notice:

The result may not be accurate since this ConversationalAI task is using neither multi-turn trained model nor LLM.

PREDICTED SESSION INTENT

check_balance

ENTITIES FOUND

No entities found.

HIDE ANALYSIS DETAILS

TURN	PARTICIPANT	UTTERANCE	EXTRACTED ENTITIES DETAILS
1	USER	"what is my account balance?"	
2	AGENT	"let me look that up for you"	

</> Raw JSON Response

```
{
  "kind": "ConversationalAIResult",
  "result": {
    "conversations": [
      {
        "id": "conversation_test_1",
        "intents": [
          {
            "name": "check_balance",
            "type": "action",
            "entities": [],
            "conversationItemRanges": [
              {
                "offset": 0,
                "count": 1
              }
            ]
          }
        ]
      }
    ]
  }
}
```

Batch Testing

Batch testing serves as a comprehensive quality assurance sweep. Rather than testing phrases manually, administrators can upload entire datasets of labeled utterances and run them against the predictive agent.

The screenshot shows the 'Batch Testing' section of the AI Studio interface. At the top, there's a header 'Evaluation Mode:' with a help icon. Below it are three tabs: 'SINGLE QUERY', 'MULTI-TURN SANDBOX', and 'BATCH TESTING' (which is active). Under the 'BATCH TESTING' tab, there's a 'DEPLOYMENT NAME *' dropdown menu with the text 'Select a deployment...'. To the right of the dropdown are two buttons: 'UPLOAD FILE' and 'START TEST'. Below these buttons is a red link that says '* Upload a dataset'. The main content area is titled 'UPLOAD A DATASET TO BEGIN' with a document icon. Below the title, it says 'We support structured CSV and JSON files for batch evaluations.' There are two columns: 'CSV Format' and 'JSON Format'. The 'CSV Format' column has a green CSV icon and says 'Your file must include 3 columns separated by commas, in this exact order:' followed by a list: 'Text', 'Labeled intent', and 'Labeled entities'. The 'JSON Format' column has a yellow JSON icon and says 'Your file must be a JSON array of objects containing the following keys:' followed by a list: 'text', 'expectedIntent', and 'expectedEntities'.

Upload a CSV or JSON file of queries to rapidly evaluate bulk performance outside of your official training sets.

Note: Your CSV or JSON files must be structured in the following formats.

CSV Format

Your file must include three columns separated by commas, in this exact order:

1. Text

2. Labeled intent
3. Labeled entities

JSON Format

Your file must be a JSON array of objects containing the following keys:

1. text
2. expectedIntent
3. expectedEntities

To upload a CSV or JSON file, follow the steps below:

1. Select your deployment name.
2. Click "Upload File" and select your file. Click Okay.
3. Click Start Test. Your results will be displayed below.

 BATCH ACCURACY SCORE: 100% The model successfully predicted the expected intent for 100% of the evaluated queries.					
#	LABELLED INTENT	PREDICTED INTENT	LABELLED ENTITIES	PREDICTED ENTITIES	MATCH
1	"My laptop screen is flickering and showing gr..." Hardware_Fault	Hardware_Fault 95.9%	-	Component_Type: "screen" (100%) Component_Type: "screen" (100%)	✓
2	"The keyboard on my workstation isn't respon..." Hardware_Fault	Hardware_Fault 95.1%	-	Component_Type: "workstation" (100%)	✓
3	"I need a new monitor for my home office set..." Service_Request	Service_Request 86.7%	-	Component_Type: "monitor" (100%)	✓
4	"Can I request a license for Adobe Creative Clo..." Service_Request	Service_Request 89.8%	-	-	✓
5	"Excel keeps crashing whenever I try to open t..." Software_Support	Software_Support 91.6%	-	Software_Name: "Excel" (100%)	✓
6	"How do I install the latest updates for Windo..." Software_Support	Software_Support 87.2%	-	-	✓
7	"I've been locked out of my account after thre..." Account_Access	Account_Access 92.1%	-	-	✓
8	"Can you help me reset my corporate login pas..." Account_Access	Account_Access 91.6%	-	Component_Type: "password" (100%) Component_Type: "password" (100%)	✓
9	"What is the weather like in Seattle today?" None	None 78.5%	-	-	✓
10	"Thanks for your help, have a great day!" None	None 93.3%	-	-	✓



Chapter 4: AI Agents

The AI Agents page allows administrators to configure agents, manage categories and view categorization reports.

It is the configuration and reporting layer that gives each agent its own identity, category structure and learning loop.

AI Agents

Under the AI Agents section, administrators can create, and configure the settings for AI agents.

Note: The number of configuration tabs will differ depending on the type of agent.

Configuration

ANDREA'S AGENT

ID #9

Configuration

Categories

Category Breakdown

Knowledge Gaps

Agent Name *

Andrea's Agent

Agent ID

9

Tags

Andrea

Agent Type *

Knowledge Agent

Answers customer questions using your knowledge base. Best for FAQ-style self-service where callers ask questions and the agent retrieves answers.

Voice Name Override

Leave blank to use tenant default

Enter an Azure TTS voice name to use for this bot. Leave blank to use the tenant default configured in AI Settings.

High Confidence Threshold (0 - 1)

0.75

Medium Confidence Threshold (0 - 1)

0.45

Minimum score for a verified answer to be returned directly without LLM assistance. Above this the Agent responds immediately.

Minimum score for a verified answer to be used with LLM assistance. Below this the Agent falls through to GAI.

Gap Analysis

Gap Lookback Window (days)

30

Minimum Interactions for Gap

10

How far back the context starts on each gap document look.

Minimum number of times a topic must appear before it is flagged as a knowledge gap.

Run gap analysis automatically

The configuration page is available for all types of agents and allows administrators to manage configuration settings including tags, agent type, and confidence threshold.

The following table describes the settings on this page:

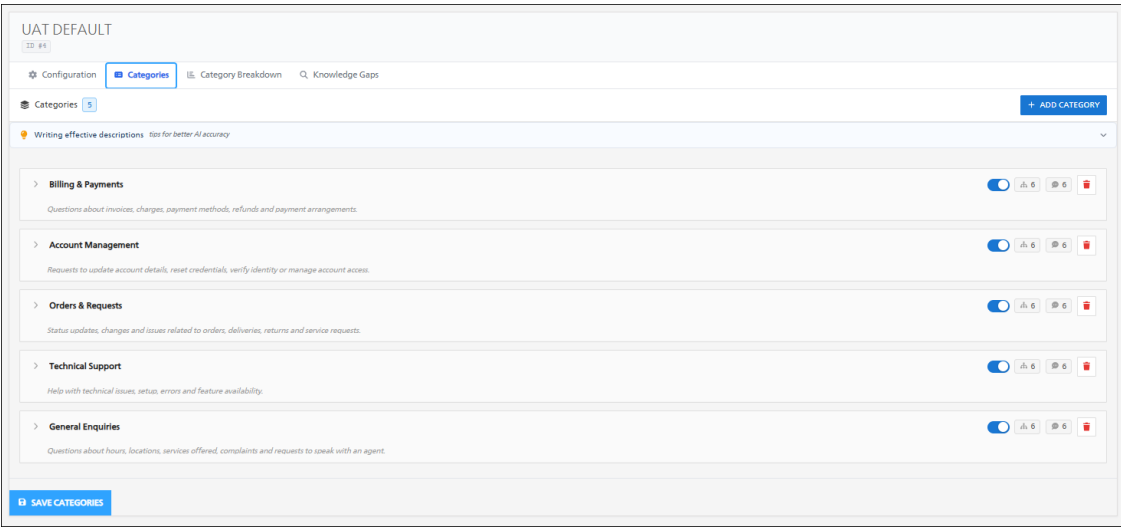
Setting	Description
Agent Name	A unique name assigned to the AI Agent.
Agent ID	A unique ID assigned to the AI Agent.

Setting	Description
Tags	The tags in the dropdown list correspond to the tags used in the knowledge base. The tags selected in this field will be used to generate AI answers using the corresponding sources.
Agent Type	AI Studio supports three types of agents: • Knowledge Agent: Answers customer questions using your knowledge base. Best for FAQ style self-service where callers ask questions and the agent retrieves answers. • Natural Language Routing: Routes contacts to the right queue or agent based on what the caller says. Best when you need intelligent intent-based routing without a full self-service flow. • Automation Agent: Handles end-to-end service requests with multi-step automation. Best for transactional flows like account lookups.
Voice Name Override	Enter an Azure TTS voice name to use for this bot. Leave this field blank to use the tenant default configured in AI Settings.
High Confidence Threshold	The minimum score for a verified answer to be returned directly without LLM assistance. When the confidence score is above the high confidence threshold, the agent responds immediately. Enter a value between 0 and 1.
Medium Confidence Threshold	The minimum score for a verified answer to be used with LLM assistance. Below this score, the Agent falls through to GAI. Enter a value between 0 and 1.
Gap Lookback Window	The number of days that the context stats look back on for each gap document.
Minimum Interactions for Gap	The minimum number of times a topic must appear before it is flagged as a knowledge gap.
Run gap analysis automatically	Toggle this setting on to run gap analysis automatically.
Run daily at (EDT)	Set the time that the gap analysis runs each day.
Custom Agent Prompt	Enter a custom agent prompt to override the default GAI prompt for this agent. Use {context} as a placeholder for the related knowledge found. If left blank, the default help-desk prompt is used.

Setting	Description
	This is an optional field. Note: Prompt length affects token usage.

Categories

The categories page is only available for knowledge agents. Categories are used primarily for reporting purposes, to allow administrators to view the number of interactions handled in each category and subcategory, and to view a more detailed breakdown for those interactions.



Categories belong to the agent, not to the knowledge base tag. A single agent can search multiple tags, and a single tag can be assigned to multiple AI agents. Two agents can share the same tag and have completely different category trees. A billing agent and a services agent using the same knowledge base interpret calls differently.

Permits & Licensing [Toggle] [3] [3] [Red X]

DESCRIPTION

Questions about building permits, business licenses, applications

EXAMPLE PHRASES [3]

how do I apply for a building permit × what is the status of my permit × how much does a business license cost ×

Add example phrase... [3]

SUBCATEGORIES [3]

NAME	DESCRIPTION
Building Permit Status	Inquiries about existing permit applications
Business License	New and renewal business license questions
Application Requirements	What documents and fees are needed
Subcategory name...	Description... [3] [ADD]

Bylaw FAQ [Toggle] [2] [4] [Red X]

General questions for bylaw

The following table describes the settings on this page:

Setting	Description
Description	A single, clear, objective sentence that explains what the category means in the context of your contact center. The clearer and more objective the description, the better the AI performs, especially on edge cases and ambiguous utterances. Focus on the caller's goal or problem, use simple unambiguous language and mention key signals that agent would look for. Avoid vague phrases, instructions to the bot, and simply repeating the category name. <u>Good Examples</u> Customer is asking about the status of an existing request or order. Caller wants to update their personal information or account details. Customer is reporting a problem with a product or service they received.



5. Add subcategories as needed.

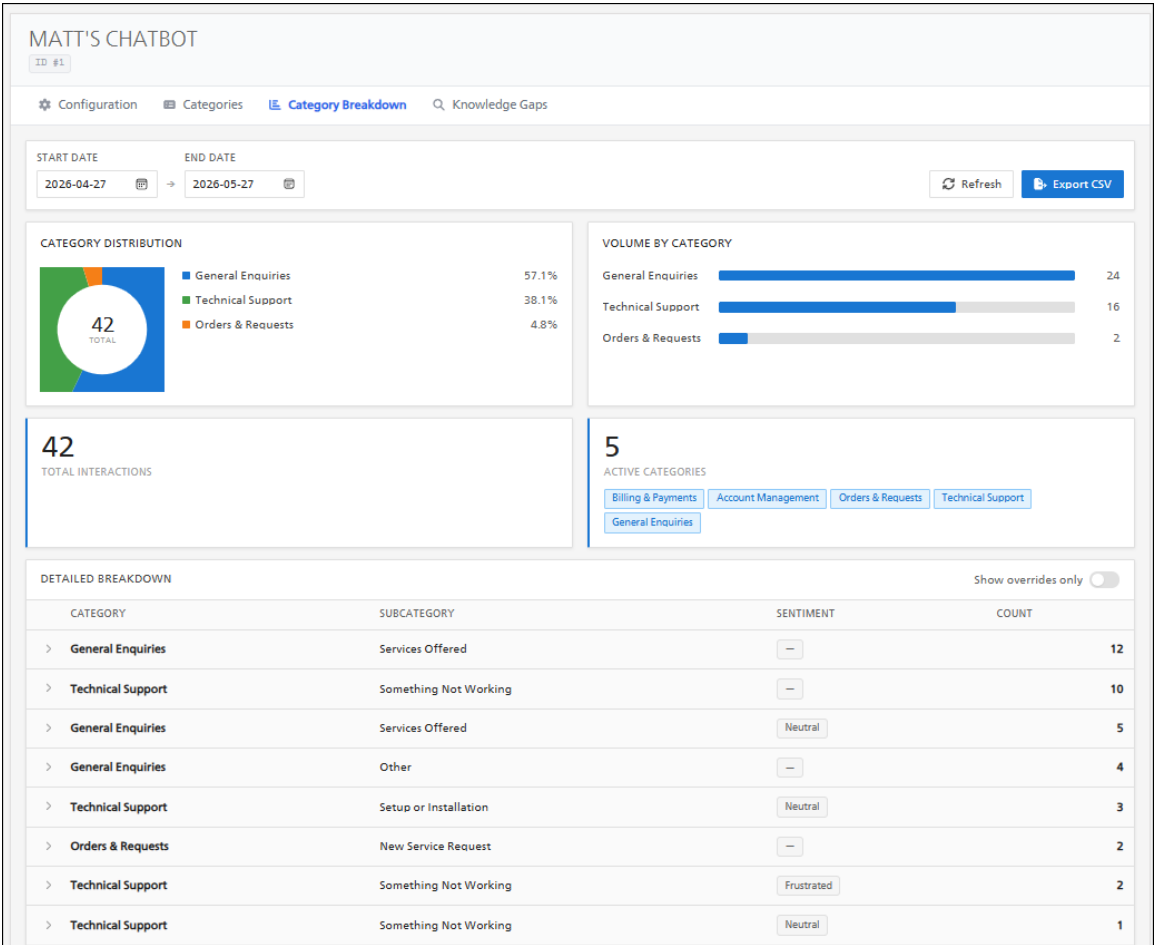


6. Click Save categories to save your changes.



Category Breakdown

The category breakdown provides detailed information on each category including interaction volume, recent utterances, override activity, and category-level trends.



A summary is displayed on the top of the page detailing the category distribution, volume by category, total interactions and active categories.

Below this is a detailed breakdown of each category.

General Enquiries	Services Offered	Neutral	5
what is the ice bot	General Enquiries	Services Offered	OVERRIDE
What ai solutions do you offer	General Enquiries	Services Offered	OVERRIDE
what is an icebot	General Enquiries	Services Offered	OVERRIDE
yes please	General Enquiries	Services Offered	OVERRIDE
what ai solutions do you offer	General Enquiries	Services Offered	OVERRIDE

When the bot assigns an incorrect category, a supervisor can override it directly in the detailed breakdown table. The override is saved and the corrected example is automatically added to the LLM categorization prompt as a few-shot example for future calls.

The top five most recent overrides per category are used as context for the LLM and the category for the next classification.

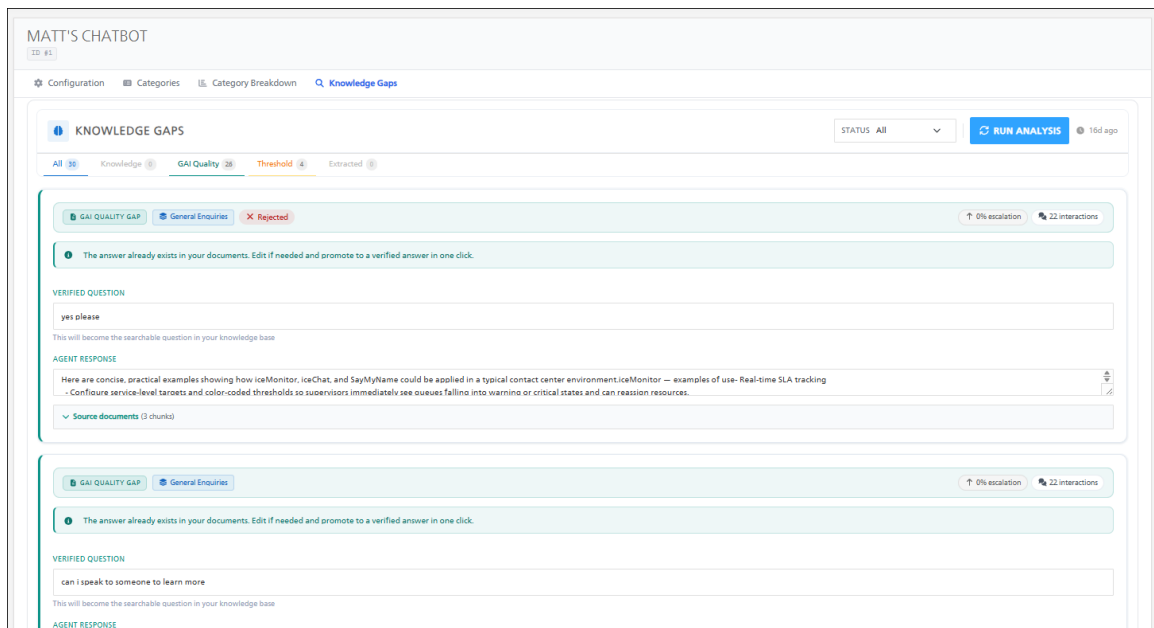
To correct an utterance classification, follow the steps below:

1. Click on an aggregated row to see individual utterances.
2. Select the correct category and subcategory from the drop down lists as needed.
3. Click Override to override the existing classification.

The “Show overrides only” toggle allows users to only see the utterances that have been overridden.

Knowledge Gaps

Gap detection analyzes caller utterances for categories where callers are escalating and classifies the root cause. Each root cause maps to a different resolution path because a missing answer, under-promoted content, a weak semantic match, and a question already solved by an agent are four different problems requiring four different solutions.



The Knowledge Gaps view is divided into the following sections by the type of gap:

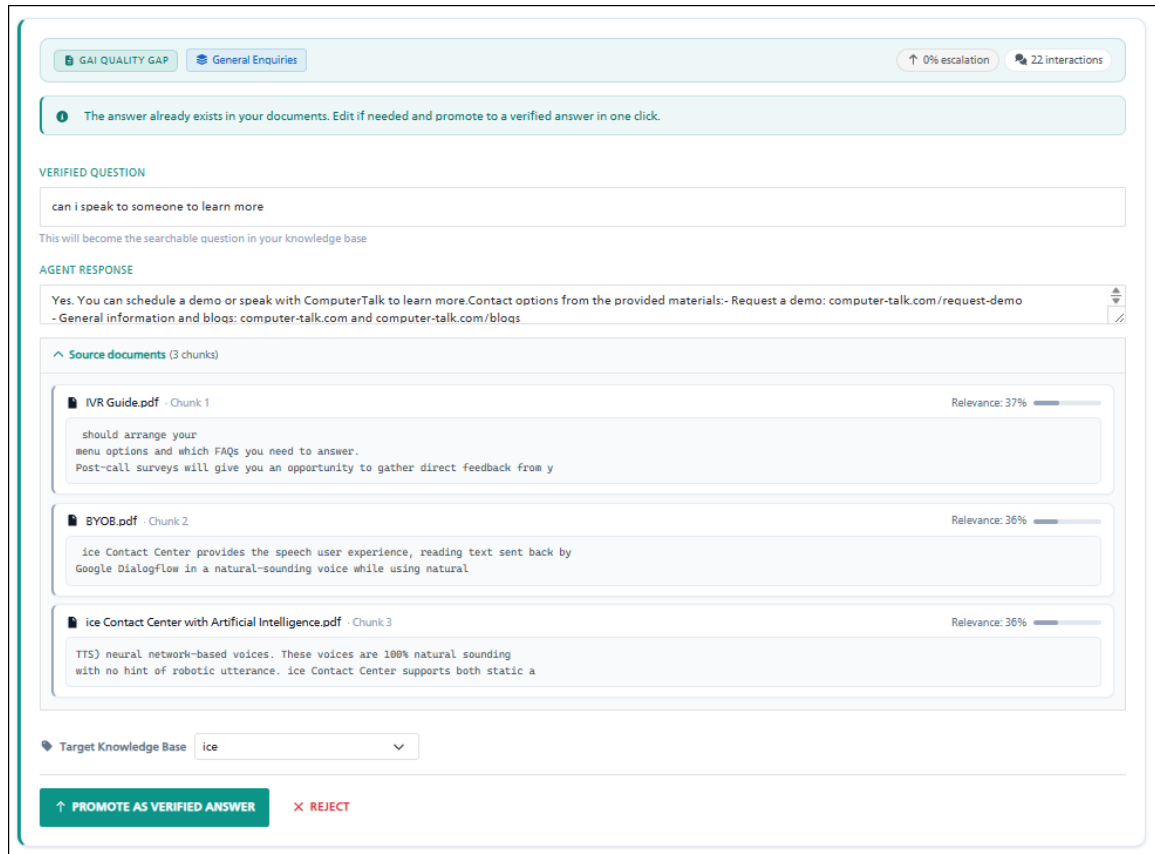
All: Displays all identified gaps.

Knowledge: Displays all identified knowledge gaps. A knowledge gap occurs when there is no matching verified answer, and GAI has not found any relevant context above the minimum confidence threshold. In these cases, the knowledge base does not cover this topic and the bot did not have any context to draw upon.

To resolve this gap, the large language model (LLM) synthesizes a canonical question and draft answer from escalated utterances. The supervisor can approve, edit or redraft the question-answer pair.

GAI Quality: Displays all identified GAI quality gaps. This occurs when there are no matching verified answers however GAI returned a response above the minimum confidence threshold. In these cases, the information exists in the document knowledge base, however it was not surfaced with enough confidence to promote it.

To resolve this gap, the supervisor can see the GAI response, the source chunks and the relevance scores. The supervisor can select the response to promote it to a verified answer, and edit the response if necessary.



Threshold: Displays all identified threshold gaps. This occurs when there is a verified answer match however the score falls between the medium and high confidence threshold. In these cases, the answer exists and is likely correct, however the semantic match is not strong enough to trigger the guaranteed verbatim return path. The LLM selection path is used instead.

To resolve these gaps, the supervisor can add additional question phrasings to the answer group or modify the answer text to improve future matches.

THRESHOLD GAP Technical Support 0% escalation 14 interactions

This answer matched callers but not at high confidence. Add their phrasings to improve future matches.

MATCHED VERIFIED ANSWER

Q. i need to change size of ice

A.

This answer matched but scored below your high confidence threshold on 1 interactions.

CALLER PHRASINGS TO ADD

These are the exact phrases callers used. Adding them strengthens the match for future interactions.

SELECT ALL | DESELECT ALL

☒ i need to change size of ice

+ ADD SELECTED PHRASINGS 1 PHRASING X DISMISS

Extracted: Displays all identified extracted gaps. This occurs when post-call transcript mining finds a question answered by an agent that has no matching verified answer above the 0.82 similarity threshold. In these cases, a human agent has already solved the problem in a call. The answer exists in contact transcripts but has not been captured in the knowledge base.

To resolve these gaps, the supervisor can promote the agent's answer directly to a verified answer with one click. The answer text is pre-filled from the transcript and may be edited before promotion.

The Run Analysis button allows administrators to re-run the analysis if changes have been made.

Configuration Categories Category Breakdown Knowledge Gaps

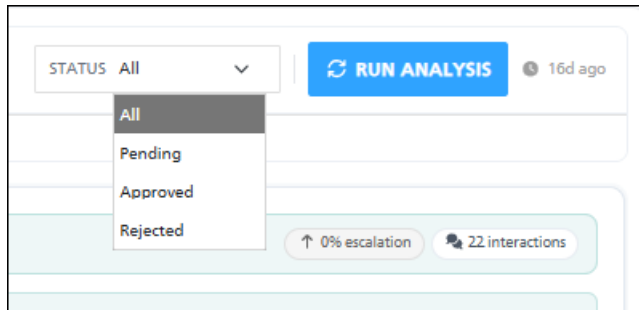
KNOWLEDGE GAPS STATUS: All RUN ANALYSIS 16d ago

All 30 Knowledge 0 GAI Quality 26 Threshold 4 Extracted 0

GAI QUALITY GAP General Enquiries Rejected 0% escalation 22 interactions

The answer already exists in your documents. Edit if needed and promote to a verified answer in one click.

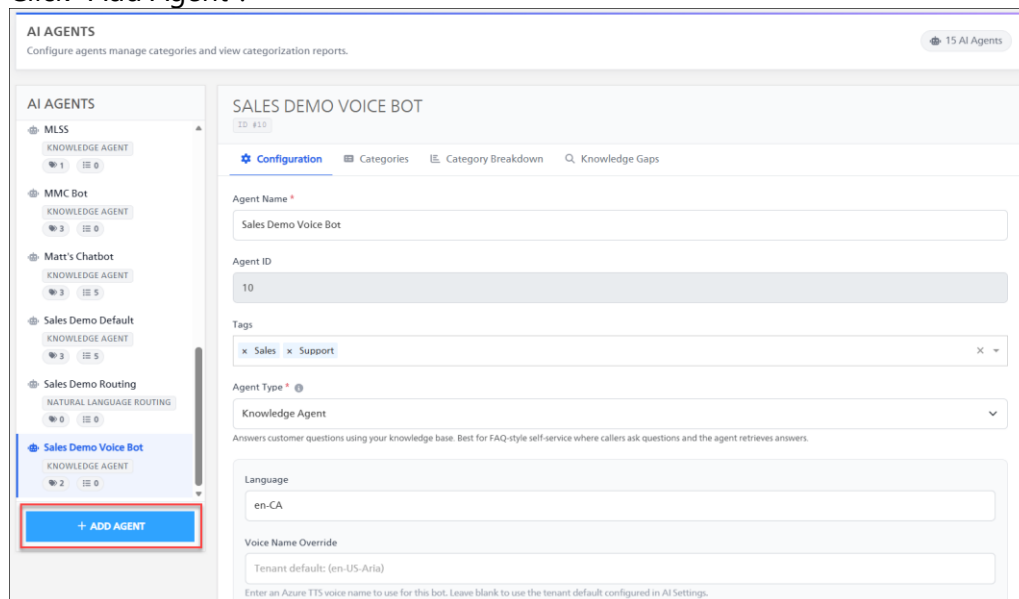
You may also filter the identified gaps by status using the status dropdown filter.



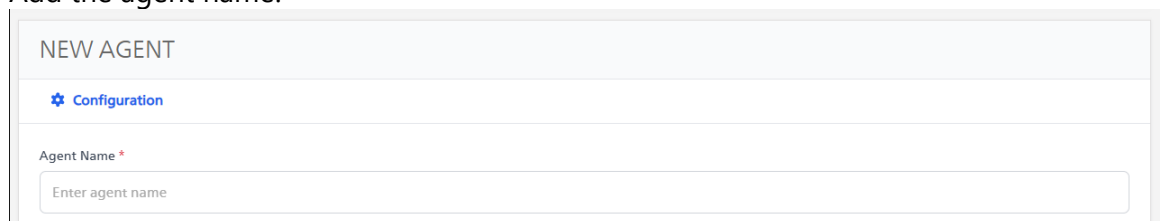
Creating an AI Agent

To create an AI Agent, follow the steps below:

1. Click "Add Agent".



2. Add the agent name.

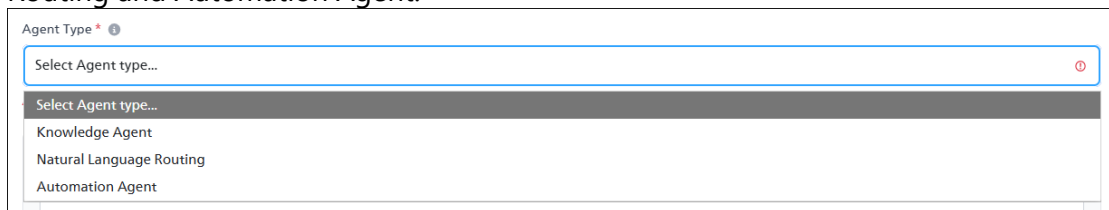


3. The Agent ID is automatically assigned. If you are creating a Knowledge Agent, enter the tags that the agent should reference from the knowledge management section.

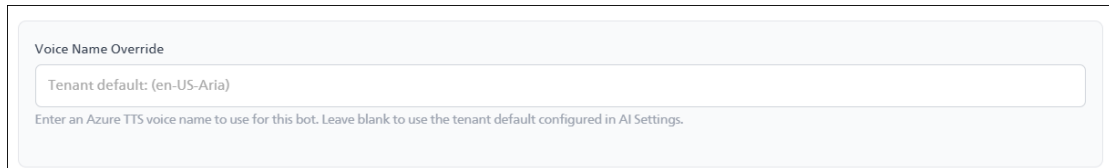
A screenshot of a form field labeled "Tags". It contains a text input with the placeholder "Select tags..." and a small downward arrow on the right side, indicating a dropdown menu.

The agent will only reference sources or verified answers assigned to tags selected in this field. Ensure that you have selected all necessary tags.

4. Select your agent type. Options include Knowledge Agent, Natural Language Routing and Automation Agent.

A screenshot of a form field labeled "Agent Type" with a red asterisk and an information icon. The dropdown menu is open, showing the following options: "Select Agent type..." (highlighted), "Knowledge Agent", "Natural Language Routing", and "Automation Agent".

5. If you wish to use a TTS voice that is different from the default voice, enter in a voice name override in this field.

A screenshot of a form field labeled "Voice Name Override". It contains a text input with the placeholder "Tenant default: (en-US-Aria)". Below the input, there is a small text note: "Enter an Azure TTS voice name to use for this bot. Leave blank to use the tenant default configured in AI Settings."

6. Click Save. If you have created a Knowledge Agent, you must also configure the Categories tab. For more information on categories and how to create a category, refer to Categories.



Chapter 5: Dashboards

The Dashboards tab in AI Studio serves as a cross-platform performance dashboard for managers and supervisors.

The dashboards module is a dedicated reporting page that provides two views depending on the number of agents selected: a cross-agent platform view for managers, and a single-agent view scoped to the specific agent type.

Platform View

When multiple agents are selected, the Dashboards page shows the cross-platform view with three tabs:

- Trends
- Categories and Sentiment
- Interactions

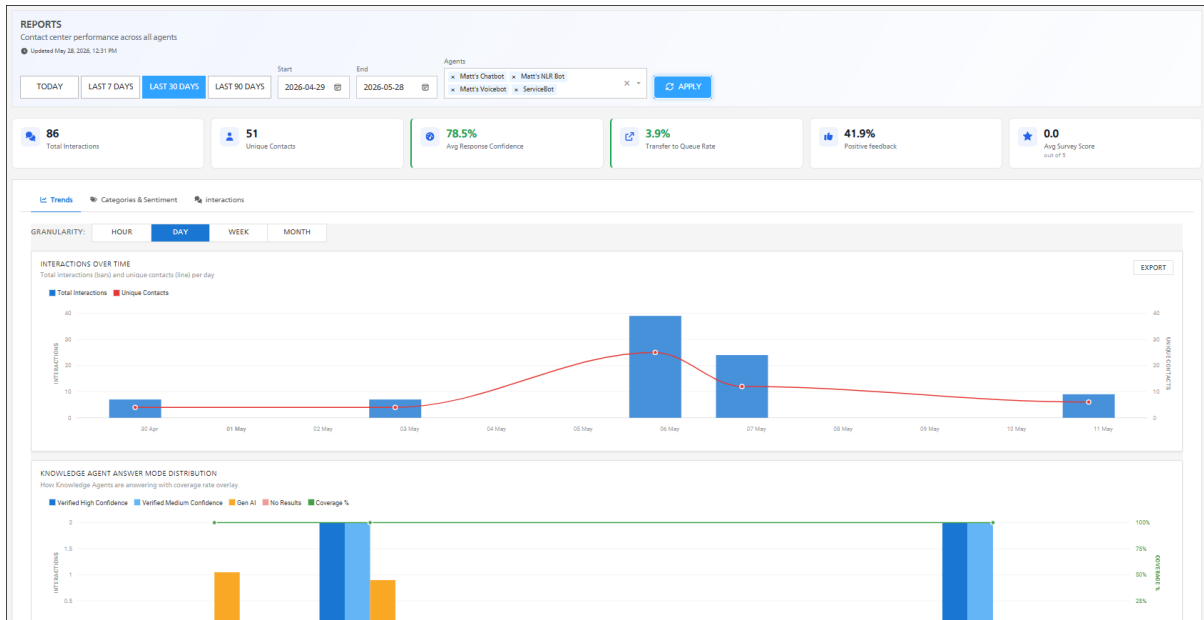
The top of the page includes a bot performance section composed of summary cards. The statistics displayed in these cards are governed by the time interval and agents selected in the filter options.

The screenshot shows the 'REPORTS' section of the AI Studio interface. It includes a sub-header 'Contact center performance across all agents' and a timestamp 'Updated Jun 2, 2026, 01:31 PM'. Below this are filter options for time intervals: 'TODAY', 'LAST 7 DAYS' (which is highlighted in blue), 'LAST 30 DAYS', and 'LAST 90 DAYS'. There are also date pickers for 'Start' (2026-05-27) and 'End' (2026-06-02). To the right, an 'Agents' dropdown menu is open, showing four selected agents: 'Matt's Chatbot', 'Matt's NLR Bot', 'Matt's Voicebot', and 'ServiceBot'. An 'APPLY' button is located at the bottom right of the filter section.

Select the interval for the displayed statistics from the options available. Administrators can select from: today, last 7 days, last 30 days, last 90 days, or a custom start and end date.

Additionally, administrators may select the agents included in the summary stats.

When complete, click "Apply" to refresh the view.



Bot Performance

The following table describes the performance stats:

Statistic	Description
Total Interactions	The total number of interactions handled across all selected Agents.
Unique Contacts	The total number of unique contacts handled across all selected Agents. Note: A contact may contain multiple interactions with a bot.
Avg Response Confidence	The average response confidence score across all selected Agents.
Transfer to Queue Rate	The percentage of contacts that are transferred to queue after an interaction with the AI Agent.
Positive Feedback	The percentage of interactions that received positive feedback. This value is only available for customers who have configured a post interaction feedback question. For example, after the bot provides a

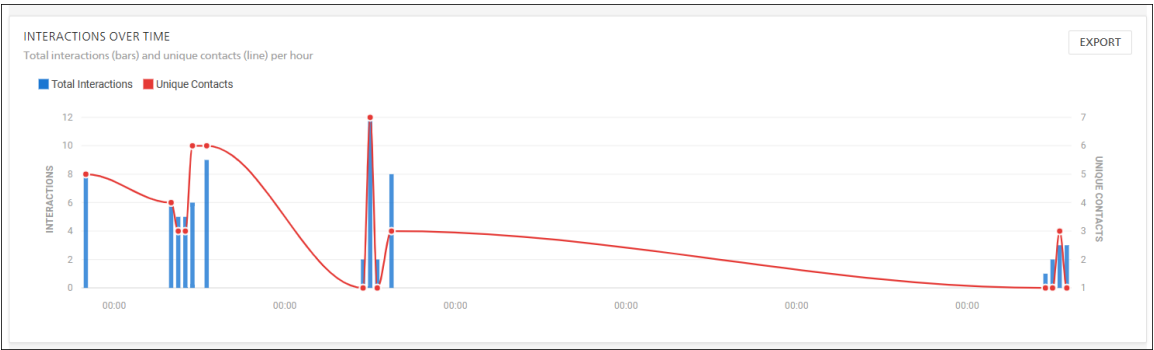
Statistic	Description
	response, the customer will be asked if the response was helpful. A response of yes would count towards the positive feedback metric.
Avg Survey Score out of 5	For customers with a post interaction survey configured, this value displays the average survey score on a scale from 1 to 5.

Trends

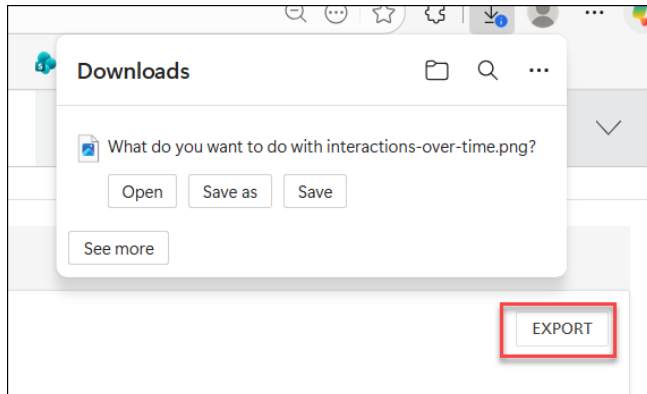
The trends tab displays four different graphs across four granularity options: hour, day, week and month. Click through the options to observe the different trends.

Interactions Over Time

This graph compares the total volume of interactions (bars) and unique contacts (line) for the selected interval.



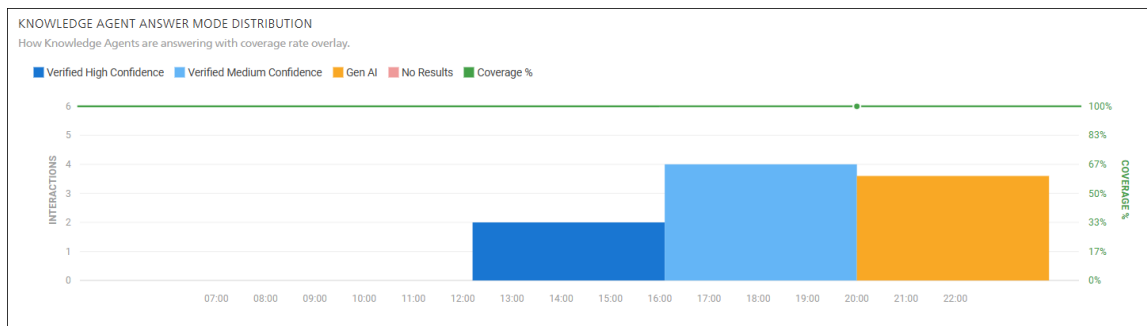
The export button allows you to export an image of the Interactions Over Time graph.



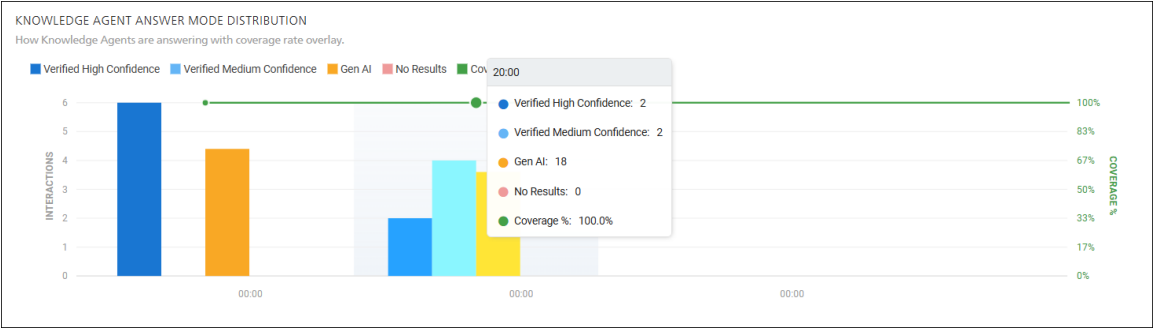
Knowledge Agent Answer Mode Distribution

The knowledge agent answer mode distribution graph shows how knowledge agents are answering with a coverage rate overlay. Coverage is defined as the percentage of interactions in which the bot was able to provide a response. This metric does not take accuracy of response into account.

The options include: verified high confidence, verified medium confidence, GenAI, and no results.

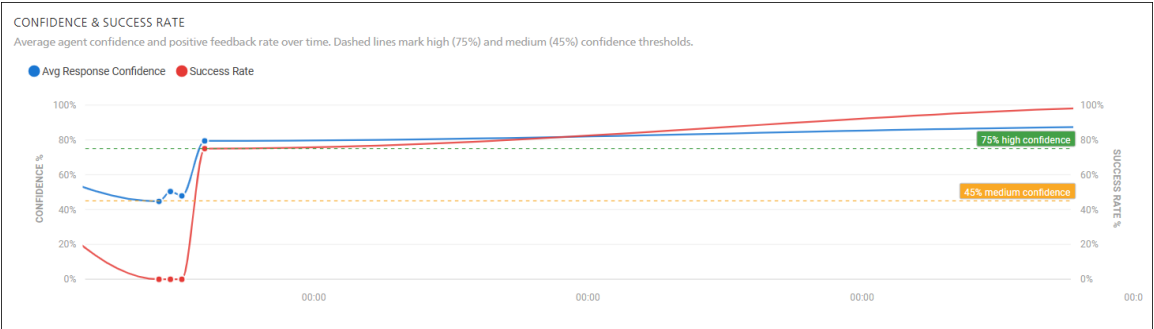


Hover over the graph to see a detailed breakdown of the number of interactions for each answer mode.

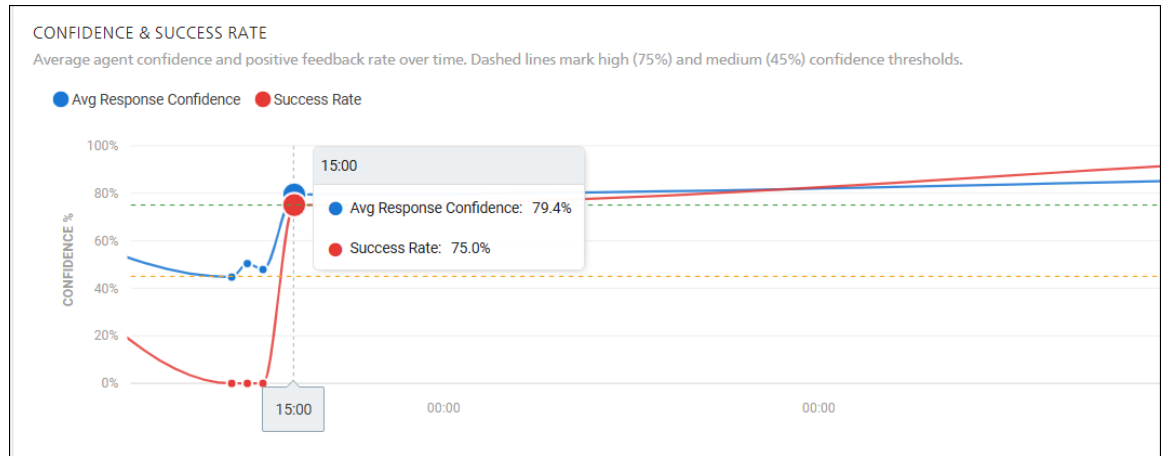


Confidence and Success Rate

The confidence and success rate graph shows the average agent confidence and positive feedback rate over time. Dashed lines mark high (75%) and medium (45%) confidence thresholds.

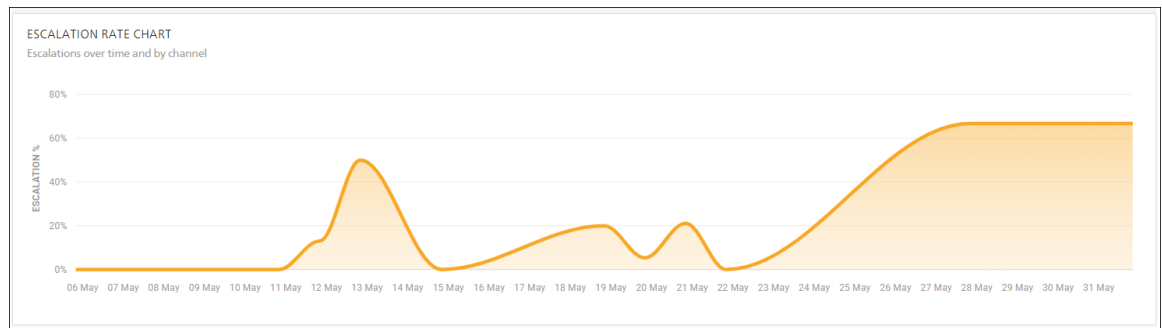


Hover over the graph to see the detailed statistic values.

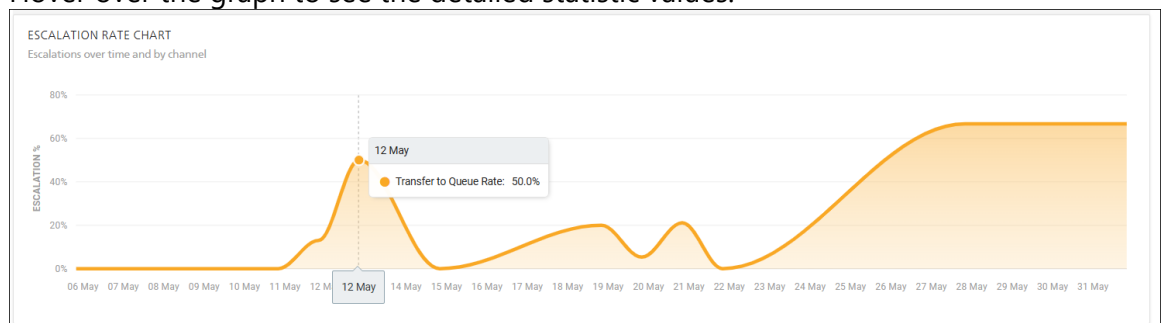


Escalation Rate Chart

The escalation rate chart shows the percentage of escalations over time and by channel.



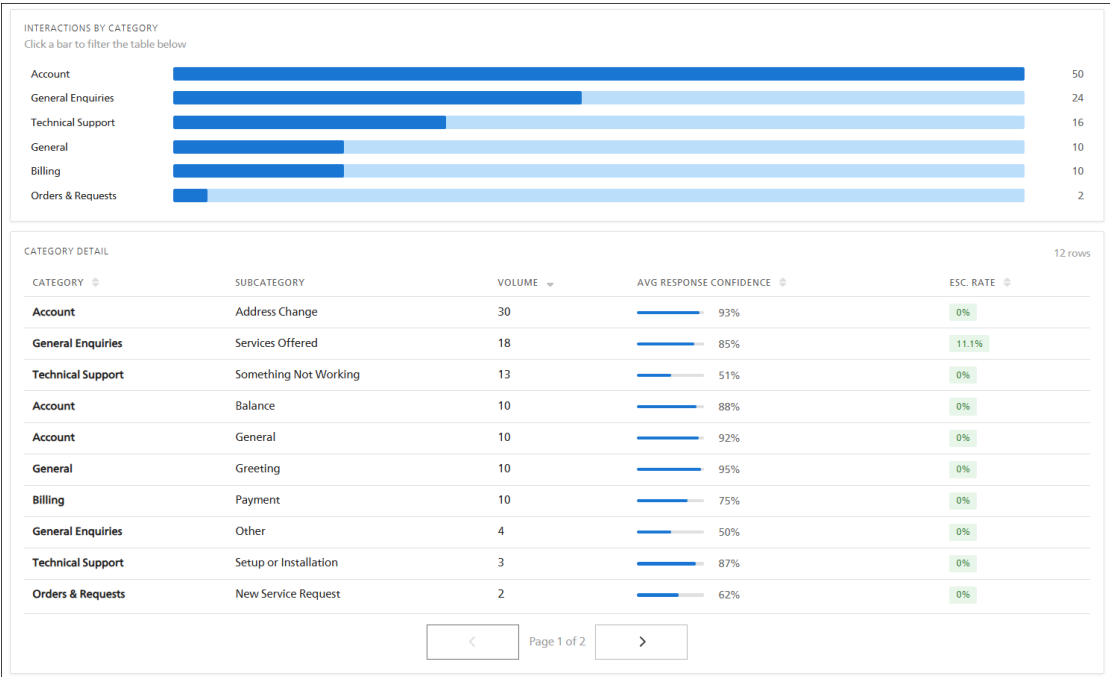
Hover over the graph to see the detailed statistic values.



Categories and Sentiment

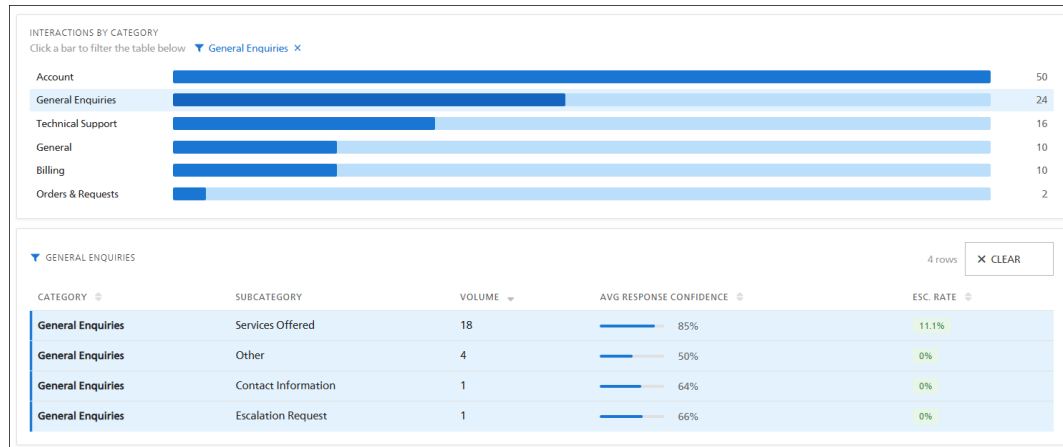
Category Breakdown

The category breakdown view is divided into two sections: interactions by category and category detail.



The Interactions by Category bar graph shows the number of interactions in each category.

Click on a bar to filter the detailed breakdown for a specific category.



The following table describes each column in the category detail.

Column	Description
Category	The name of the category selected.
Subcategory	The name of the subcategory.
Volume	The number of interactions in each subcategory.
Avg Response Confidence	The average response confidence score for the interactions in each subcategory.
Knowledge Gaps	The number of identified knowledge gaps.
Esc Rate	The escalation rate for the interactions in each subcategory.

Sentiment Breakdown

The sentiment breakdown shows the sentiment distribution for the interactions



Interactions

The interactions tab shows a list of all the interactions handled by the bot.

Click on an utterance to view the full conversation timeline.

The screenshot displays the 'Interactions' tab in AI Studio. On the left, a table lists 100 interactions with columns for Contact ID, Utterance, Category, Sentiment, Confidence, and Feedback. The first row shows Contact ID 206031 with the utterance 'I want to speak with an ag...'. On the right, a detailed view of interaction 206031 is shown, titled 'Sales Demo Routing'. It displays a timeline of events: 'Entrance' at 11:12 AM, followed by 'Utterance' and 'Escalated' events. A blue box highlights the utterance 'I want to demo the service agent' at 03:12 PM, which is routed to 'Sales Demo Routing' with 93% confidence. Another blue box highlights the utterance 'I want to speak with an agent' at 03:13 PM, which is routed to 'Agent' with 97% confidence.

CONTACT ID	UTTERANCE	CATEGORY	SENTIMENT	CONFIDENCE	FEEDBACK
206031	I want to speak with an ag...		—	97%	—
206031	I want to demo the service ...		—	93%	—
206021	I want to speak with an ag...		—	97%	—
205991	goodbye	General Enquiries Other	Neutral	33%	—
205991	how does iceBar work?	Technical Support Setup or Installation	Neutral	100%	—
205121	Facebook Messenger: 3800...	Account Management Password Reset	Neutral	51%	—
205111	sip_iceMessaging@icescap...	General Enquiries Escalation Request	Neutral	74%	—
205121	Speak to an agent		—	92%	—
205111	Speak to an agent		—	92%	—
205081	Suggested response	General Enquiries Other	Neutral	45%	—
204991	Show my account balance		—	87%	—
204991	How long does it take for a...	Transfers & Payments Payment Status	Neutral	98%	—
204991	SHow the online banking d...		—	94%	—
204981	{1:description}	Orders & Requests Request a Document	Neutral	45%	—

The following table describes each column.

Column	Description
Contact ID	The contact ID associated with the interaction. Note: Multiple interactions may be part of the same conversation, and therefore use the same contact ID.
Utterance	The captured customer utterance.
Category	Displays the category identified for the interaction. Note: This is only applicable for knowledge agents.
Sentiment	Displays the sentiment of the interaction. Options include: Positive

Column	Description
	- Neutral - Frustrated Note: This is only applicable for knowledge agents.
Confidence	Displays the confidence level for the interaction.
Feedback	Displays the feedback given by the customer if applicable. Note: If the feedback question has not been configured in the workflow, this column will remain empty.
Date	The date of the interaction.

Search and Filter Options

The search and filter options allow users to filter the interactions table as needed.

The screenshot shows a search and filter interface. At the top, there is a search bar with the placeholder text "Search utterances & responses...". To the right of the search bar are three dropdown menus labeled "Sentiment", "Category", and "Answer type". Below these elements is a horizontal slider bar labeled "Min confidence: 0%" on the left, with a blue dot indicating the current confidence level.

The following table describes the search and filter options.

Option	Tabs
Search Bar	Search for key phrases, or terms within the utterances and responses.
Sentiment	Filter the interactions by the sentiment identified. Options include: - Positive - Neutral - Frustrated
Category	Filter the interactions by category.
Answer Type	Filter the interactions by answer type. Options include: - Verified High Confidence - Verified Medium Confidence

Option	Tabs
	• genAI • No Results For more information on answer types, refer to Answer Quality Breakdown.
Minimum confidence	Filter by the minimum confidence level for the interaction.

Single Agent View

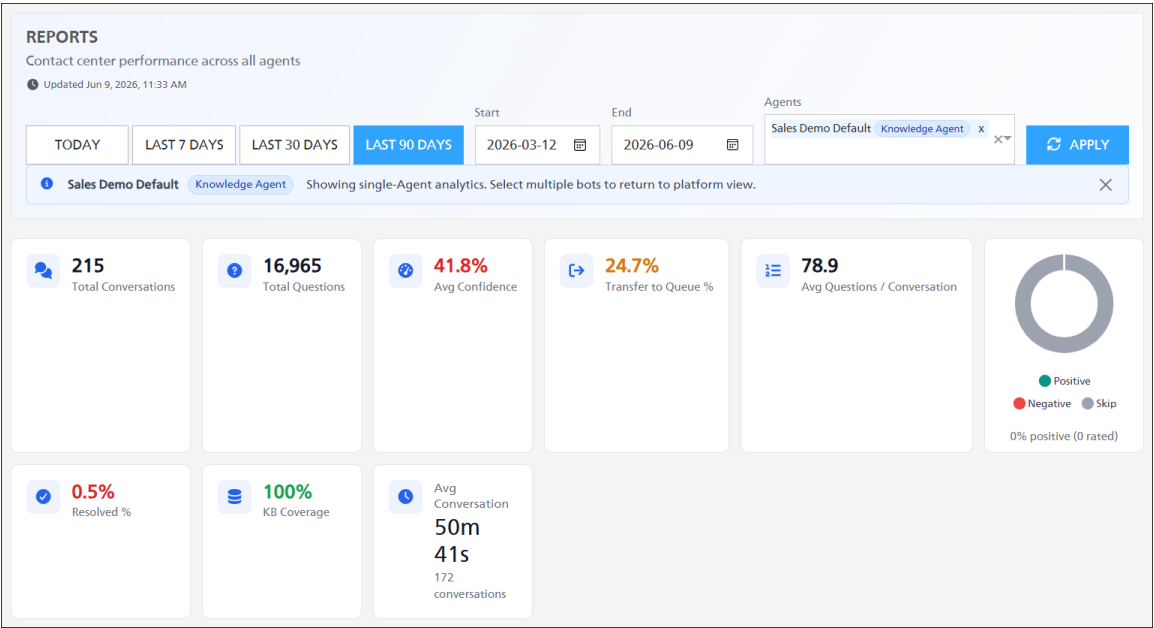
When exactly one agent is selected, the page transitions to an agent-type-specific view with tabs and KPI cards tailored to the agent’s behaviour pattern.

Agent Type	Tabs
FAQ	Overview Answer Quality Categories Interactions
Natural Language Routing	Overview Intent Breakdown Routing Interactions
Automation Agent	Overview Journey Outcomes Interactions

FAQ Agent Metrics

Overview

The overview provides performance statistics for the knowledge agent.



The following table describes each statistic.

Statistic	Description
Total Conversations	The total number of conversations handled by the bot in the selected interval.
Total Questions	The total number of questions posed to the bot in the selected interval.
Avg Confidence	The average confidence score of all interactions in the selected interval.
Transfer to Queue %	The percentage of contacts that are transferred to queue after an interaction with the AI Agent.

Statistic	Description
Avg Questions / Conversation	The average number of questions in each conversation.
Positive Negative Skip	The feedback breakdown. Displays the percentage of interactions that received positive feedback, negative feedback, and the percentage of interactions where the customer skipped the feedback question. Note: This will only generate data if the post-interaction feedback question has been configured in the workflow.
Resolved %	The percentage of interactions resolved by the bot.
KB Coverage	The knowledge base coverage percentage. Coverage is defined by the bot's ability to respond to a response. If a bot is able to provide a response to the interaction, it is covered.
Avg Conversation	The amount of time spent on average for a conversation.

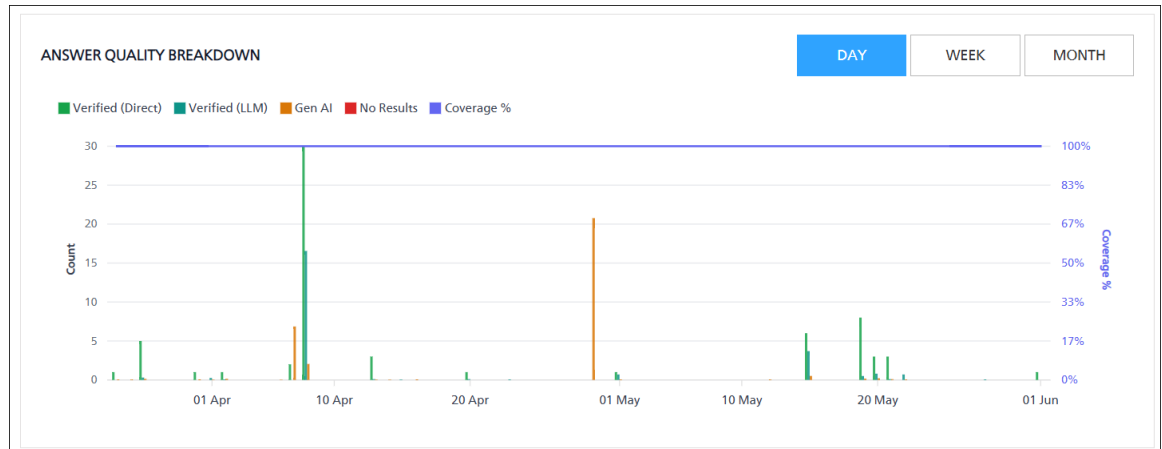
Answer Quality

Answer Quality Breakdown

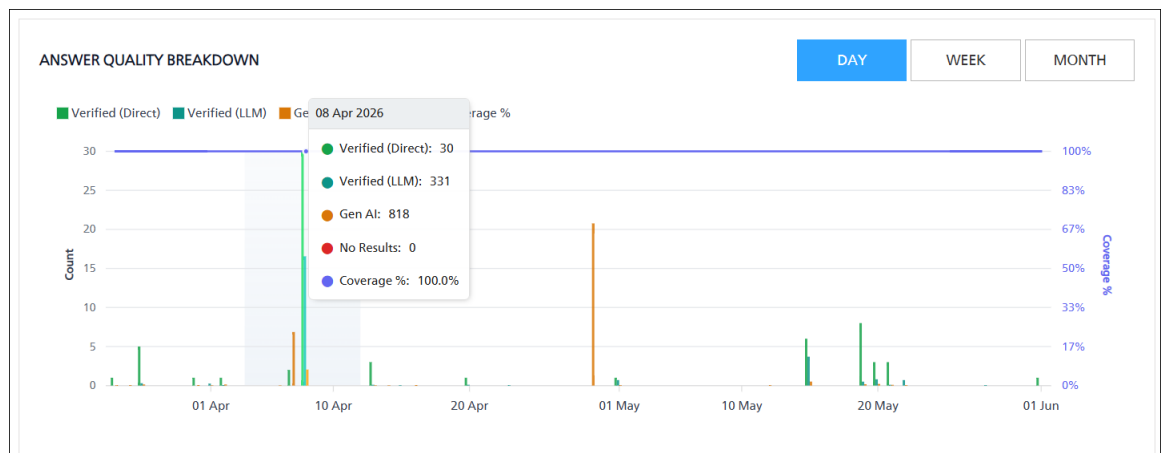
The Answer Quality Breakdown graph categorizes the response returned by the bot into the following categories:

- Verified (Direct) – Returns a verified answer that meets or exceeds the high confidence threshold without genAI.
- Verified (LLM) – Verified answer does not meet the high confidence threshold, but exceeds the medium confidence threshold. The bot evaluates the top responses and returns the best fit.
- genAI – Returns an AI generated response.
- No Results – No matching responses meeting the required confidence levels have been found. The bot will respond with a generic no results message.

The coverage % is overlaid on the graph to allow comparison between the answer types and the coverage provided by the bot.

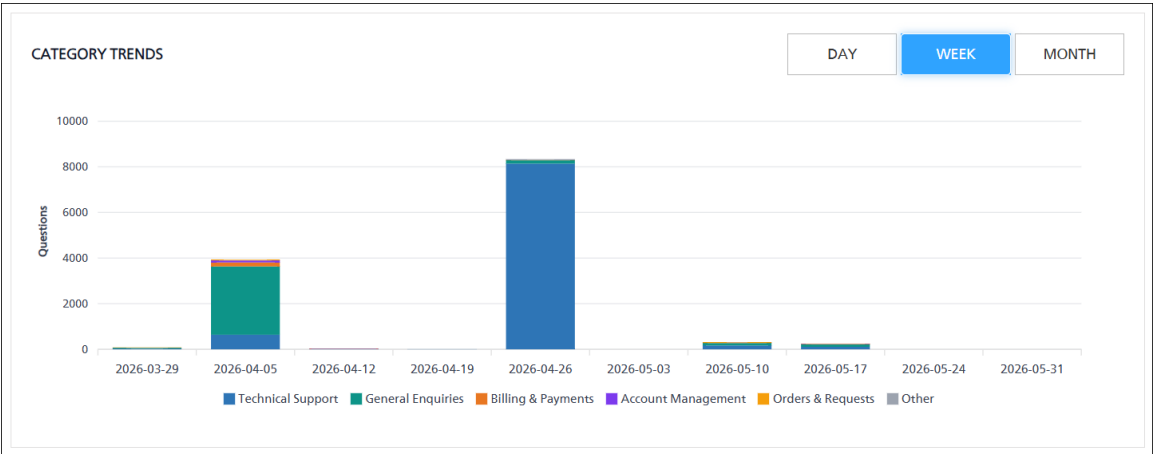


Hover over the graph to see the detailed breakdown for a specific day.

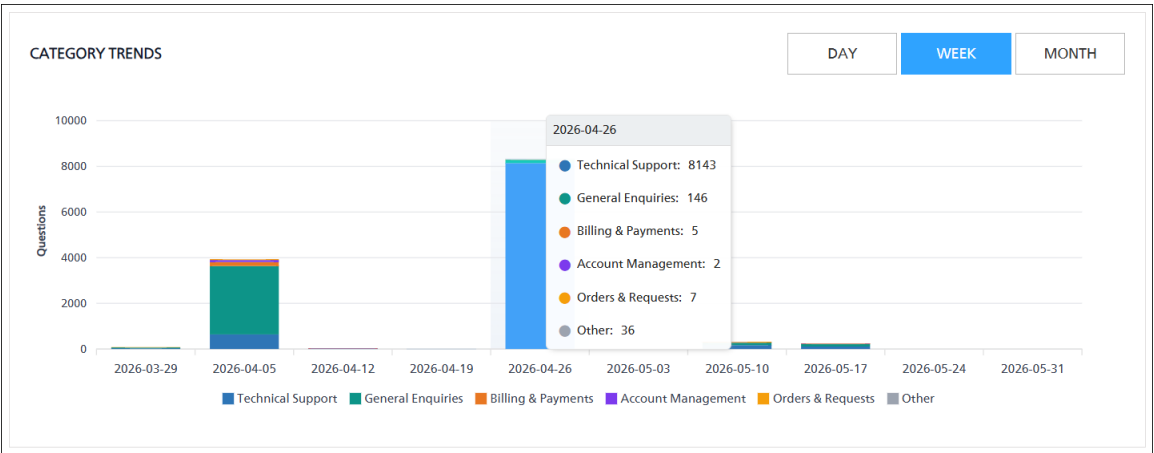


Category Trends

The category trends bar graph shows the distribution of interactions per category over the selected interval.

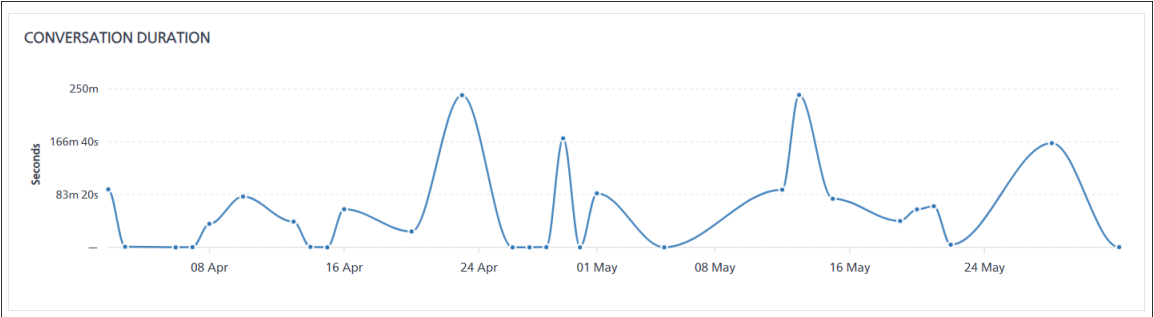


Hover over the graph to see the detailed breakdown for a specific day.

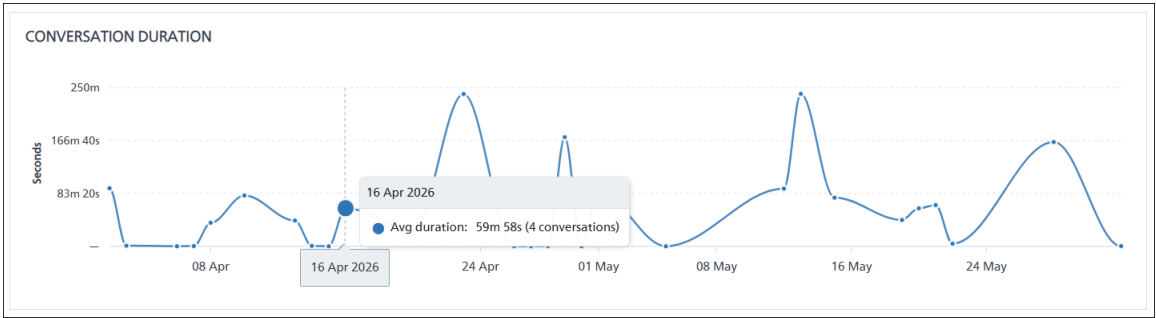


Conversation Duration

The conversation duration graph maps the average conversation duration for each day.



Hover over the graph to see the detailed breakdown for a specific day.



Categories

For more information on categories, refer to Category Breakdown.

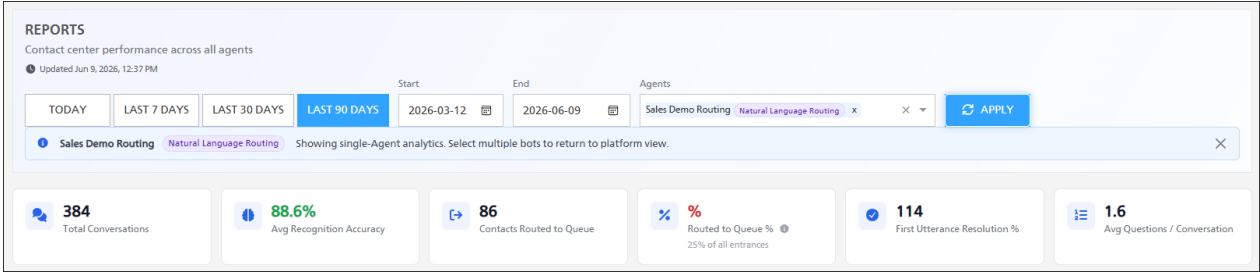
Interactions

For more information on interactions, refer to Interactions.

Natural Language Routing Metrics

Overview

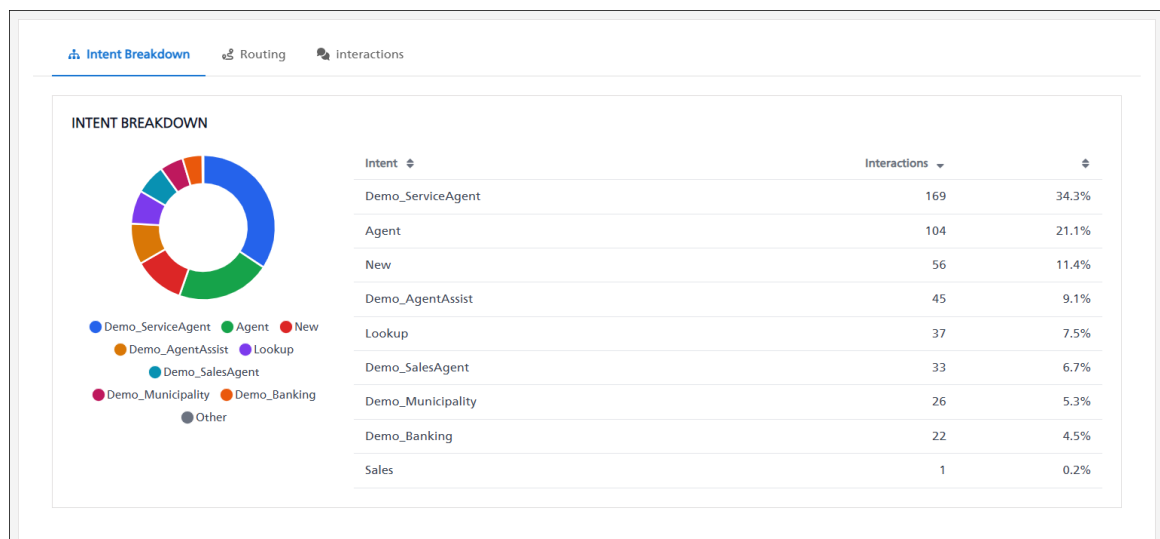
The overview provides performance statistics for the natural language routing agent.



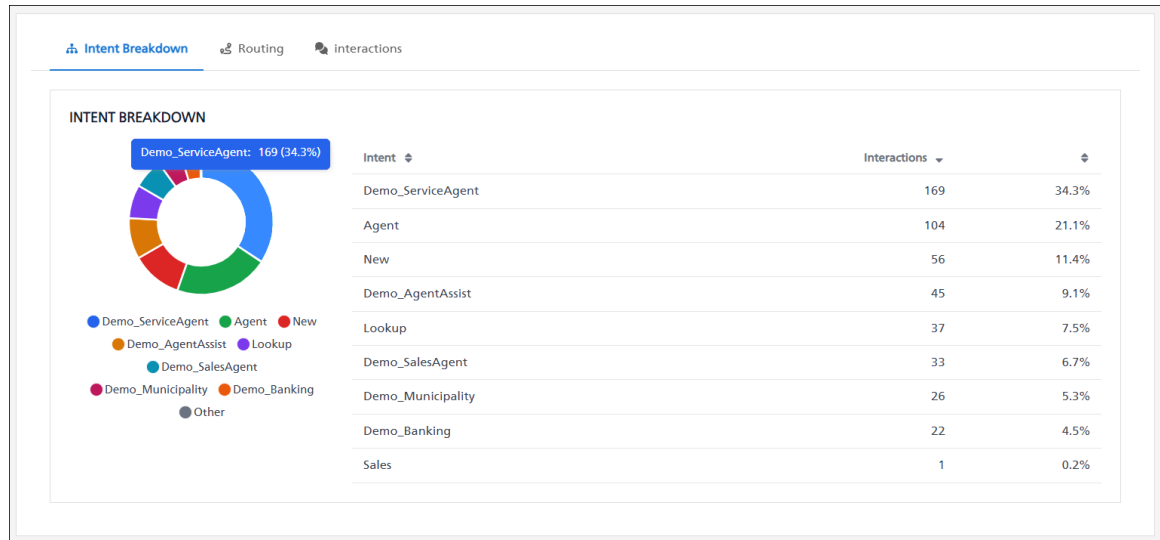
Statistic	Description
Total Conversations	The total number of conversations handled by the bot in the selected interval.
Avg Recognition Accuracy	Displays the average intent recognition accuracy as a percent.
Contacts Routed to Queue	Displays the number of contacts routed to queue after being handled by the bot, included those routed immediately.
Routed to Queue %	Displays the percentage of contacts routed to a queue.
First Utterance Resolution	The percentage of contacts that were resolved after the first utterance.
Avg Questions / Conversation	The average number of questions posed to the bot per conversation in the selected interval.

Intent Breakdown

The intent breakdown graph shows the distribution of interactions by intent, displaying both the number of interactions and the percentage.

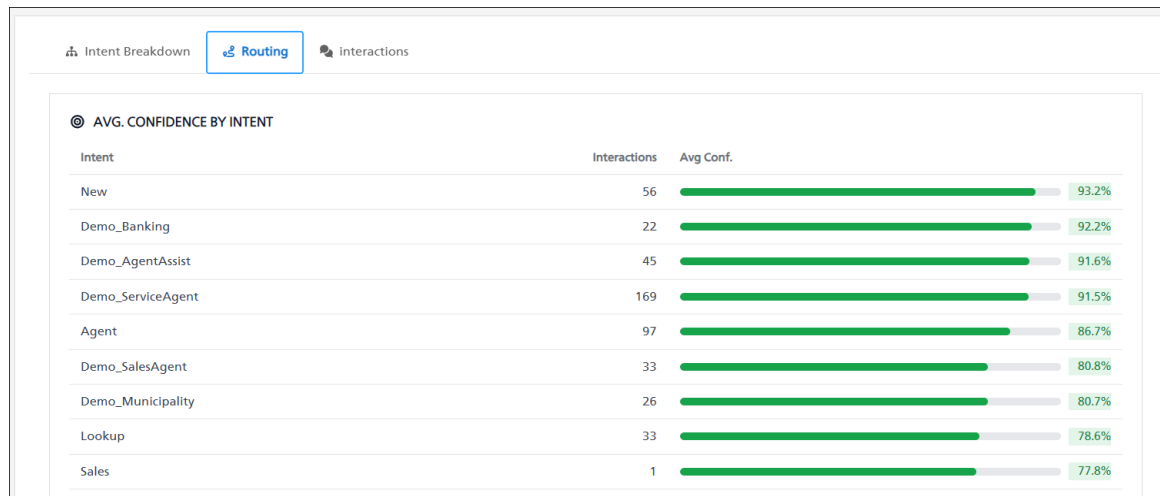


Hover over the graph to see the statistics for a specific intent.



Routing

The first table displays the average confidence percentage by intent. For each intent, the table displays the number of interactions identified, and the average confidence score for these interactions.



The second table displays the routing success rate divided by intent. For each intent, it shows the number of interactions, and the routing success rate. **Note:** Routing success is based on workflow configuration.

● ROUTING SUCCESS RATE BY INTENT

Intent	Interactions	Success Rate
New	56	100.0%
Lookup	33	100.0%
Demo_Municipality	26	100.0%
Sales	1	100.0%
Demo_SalesAgent	33	93.9%
Demo_Banking	22	90.9%
Demo_ServiceAgent	169	84.0%
Agent	97	56.7%
Demo_AgentAssist	45	8.9%

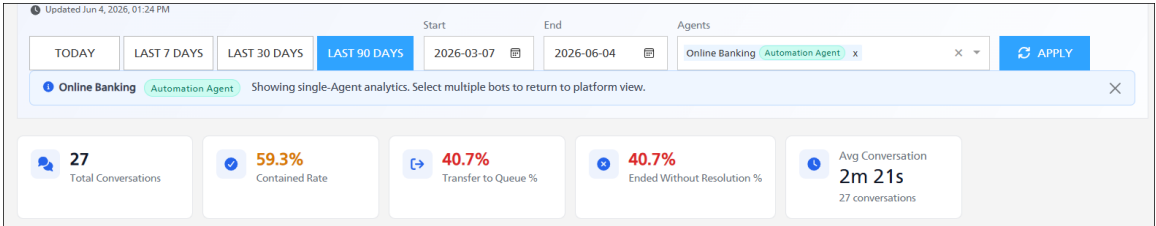
Interactions

For more information on interactions, refer to Interactions.

Automation Agent

Overview

The overview provides performance statistics for the automation agent.

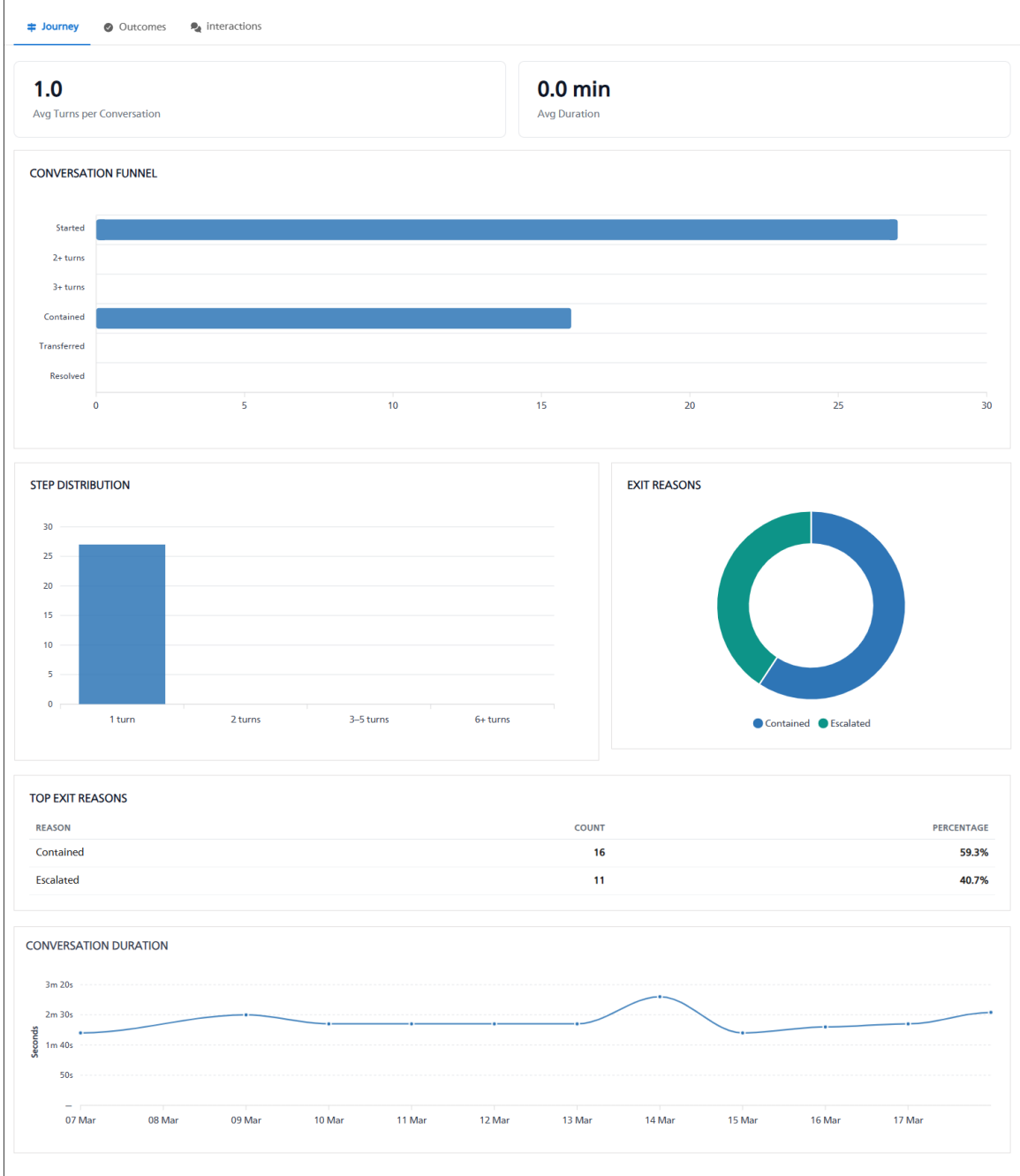


The following table describes each statistic.

Statistic	Description
Total Conversations	The total number of conversations handled by the bot during the selected interval.
Contained Rate	The number of contacts completely handled by the bot without escalations.
Transfer to Queue %	The percentage of contacts handled by the bot that were transferred to a queue.
Ended without resolution %	The percentage of contacts handled by the bot that ended without a resolution.
Avg Conversation	The average amount of time spent on a conversation.

Journey

The journey tab shows details about the number of turns, duration, step distribution and exit reasons per conversation.



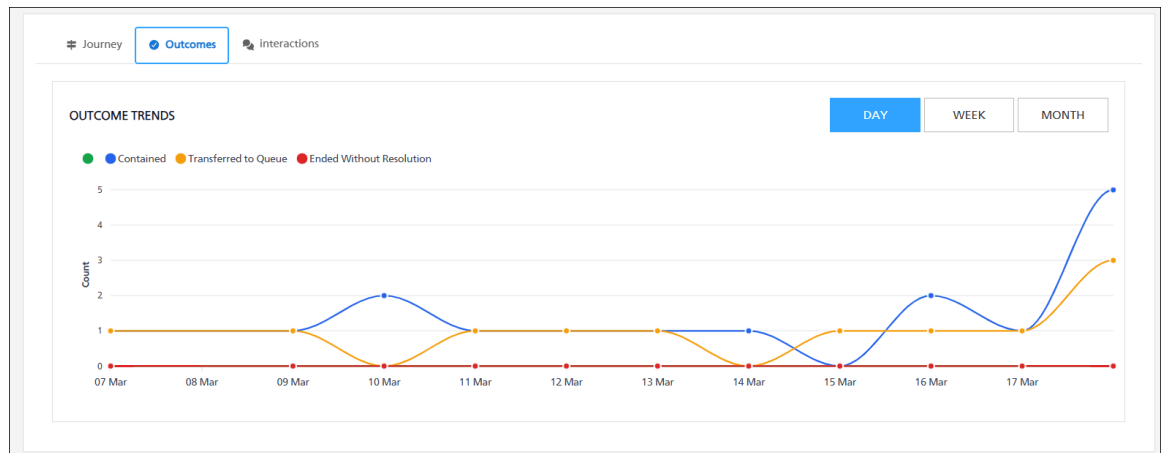
The following table describes each statistic.

Statistic	Description
Avg turns per conversation	The average number of turns per conversation. A turn refers to each time the customer speaks to the bot.
Avg duration	The average amount of time spent interacting with the bot.
Conversation funnel	This graph details the number of conversations that fit into the following categories: - Started: the number of conversations that have started with the agent. 2+ turns: the number of conversations that have 2 or more turns. 3+ turns: the number of conversations that have 3 or more turns. - Contained: the number of conversations that have been contained with the bot. These conversations have not required a transfer to queue. - Transferred: the number of conversations that have been transferred to a queue. - Resolved: the number of conversations that have been resolved by the bot.
Step distribution	This graph shows the number of conversations that have had turn, 2 turns, 3-5 turns, and 6 or more turns.
Exit reasons	This graph shows the distribution of contained vs escalated conversations handled by the bot.
Top exit reasons	This table shows the top exit reasons, the number of conversations for each, and the percentage.
Conversation duration	This graph shows the average conversation duration handled by the bot per day.

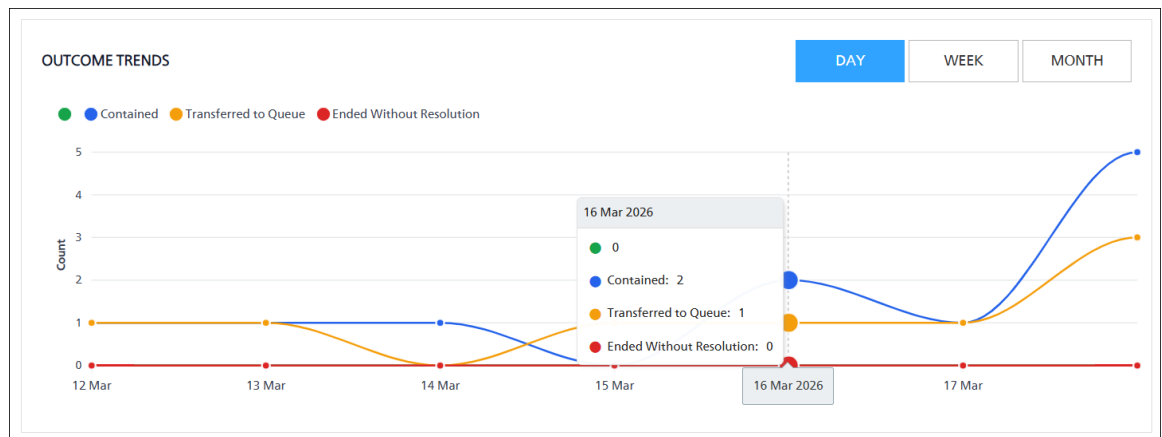
Outcomes

The outcomes graph shows the number of conversations for each outcome:

- Contained
- Transferred to Queue
- Ended Without Resolution



Hover over the graph to view the detailed breakdown for a specific day.



Interactions

For more information on interactions, refer to Interactions.