

Package ‘ellmer’

September 4, 2025

Title Chat with Large Language Models

Version 0.3.2

Description Chat with large language models from a range of providers including 'Claude' <<https://claude.ai>>, 'OpenAI' <<https://chatgpt.com>>, and more. Supports streaming, asynchronous calls, tool calling, and structured data extraction.

License MIT + file LICENSE

URL <https://ellmer.tidyverse.org>, <https://github.com/tidyverse/ellmer>

BugReports <https://github.com/tidyverse/ellmer/issues>

Depends R (>= 4.1)

Imports cli, coro (>= 1.1.0), glue, httr2 (>= 1.2.1), jsonlite, later (>= 1.4.0), lifecycle, promises (>= 1.3.1), R6, rlang (>= 1.1.0), S7 (>= 0.2.0)

Suggests connectcreds, curl (>= 6.0.1), gargle, gitcreds, jose, knitr, magick, openssl, paws.common, rmarkdown, shiny, shinychat (>= 0.2.0), testthat (>= 3.0.0), vcr (>= 2.0.0), withr

VignetteBuilder knitr

Config/Needs/website tidyverse/tidytemplate, rmarkdown

Config/testthat/edition 3

Config/testthat/parallel true

Config/testthat/start-first chat, provider*

Encoding UTF-8

RoxygenNote 7.3.2

Collate 'utils-S7.R' 'types.R' 'ellmer-package.R' 'tools-def.R' 'content.R' 'provider.R' 'as-json.R' 'batch-chat.R' 'chat-structured.R' 'turns.R' 'chat-tools.R' 'utils-coro.R' 'chat.R' 'content-image.R' 'content-pdf.R' 'content-replay.R' 'deprecated.R' 'httr2.R' 'import-standalone-obj-type.R' 'import-standalone-purrr.R' 'import-standalone-types-check.R' 'interpolate.R' 'live.R' 'parallel-chat.R' 'params.R'

'provider-anthropic.R' 'provider-any.R' 'provider-aws.R'
 'provider-openai.R' 'provider-azure.R' 'provider-google.R'
 'provider-cloudflare.R' 'provider-databricks.R'
 'provider-deepseek.R' 'provider-github.R'
 'provider-google-upload.R' 'provider-groq.R'
 'provider-huggingface.R' 'provider-mistral.R'
 'provider-ollama.R' 'provider-openrouter.R'
 'provider-perplexity.R' 'provider-portkey.R'
 'provider-snowflake-cortex.R' 'provider-snowflake.R'
 'provider-vllm.R' 'tokens.R' 'tools-def-auto.R'
 'utils-callbacks.R' 'utils-cat.R' 'utils-merge.R'
 'utils-prettytime.R' 'utils.R' 'zzz.R'

NeedsCompilation no

Author Hadley Wickham [aut, cre] (ORCID:

<<https://orcid.org/0000-0003-4757-117X>>),

Joe Cheng [aut],

Aaron Jacobs [aut],

Garrick Aden-Buie [aut] (ORCID:

<<https://orcid.org/0000-0002-7111-0077>>),

Barret Schloerke [aut] (ORCID: <<https://orcid.org/0000-0001-9986-114X>>),

Posit Software, PBC [cph, fnd] (ROR: <<https://ror.org/03wc8by49>>)

Maintainer Hadley Wickham <hadley@posit.co>

Repository CRAN

Date/Publication 2025-09-03 22:00:09 UTC

Contents

batch_chat	3
Chat	5
chat	11
chat_anthropic	12
chat_aws_bedrock	14
chat_azure_openai	15
chat_cloudflare	17
chat_cortex_analyst	19
chat_databricks	20
chat_deepseek	22
chat_github	23
chat_google_gemini	25
chat_groq	27
chat_huggingface	29
chat_mistral	31
chat_ollama	32
chat_openai	34
chat_openrouter	35
chat_perplexity	37

chat_portkey	38
chat_snowflake	40
chat_vllm	41
Content	43
contents_text	45
content_image_url	46
content_pdf_file	47
create_tool_def	48
google_upload	49
interpolate	50
live_console	51
parallel_chat	52
params	53
Provider	54
token_usage	55
tool	56
tool_annotations	58
tool_reject	60
Turn	62
Type	63
type_boolean	64

Index

67

batch_chat

Submit multiple chats in one batch

Description

[Experimental]

`batch_chat()` and `batch_chat_structured()` currently only work with `chat_openai()` and `chat_anthropic()`. They use the **OpenAI** and **Anthropic** batch APIs which allow you to submit multiple requests simultaneously. The results can take up to 24 hours to complete, but in return you pay 50% less than usual (but note that `ellmer` doesn't include this discount in its pricing metadata). If you want to get results back more quickly, or you're working with a different provider, you may want to use `parallel_chat()` instead.

Since batched requests can take a long time to complete, `batch_chat()` requires a file path that is used to store information about the batch so you never lose any work. You can either set `wait = FALSE` or simply interrupt the waiting process, then later, either call `batch_chat()` to resume where you left off or call `batch_chat_completed()` to see if the results are ready to retrieve. `batch_chat()` will store the chat responses in this file, so you can either keep it around to cache the results, or delete it to free up disk space.

This API is marked as experimental since I don't yet know how to handle errors in the most helpful way. Fortunately they don't seem to be common, but if you have ideas, please let me know!

Usage

```
batch_chat(chat, prompts, path, wait = TRUE)

batch_chat_text(chat, prompts, path, wait = TRUE)

batch_chat_structured(
  chat,
  prompts,
  path,
  type,
  wait = TRUE,
  convert = TRUE,
  include_tokens = FALSE,
  include_cost = FALSE
)

batch_chat_completed(chat, prompts, path)
```

Arguments

chat	A base chat object.
prompts	A vector created by <code>interpolate()</code> or a list of character vectors.
path	Path to file (with <code>.json</code> extension) to store state. The file records a hash of the provider, the prompts, and the existing chat turns. If you attempt to reuse the same file with any of these being different, you'll get an error.
wait	If TRUE, will wait for batch to complete. If FALSE, it will return NULL if the batch is not complete, and you can retrieve the results later by re-running <code>batch_chat()</code> when <code>batch_chat_completed()</code> is TRUE.
type	A type specification for the extracted data. Should be created with a <code>type_()</code> function.
convert	If TRUE, automatically convert from JSON lists to R data types using the schema. This typically works best when type is <code>type_object()</code> as this will give you a data frame with one column for each property. If FALSE, returns a list.
include_tokens	If TRUE, and the result is a data frame, will add <code>input_tokens</code> and <code>output_tokens</code> columns giving the total input and output tokens for each prompt.
include_cost	If TRUE, and the result is a data frame, will add <code>cost</code> column giving the cost of each prompt.

Value

For `batch_chat()`, a list of `Chat` objects, one for each prompt. For `batch_chat_test()`, a character vector of text responses. For `batch_chat_structured()`, a single structured data object with one element for each prompt. Typically, when type is an object, this will be a data frame with one row for each prompt, and one column for each property.

For any of the above, will return NULL if `wait = FALSE` and the job is not complete.

Examples

```
chat <- chat_openai(model = "gpt-4.1-nano")

# Chat -----

prompts <- interpolate("What do people from {{state.name}} bring to a potluck dinner?")
## Not run:
chats <- batch_chat(chat, prompts, path = "potluck.json")
chats

## End(Not run)

# Structured data -----
prompts <- list(
  "I go by Alex. 42 years on this planet and counting.",
  "Pleased to meet you! I'm Jamal, age 27.",
  "They call me Li Wei. Nineteen years young.",
  "Fatima here. Just celebrated my 35th birthday last week.",
  "The name's Robert - 51 years old and proud of it.",
  "Kwame here - just hit the big 5-0 this year."
)
type_person <- type_object(name = type_string(), age = type_number())
## Not run:
data <- batch_chat_structured(
  chat = chat,
  prompts = prompts,
  path = "people-data.json",
  type = type_person
)
data

## End(Not run)
```

Chat

The Chat object

Description

A Chat is a sequence of user and assistant [Turns](#) sent to a specific [Provider](#). A Chat is a mutable R6 object that takes care of managing the state associated with the chat; i.e. it records the messages that you send to the server, and the messages that you receive back. If you register a tool (i.e. an R function that the assistant can call on your behalf), it also takes care of the tool loop.

You should generally not create this object yourself, but instead call `chat_openai()` or friends instead.

Value

A Chat object

Methods

Public methods:

- `Chat$new()`
- `Chat$get_turns()`
- `Chat$set_turns()`
- `Chat$add_turn()`
- `Chat$get_system_prompt()`
- `Chat$get_model()`
- `Chat$set_system_prompt()`
- `Chat$get_tokens()`
- `Chat$get_cost()`
- `Chat$last_turn()`
- `Chat$chat()`
- `Chat$chat_structured()`
- `Chat$chat_structured_async()`
- `Chat$chat_async()`
- `Chat$stream()`
- `Chat$stream_async()`
- `Chat$register_tool()`
- `Chat$register_tools()`
- `Chat$get_provider()`
- `Chat$get_tools()`
- `Chat$set_tools()`
- `Chat$on_tool_request()`
- `Chat$on_tool_result()`
- `Chat$extract_data()`
- `Chat$extract_data_async()`
- `Chat$clone()`

Method `new()`:

Usage:

```
Chat$new(provider, system_prompt = NULL, echo = "none")
```

Arguments:

`provider` A provider object.

`system_prompt` System prompt to start the conversation with.

`echo` One of the following options:

- `none`: don't emit any output (default when running in a function).
- `output`: echo text and tool-calling output as it streams in (default when running at the console).
- `all`: echo all input and output.

Note this only affects the `chat()` method. You can override the default by setting the `ellmer_echo` option.

Method `get_turns()`: Retrieve the turns that have been sent and received so far (optionally starting with the system prompt, if any).

Usage:

```
Chat$get_turns(include_system_prompt = FALSE)
```

Arguments:

`include_system_prompt` Whether to include the system prompt in the turns (if any exists).

Method `set_turns()`: Replace existing turns with a new list.

Usage:

```
Chat$set_turns(value)
```

Arguments:

`value` A list of [Turns](#).

Method `add_turn()`: Add a pair of turns to the chat.

Usage:

```
Chat$add_turn(user, system)
```

Arguments:

`user` The user [Turn](#).

`system` The system [Turn](#).

Method `get_system_prompt()`: If set, the system prompt, if not, NULL.

Usage:

```
Chat$get_system_prompt()
```

Method `get_model()`: Retrieve the model name

Usage:

```
Chat$get_model()
```

Method `set_system_prompt()`: Update the system prompt

Usage:

```
Chat$set_system_prompt(value)
```

Arguments:

`value` A character vector giving the new system prompt

Method `get_tokens()`: A data frame with a `tokens` column that provides the number of input tokens used by user turns and the number of output tokens used by assistant turns.

Usage:

```
Chat$get_tokens(include_system_prompt = FALSE)
```

Arguments:

`include_system_prompt` Whether to include the system prompt in the turns (if any exists).

Method `get_cost()`: The cost of this chat

Usage:

```
Chat$get_cost(include = c("all", "last"))
```

Arguments:

`include` The default, "all", gives the total cumulative cost of this chat. Alternatively, use "last" to get the cost of just the most recent turn.

Method `last_turn()`: The last turn returned by the assistant.

Usage:

```
Chat$last_turn(role = c("assistant", "user", "system"))
```

Arguments:

`role` Optionally, specify a role to find the last turn with for the role.

Returns: Either a Turn or NULL, if no turns with the specified role have occurred.

Method `chat()`: Submit input to the chatbot, and return the response as a simple string (probably Markdown).

Usage:

```
Chat$chat(..., echo = NULL)
```

Arguments:

`...` The input to send to the chatbot. Can be strings or images (see [content_image_file\(\)](#) and [content_image_url\(\)](#)).

`echo` Whether to emit the response to stdout as it is received. If NULL, then the value of `echo` set when the chat object was created will be used.

Method `chat_structured()`: Extract structured data

Usage:

```
Chat$chat_structured(..., type, echo = "none", convert = TRUE)
```

Arguments:

`...` The input to send to the chatbot. This is typically the text you want to extract data from, but it can be omitted if the data is obvious from the existing conversation.

`type` A type specification for the extracted data. Should be created with a [type_\(\)](#) function.

`echo` Whether to emit the response to stdout as it is received. Set to "text" to stream JSON data as it's generated (not supported by all providers).

`convert` Automatically convert from JSON lists to R data types using the schema. For example, this will turn arrays of objects into data frames and arrays of strings into a character vector.

Method `chat_structured_async()`: Extract structured data, asynchronously. Returns a promise that resolves to an object matching the type specification.

Usage:

```
Chat$chat_structured_async(..., type, echo = "none", convert = TRUE)
```

Arguments:

`...` The input to send to the chatbot. Will typically include the phrase "extract structured data".

`type` A type specification for the extracted data. Should be created with a [type_\(\)](#) function.

`echo` Whether to emit the response to stdout as it is received. Set to "text" to stream JSON data as it's generated (not supported by all providers).

`convert` Automatically convert from JSON lists to R data types using the schema. For example, this will turn arrays of objects into data frames and arrays of strings into a character vector.

Method `chat_async()`: Submit input to the chatbot, and receive a promise that resolves with the response all at once. Returns a promise that resolves to a string (probably Markdown).

Usage:

```
Chat$chat_async(..., tool_mode = c("concurrent", "sequential"))
```

Arguments:

... The input to send to the chatbot. Can be strings or images.

`tool_mode` Whether tools should be invoked one-at-a-time ("sequential") or concurrently ("concurrent"). Sequential mode is best for interactive applications, especially when a tool may involve an interactive user interface. Concurrent mode is the default and is best suited for automated scripts or non-interactive applications.

Method `stream()`: Submit input to the chatbot, returning streaming results. Returns A **coro generator** that yields strings. While iterating, the generator will block while waiting for more content from the chatbot.

Usage:

```
Chat$stream(..., stream = c("text", "content"))
```

Arguments:

... The input to send to the chatbot. Can be strings or images.

`stream` Whether the stream should yield only "text" or ellmer's rich content types. When `stream = "content"`, `stream()` yields **Content** objects.

Method `stream_async()`: Submit input to the chatbot, returning asynchronously streaming results. Returns a **coro async generator** that yields string promises.

Usage:

```
Chat$stream_async(
  ...,
  tool_mode = c("concurrent", "sequential"),
  stream = c("text", "content")
)
```

Arguments:

... The input to send to the chatbot. Can be strings or images.

`tool_mode` Whether tools should be invoked one-at-a-time ("sequential") or concurrently ("concurrent"). Sequential mode is best for interactive applications, especially when a tool may involve an interactive user interface. Concurrent mode is the default and is best suited for automated scripts or non-interactive applications.

`stream` Whether the stream should yield only "text" or ellmer's rich content types. When `stream = "content"`, `stream()` yields **Content** objects.

Method `register_tool()`: Register a tool (an R function) that the chatbot can use. Learn more in `vignette("tool-calling")`.

Usage:

```
Chat$register_tool(tool)
```

Arguments:

`tool` A tool definition created by `tool()`.

Method `register_tools()`: Register a list of tools. Learn more in `vignette("tool-calling")`.

Usage:

```
Chat$register_tools(tools)
```

Arguments:

`tools` A list of tool definitions created by `tool()`.

Method `get_provider()`: Get the underlying provider object. For expert use only.

Usage:

```
Chat$get_provider()
```

Method `get_tools()`: Retrieve the list of registered tools.

Usage:

```
Chat$get_tools()
```

Method `set_tools()`: Sets the available tools. For expert use only; most users should use `register_tool()`.

Usage:

```
Chat$set_tools(tools)
```

Arguments:

`tools` A list of tool definitions created with `tool()`.

Method `on_tool_request()`: Register a callback for a tool request event.

Usage:

```
Chat$on_tool_request(callback)
```

Arguments:

`callback` A function to be called when a tool request event occurs, which must have `request` as its only argument.

Returns: A function that can be called to remove the callback.

Method `on_tool_result()`: Register a callback for a tool result event.

Usage:

```
Chat$on_tool_result(callback)
```

Arguments:

`callback` A function to be called when a tool result event occurs, which must have `result` as its only argument.

Returns: A function that can be called to remove the callback.

Method `extract_data()`: **[Deprecated]** Deprecated in favour of `$chat_structured()`.

Usage:

```
Chat$extract_data(...)
```

Arguments:

... See `$chat_structured()`

Method `extract_data_async()`: **[Deprecated]**

Usage:

`Chat$extract_data_async(...)`

Arguments:

... See `$chat_structured_async()`

Method `clone()`: The objects of this class are cloneable with this method.

Usage:

`Chat$clone(deep = FALSE)`

Arguments:

`deep` Whether to make a deep clone.

Examples

```
chat <- chat_openai()
chat$chat("Tell me a funny joke")
```

chat

Chat with any provider

Description

This is a generic interface to all the other `chat_` functions that allow to you pick the provider and the model with a simple string.

Usage

```
chat(
  name,
  ...,
  system_prompt = NULL,
  params = NULL,
  echo = c("none", "output", "all")
)
```

Arguments

name	Provider (and optionally model) name in the form "provider/model" or "provider" (which will use the default model for that provider).
...	Arguments passed to the provider function.
system_prompt	A system prompt to set the behavior of the assistant.
params	Common model parameters, usually created by <code>params()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.

chat_anthropic	<i>Chat with an Anthropic Claude model</i>
----------------	--

Description

Anthropic provides a number of chat based models under the **Claude** moniker. Note that a Claude Pro membership does not give you the ability to call models via the API; instead, you will need to sign up (and pay for) a **developer account**.

Usage

```
chat_anthropic(
  system_prompt = NULL,
  params = NULL,
  max_tokens = deprecated(),
  model = NULL,
  api_args = list(),
  base_url = "https://api.anthropic.com/v1",
  beta_headers = character(),
  api_key = anthropic_key(),
  api_headers = character(),
  echo = NULL
)

models_anthropic(
  base_url = "https://api.anthropic.com/v1",
  api_key = anthropic_key()
)
```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
params	Common model parameters, usually created by params() .
max_tokens	Maximum number of tokens to generate before stopping.
model	The model to use for the chat (defaults to "claude-sonnet-4-20250514"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use. Use models_anthropic() to see all options.
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with modifyList() .
base_url	The base URL to the endpoint; the default uses OpenAI.
beta_headers	Optionally, a character vector of beta headers to opt-in claude features that are still in beta.
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the ANTHROPIC_API_KEY environment variable. The best place to set this is in <code>.Renviro</code> , which you can easily edit by calling <code>usethis::edit_r_environ()</code> .
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
chat <- chat_anthropic()
chat$chat("Tell me three jokes about statisticians")
```

chat_aws_bedrock

*Chat with an AWS bedrock model***Description**

AWS Bedrock provides a number of language models, including those from Anthropic's **Claude**, using the Bedrock **Converse API**.

Authentication:

Authentication is handled through `{paws.common}`, so if authentication does not work for you automatically, you'll need to follow the advice at <https://www.paws-r-sdk.com/#credentials>. In particular, if your org uses AWS SSO, you'll need to run `aws sso login` at the terminal.

Usage

```
chat_aws_bedrock(
  system_prompt = NULL,
  base_url = NULL,
  model = NULL,
  profile = NULL,
  params = NULL,
  api_args = list(),
  api_headers = character(),
  echo = NULL
)

models_aws_bedrock(profile = NULL, base_url = NULL)
```

Arguments

<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>base_url</code>	The base URL to the endpoint; the default uses OpenAI.
<code>model</code>	The model to use for the chat (defaults to "anthropic.claude-3-5-sonnet-20240620-v1:0"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use. Use <code>models_models_aws_bedrock()</code> to see all options. . While <code>ellmer</code> provides a default model, there's no guarantee that you'll have access to it, so you'll need to specify a model that you can. If you're using cross-region inference , you'll need to use the inference profile ID, e.g. <code>model="us.anthropic.claude-3-5-so</code>
<code>profile</code>	AWS profile to use.
<code>params</code>	Common model parameters, usually created by <code>params()</code> .
<code>api_args</code>	Named list of arbitrary extra arguments appended to the body of every chat API call. Some useful arguments include:

```

    api_args = list(
      inferenceConfig = list(
        maxTokens = 100,
        temperature = 0.7,
        topP = 0.9,
        topK = 20
      )
    )
  )
}

api_headers Named character vector of arbitrary extra headers appended to every chat API
call.

echo One of the following options:
  • none: don't emit any output (default when running in a function).
  • output: echo text and tool-calling output as it streams in (default when
    running at the console).
  • all: echo all input and output.
Note this only affects the chat() method.

```

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```

## Not run:
# Basic usage
chat <- chat_aws_bedrock()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)

```

chat_azure_openai

Chat with a model hosted on Azure OpenAI

Description

The [Azure OpenAI server](#) hosts a number of open source models as well as proprietary models from OpenAI.

Authentication:

chat_azure_openai() supports API keys and the credentials parameter, but it also makes use of:

- Azure service principals (when the AZURE_TENANT_ID, AZURE_CLIENT_ID, and AZURE_CLIENT_SECRET environment variables are set).
- Interactive Entra ID authentication, like the Azure CLI.
- Viewer-based credentials on Posit Connect. Requires the **connectcreds** package.

Usage

```
chat_azure_openai(
  endpoint = azure_endpoint(),
  model,
  params = NULL,
  api_version = NULL,
  system_prompt = NULL,
  api_key = NULL,
  token = deprecated(),
  credentials = NULL,
  api_args = list(),
  echo = c("none", "output", "all"),
  api_headers = character(),
  deployment_id = deprecated()
)
```

Arguments

endpoint	Azure OpenAI endpoint url with protocol and hostname, i.e. <code>https://{your-resource-name}.openai.</code> Defaults to using the value of the <code>AZURE_OPENAI_ENDPOINT</code> environment variable.
model	The deployment id for the model you want to use.
params	Common model parameters, usually created by <code>params()</code> .
api_version	The API version to use.
system_prompt	A system prompt to set the behavior of the assistant.
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the <code>AZURE_OPENAI_API_KEY</code> environment variable. The best place to set this is in <code>.Renvirom</code> , which you can easily edit by calling <code>usethis::edit_r_envirom()</code> .
token	[Deprecated] A literal Azure token to use for authentication. Deprecated in favour of ambient Azure credentials or an explicit <code>credentials</code> argument.
credentials	A list of authentication headers to pass into <code>httr2::req_headers()</code> , a function that returns them, or <code>NULL</code> to use <code>token</code> or <code>api_key</code> to generate these headers instead. This is an escape hatch that allows users to incorporate Azure credentials generated by other packages into ellmer , or to manage the lifetime of credentials that need to be refreshed.

api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.
deployment_id	[Deprecated] Use <code>model</code> instead.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_azure_openai(deployment_id = "gpt-4o-mini")
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_cloudflare	<i>Chat with a model hosted on CloudFlare</i>
-----------------	---

Description

Cloudflare works AI hosts a variety of open-source AI models. To use the Cloudflare API, you must have an Account ID and an Access Token, which you can obtain [by following these instructions](#).

Known limitations:

- Tool calling does not appear to work.
- Images don't appear to work.

Usage

```
chat_cloudflare(
  account = cloudflare_account(),
  system_prompt = NULL,
  params = NULL,
  api_key = cloudflare_key(),
  model = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)
```

Arguments

account	The Cloudflare account ID. Taken from the CLOUDFLARE_ACCOUNT_ID env var, if defined.
system_prompt	A system prompt to set the behavior of the assistant.
params	Common model parameters, usually created by params() .
api_key	The API key to use for authentication. You generally should not supply this directly, but instead set the HUGGINGFACE_API_KEY environment variable.
model	The model to use for the chat (defaults to "meta-llama/Llama-3.3-70b-instruct-fp8-fast"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with modifyList() .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the chat() method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_cloudflare()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_cortex_analyst	<i>Create a chatbot that speaks to the Snowflake Cortex Analyst</i>
---------------------	---

Description

[Deprecated]

Please use [chat_snowflake\(\)](#) instead as that appears to be where Snowflake is putting their efforts.

Chat with the LLM-powered **Snowflake Cortex Analyst**.

Authentication:

chat_cortex_analyst() picks up the following ambient Snowflake credentials:

- A static OAuth token defined via the SNOWFLAKE_TOKEN environment variable.
- Key-pair authentication credentials defined via the SNOWFLAKE_USER and SNOWFLAKE_PRIVATE_KEY (which can be a PEM-encoded private key or a path to one) environment variables.
- Posit Workbench-managed Snowflake credentials for the corresponding account.
- Viewer-based credentials on Posit Connect. Requires the **connectcreds** package.

Known limitations:

Unlike most comparable model APIs, Cortex does not take a system prompt. Instead, the caller must provide a "semantic model" describing available tables, their meaning, and verified queries that can be run against them as a starting point. The semantic model can be passed as a YAML string or via reference to an existing file in a Snowflake Stage.

Note that Cortex does not support multi-turn, so it will not remember previous messages. Nor does it support registering tools, and attempting to do so will result in an error.

See [chat_snowflake\(\)](#) to chat with more general-purpose models hosted on Snowflake.

Usage

```
chat_cortex_analyst(
  account = snowflake_account(),
  credentials = NULL,
  model_spec = NULL,
  model_file = NULL,
  api_args = list(),
  echo = c("none", "output", "all"),
  api_headers = character()
)
```

Arguments

account	A Snowflake account identifier , e.g. "testorg-test_account". Defaults to the value of the SNOWFLAKE_ACCOUNT environment variable.
credentials	A list of authentication headers to pass into <code>httr2:req_headers()</code> , a function that returns them when called, or NULL, the default, to use ambient credentials.
model_spec	A semantic model specification, or NULL when using model_file instead.
model_file	Path to a semantic model file stored in a Snowflake Stage, or NULL when using model_spec instead.
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the chat() method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A `Chat` object.

chat_databricks	<i>Chat with a model hosted on Databricks</i>
-----------------	---

Description

Databricks provides out-of-the-box access to a number of **foundation models** and can also serve as a gateway for external models hosted by a third party.

Authentication:

chat_databricks() picks up on ambient Databricks credentials for a subset of the **Databricks client unified authentication** model. Specifically, it supports:

- Personal access tokens
- Service principals via OAuth (OAuth M2M)
- User account via OAuth (OAuth U2M)
- Authentication via the Databricks CLI
- Posit Workbench-managed credentials
- Viewer-based credentials on Posit Connect. Requires the **connectcreds** package.

Known limitations:

Databricks models do not support images, but they do support structured outputs and tool calls for most models.

Usage

```
chat_databricks(
  workspace = databricks_workspace(),
  system_prompt = NULL,
  model = NULL,
  token = NULL,
  params = NULL,
  api_args = list(),
  echo = c("none", "output", "all"),
  api_headers = character()
)
```

Arguments

workspace	The URL of a Databricks workspace, e.g. "https://example.cloud.databricks.com". Will use the value of the environment variable DATABRICKS_HOST, if set.
system_prompt	A system prompt to set the behavior of the assistant.
model	The model to use for the chat (defaults to "databricks-claude-3-7-sonnet"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use. Available foundational models include: • databricks-claude-3-7-sonnet (the default) • databricks-mixtral-8x7b-instruct • databricks-meta-llama-3-1-70b-instruct • databricks-meta-llama-3-1-405b-instruct
token	An authentication token for the Databricks workspace, or NULL to use ambient credentials.
params	Common model parameters, usually created by params() .
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with modifyList() .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: `chat_anthropic()`, `chat_aws_bedrock()`, `chat_azure_openai()`, `chat_cloudflare()`, `chat_deepseek()`, `chat_github()`, `chat_google_gemini()`, `chat_groq()`, `chat_huggingface()`, `chat_mistral()`, `chat_ollama()`, `chat_openai()`, `chat_openrouter()`, `chat_perplexity()`, `chat_portkey()`

Examples

```
## Not run:
chat <- chat_databricks()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_deepseek

Chat with a model hosted on DeepSeek

Description

Sign up at <https://platform.deepseek.com>.

Known limitations:

- Structured data extraction is not supported.
- Images are not supported.

Usage

```
chat_deepseek(
  system_prompt = NULL,
  base_url = "https://api.deepseek.com",
  api_key = deepseek_key(),
  model = NULL,
  params = NULL,
  seed = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)
```

Arguments

<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>base_url</code>	The base URL to the endpoint; the default uses DeepSeek.
<code>api_key</code>	API key to use for authentication. You generally should not supply this directly, but instead set the <code>DEEPSEEK_API_KEY</code> environment variable. The best place to set this is in <code>.Renvirom</code> , which you can easily edit by calling <code>usethis::edit_r_envirom()</code> .

model	The model to use for the chat (defaults to "deepseek-chat"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
params	Common model parameters, usually created by params() .
seed	Optional integer seed that ChatGPT uses to try and make output more reproducible.
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with modifyList() .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_deepseek()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

Description

GitHub Models hosts a number of open source and OpenAI models. To access the GitHub model marketplace, you will need to apply for and be accepted into the beta access program. See <https://github.com/marketplace/models> for details.

This function is a lightweight wrapper around `chat_openai()` with the defaults tweaked for the GitHub Models marketplace.

GitHub also supports the Azure AI Inference SDK, which you can use by setting `base_url` to `"https://models.inference.ai.azure.com/"`. This endpoint was used in **ellmer** v0.3.0 and earlier.

Usage

```
chat_github(
  system_prompt = NULL,
  base_url = "https://models.github.ai/inference/",
  api_key = github_key(),
  model = NULL,
  params = NULL,
  seed = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)

models_github(base_url = "https://models.github.ai/", api_key = github_key())
```

Arguments

<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>base_url</code>	The base URL to the endpoint; the default uses OpenAI.
<code>api_key</code>	API key to use for authentication. You generally should not supply this directly, but instead set the <code>GITHUB_PAT</code> environment variable. The best place to set this is in <code>.Renvirom</code> , which you can easily edit by calling <code>usethis::edit_r_envirom()</code> .
<code>model</code>	The model to use for the chat (defaults to "gpt-4o"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
<code>params</code>	Common model parameters, usually created by <code>params()</code> .
<code>seed</code>	Optional integer seed that ChatGPT uses to try and make output more reproducible.
<code>api_args</code>	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with <code>modifyList()</code> .
<code>echo</code>	One of the following options: • <code>none</code>: don't emit any output (default when running in a function).

- output: echo text and tool-calling output as it streams in (default when running at the console).
- all: echo all input and output.

Note this only affects the `chat()` method.

`api_headers` Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_github()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_google_gemini *Chat with a Google Gemini or Vertex AI model*

Description

Google's AI offering is broken up into two parts: Gemini and Vertex AI. Most enterprises are likely to use Vertex AI, and individuals are likely to use Gemini.

Use [google_upload\(\)](#) to upload files (PDFs, images, video, audio, etc.)

Authentication:

By default, `chat_google_gemini()` will use Google's default application credentials if there is no API key provided. This requires the **gargle** package.

It can also pick up on viewer-based credentials on Posit Connect. This in turn requires the **connectcreds** package.

Usage

```

chat_google_gemini(
  system_prompt = NULL,
  base_url = "https://generativelanguage.googleapis.com/v1beta/",
  api_key = NULL,
  model = NULL,
  params = NULL,
  api_args = list(),
  api_headers = character(),
  echo = NULL
)

chat_google_vertex(
  location,
  project_id,
  system_prompt = NULL,
  model = NULL,
  params = NULL,
  api_args = list(),
  api_headers = character(),
  echo = NULL
)

models_google_gemini(
  base_url = "https://generativelanguage.googleapis.com/v1beta/",
  api_key = NULL
)

models_google_vertex(location, project_id)

```

Arguments

<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>base_url</code>	The base URL to the endpoint; the default uses OpenAI.
<code>api_key</code>	API key to use for authentication. You generally should not supply this directly, but instead set the <code>GOOGLE_API_KEY</code> environment variable. The best place to set this is in <code>.Renv</code> , which you can easily edit by calling <code>usethis::edit_r_environ()</code> . For Gemini, you can alternatively set <code>GEMINI_API_KEY</code> .
<code>model</code>	The model to use for the chat (defaults to "gemini-2.5-flash"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use. Use <code>models_google_gemini()</code> to see all options.
<code>params</code>	Common model parameters, usually created by <code>params()</code> .
<code>api_args</code>	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with <code>modifyList()</code> .

api_headers	Named character vector of arbitrary extra headers appended to every chat API call.
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the chat() method.
location	Location, e.g. us-east1, me-central1, africa-south1 or global.
project_id	Project ID.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_google_gemini()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_groq	<i>Chat with a model hosted on Groq</i>
-----------	---

Description

Sign up at <https://groq.com>.
This function is a lightweight wrapper around [chat_openai\(\)](#) with the defaults tweaked for groq.

Known limitations:
groq does not currently support structured data extraction.

Usage

```
chat_groq(
  system_prompt = NULL,
  base_url = "https://api.groq.com/openai/v1",
  api_key = groq_key(),
  model = NULL,
  params = NULL,
  seed = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)
```

Arguments

<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>base_url</code>	The base URL to the endpoint; the default uses OpenAI.
<code>api_key</code>	API key to use for authentication. You generally should not supply this directly, but instead set the <code>GROQ_API_KEY</code> environment variable. The best place to set this is in <code>.Renviro</code> , which you can easily edit by calling <code>usethis::edit_r_enviro()</code> .
<code>model</code>	The model to use for the chat (defaults to "llama3-8b-8192"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
<code>params</code>	Common model parameters, usually created by <code>params()</code> .
<code>seed</code>	Optional integer seed that ChatGPT uses to try and make output more reproducible.
<code>api_args</code>	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with <code>modifyList()</code> .
<code>echo</code>	One of the following options: • <code>none</code>: don't emit any output (default when running in a function). • <code>output</code>: echo text and tool-calling output as it streams in (default when running at the console). • <code>all</code>: echo all input and output. Note this only affects the <code>chat()</code> method.
<code>api_headers</code>	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_groq()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_huggingface

Chat with a model hosted on Hugging Face Serverless Inference API

Description

Hugging Face hosts a variety of open-source and proprietary AI models available via their Inference API. To use the Hugging Face API, you must have an Access Token, which you can obtain from your **Hugging Face account** (ensure that at least "Make calls to Inference Providers" and "Make calls to your Inference Endpoints" is checked).

This function is a lightweight wrapper around [chat_openai\(\)](#), with the defaults adjusted for Hugging Face.

Known limitations:

- Some models do not support the chat interface or parts of it, for example google/gemma-2-2b-it does not support a system prompt. You will need to carefully choose the model.

Usage

```
chat_huggingface(
  system_prompt = NULL,
  params = NULL,
  api_key = hf_key(),
  model = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)
```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
params	Common model parameters, usually created by params() .
api_key	The API key to use for authentication. You generally should not supply this directly, but instead set the HUGGINGFACE_API_KEY environment variable.
model	The model to use for the chat (defaults to "meta-llama/Llama-3.1-8B-Instruct"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with modifyList() .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the chat() method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_huggingface()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_mistral

*Chat with a model hosted on Mistral's La Plateforme***Description**

Get your API key from <https://console.mistral.ai/api-keys>.

Known limitations:

- Tool calling is unstable.
- Images require a model that supports images.

Usage

```
chat_mistral(
  system_prompt = NULL,
  params = NULL,
  api_key = mistral_key(),
  model = NULL,
  seed = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)
```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
params	Common model parameters, usually created by <code>params()</code> .
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the MISTRAL_API_KEY environment variable. The best place to set this is in <code>.Renvirom</code> , which you can easily edit by calling <code>usethis::edit_r_environ()</code> .
model	The model to use for the chat (defaults to "mistral-large-latest"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
seed	Optional integer seed that ChatGPT uses to try and make output more reproducible.
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output.

Note this only affects the `chat()` method.

`api_headers` Named character vector of arbitrary extra headers appended to every chat API call.

Value

A `Chat` object.

See Also

Other chatbots: `chat_anthropic()`, `chat_aws_bedrock()`, `chat_azure_openai()`, `chat_cloudflare()`, `chat_databricks()`, `chat_deepseek()`, `chat_github()`, `chat_google_gemini()`, `chat_groq()`, `chat_huggingface()`, `chat_ollama()`, `chat_openai()`, `chat_openrouter()`, `chat_perplexity()`, `chat_portkey()`

Examples

```
## Not run:
chat <- chat_mistral()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_ollama

Chat with a local Ollama model

Description

To use `chat_ollama()` first download and install **Ollama**. Then install some models either from the command line (e.g. with `ollama pull llama3.1`) or within R using **ollamar** (e.g. `ollamar::pull("llama3.1")`). This function is a lightweight wrapper around `chat_openai()` with the defaults tweaked for ollama.

Known limitations:

- Tool calling is not supported with streaming (i.e. when `echo` is "text" or "all")
- Models can only use 2048 input tokens, and there's no way to get them to use more, except by creating a custom model with a different default.
- Tool calling generally seems quite weak, at least with the models I have tried it with.

Usage

```
chat_ollama(
  system_prompt = NULL,
  base_url = Sys.getenv("OLLAMA_BASE_URL", "http://localhost:11434"),
  model,
  seed = NULL,
  params = NULL,
  api_args = list(),
```



```

    echo = NULL,
    api_key = NULL,
    api_headers = character()
)

models_ollama(base_url = "http://localhost:11434")

```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
base_url	The base URL to the endpoint; the default uses OpenAI.
model	The model to use for the chat. Use <code>models_ollama()</code> to see all options.
seed	Optional integer seed that ChatGPT uses to try and make output more reproducible.
params	Common model parameters, usually created by <code>params()</code> .
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_key	Ollama doesn't require an API key for local usage and in most cases you do not need to provide an <code>api_key</code> . However, if you're accessing an Ollama instance hosted behind a reverse proxy or secured endpoint that enforces bearer-token authentication, you can set <code>api_key</code> (or the <code>OLLAMA_API_KEY</code> environment variable).
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A `Chat` object.

See Also

Other chatbots: `chat_anthropic()`, `chat_aws_bedrock()`, `chat_azure_openai()`, `chat_cloudflare()`, `chat_databricks()`, `chat_deepseek()`, `chat_github()`, `chat_google_gemini()`, `chat_groq()`, `chat_huggingface()`, `chat_mistral()`, `chat_openai()`, `chat_openrouter()`, `chat_perplexity()`, `chat_portkey()`

Examples

```
## Not run:
chat <- chat_ollama(model = "llama3.2")
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_openai

Chat with an OpenAI model

Description

OpenAI provides a number of chat-based models, mostly under the **ChatGPT** brand. Note that a ChatGPT Plus membership does not grant access to the API. You will need to sign up for a developer account (and pay for it) at the [developer platform](#).

Usage

```
chat_openai(
  system_prompt = NULL,
  base_url = Sys.getenv("OPENAI_BASE_URL", "https://api.openai.com/v1"),
  api_key = openai_key(),
  model = NULL,
  params = NULL,
  seed = lifecycle::deprecated(),
  api_args = list(),
  api_headers = character(),
  echo = c("none", "output", "all")
)

models_openai(base_url = "https://api.openai.com/v1", api_key = openai_key())
```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
base_url	The base URL to the endpoint; the default uses OpenAI.
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the OPENAI_API_KEY environment variable. The best place to set this is in .Renvirom, which you can easily edit by calling <code>usethis::edit_r_environ()</code> .
model	The model to use for the chat (defaults to "gpt-4.1"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use. Use <code>models_openai()</code> to see all options.
params	Common model parameters, usually created by <code>params()</code> .
seed	Optional integer seed that ChatGPT uses to try and make output more reproducible.

api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with <code>modifyList()</code> .
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
chat <- chat_openai()
chat$chat("
  What is the difference between a tibble and a data frame?
  Answer with a bulleted list
")

chat$chat("Tell me three funny jokes about statisticians")
```

chat_openrouter

Chat with one of the many models hosted on OpenRouter

Description

Sign up at <https://openrouter.ai>.

Support for features depends on the underlying model that you use; see <https://openrouter.ai/models> for details.

Usage

```
chat_openrouter(
  system_prompt = NULL,
  api_key = openrouter_key(),
  model = NULL,
  seed = NULL,
  params = NULL,
  api_args = list(),
  echo = c("none", "output", "all"),
  api_headers = character()
)
```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the OPENROUTER_API_KEY environment variable. The best place to set this is in .Renvirom, which you can easily edit by calling <code>usethis::edit_r_environ()</code> .
model	The model to use for the chat (defaults to "gpt-4o"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
seed	Optional integer seed that ChatGPT uses to try and make output more reproducible.
params	Common model parameters, usually created by <code>params()</code> .
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_perplexity\(\)](#), [chat_portkey\(\)](#)

Examples

```
## Not run:
chat <- chat_openrouter()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_perplexity	<i>Chat with a model hosted on perplexity.ai</i>
-----------------	--

Description

Sign up at <https://www.perplexity.ai>.

Perplexity AI is a platform for running LLMs that are capable of searching the web in real-time to help them answer questions with information that may not have been available when the model was trained.

This function is a lightweight wrapper around `chat_openai()` with the defaults tweaked for Perplexity AI.

Usage

```
chat_perplexity(
  system_prompt = NULL,
  base_url = "https://api.perplexity.ai/",
  api_key = perplexity_key(),
  model = NULL,
  seed = NULL,
  params = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)
```

Arguments

<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>base_url</code>	The base URL to the endpoint; the default uses OpenAI.
<code>api_key</code>	API key to use for authentication. You generally should not supply this directly, but instead set the <code>PERPLEXITY_API_KEY</code> environment variable. The best place to set this is in <code>.Renvi</code> , which you can easily edit by calling <code>usethis::edit_r_environ()</code> .
<code>model</code>	The model to use for the chat (defaults to "llama-3.1-sonar-small-128k-online"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.

seed	Optional integer seed that ChatGPT uses to try and make output more reproducible.
params	Common model parameters, usually created by <code>params()</code> .
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A `Chat` object.

See Also

Other chatbots: `chat_anthropic()`, `chat_aws_bedrock()`, `chat_azure_openai()`, `chat_cloudflare()`, `chat_databricks()`, `chat_deepseek()`, `chat_github()`, `chat_google_gemini()`, `chat_groq()`, `chat_huggingface()`, `chat_mistral()`, `chat_ollama()`, `chat_openai()`, `chat_openrouter()`, `chat_portkey()`

Examples

```
## Not run:
chat <- chat_perplexity()
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_portkey

Chat with a model hosted on PortkeyAI

Description

PortkeyAI provides an interface (AI Gateway) to connect through its Universal API to a variety of LLMs providers with a single endpoint.

Authentication:

API keys together with configurations of LLM providers are stored inside Portkey application.

Usage

```
chat_portkey(
  system_prompt = NULL,
  base_url = "https://api.portkey.ai/v1",
  api_key = portkey_key(),
  virtual_key = portkey_virtual_key(),
  model = NULL,
  params = NULL,
  api_args = list(),
  echo = NULL,
  api_headers = character()
)

models_portkey(
  base_url = "https://api.portkey.ai/v1",
  api_key = portkey_key(),
  virtual_key = NULL
)
```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
base_url	The base URL to the endpoint; the default uses OpenAI.
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the PORTKEY_API_KEY environment variable. The best place to set this is in <code>.Renviro</code> , which you can easily edit by calling <code>usethis::edit_r_enviro()</code> .
virtual_key	A virtual identifier storing LLM provider's API key. See documentation . Can be read from the PORTKEY_VIRTUAL_KEY environment variable.
model	The model to use for the chat (defaults to "gpt-4o"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use. Use <code>models_openai()</code> to see all options.
params	Common model parameters, usually created by params() .
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with modifyList() .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

See Also

Other chatbots: [chat_anthropic\(\)](#), [chat_aws_bedrock\(\)](#), [chat_azure_openai\(\)](#), [chat_cloudflare\(\)](#), [chat_databricks\(\)](#), [chat_deepseek\(\)](#), [chat_github\(\)](#), [chat_google_gemini\(\)](#), [chat_groq\(\)](#), [chat_huggingface\(\)](#), [chat_mistral\(\)](#), [chat_ollama\(\)](#), [chat_openai\(\)](#), [chat_openrouter\(\)](#), [chat_perplexity\(\)](#)

Examples

```
## Not run:
chat <- chat_portkey(virtual_key = Sys.getenv("PORTKEY_VIRTUAL_KEY"))
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

chat_snowflake

Chat with a model hosted on Snowflake

Description

The Snowflake provider allows you to interact with LLM models available through the [Cortex LLM REST API](#).

Authentication:

chat_snowflake() picks up the following ambient Snowflake credentials:

- A static OAuth token defined via the SNOWFLAKE_TOKEN environment variable.
- Key-pair authentication credentials defined via the SNOWFLAKE_USER and SNOWFLAKE_PRIVATE_KEY (which can be a PEM-encoded private key or a path to one) environment variables.
- Posit Workbench-managed Snowflake credentials for the corresponding account.
- Viewer-based credentials on Posit Connect. Requires the **connectcreds** package.

Known limitations:

Note that Snowflake-hosted models do not support images.

See [chat_cortex_analyst\(\)](#) to chat with the Snowflake Cortex Analyst rather than a general-purpose model.

Usage

```
chat_snowflake(
  system_prompt = NULL,
  account = snowflake_account(),
  credentials = NULL,
  model = NULL,
  params = NULL,
```



```

    api_args = list(),
    echo = c("none", "output", "all"),
    api_headers = character()
  )

```

Arguments

system_prompt	A system prompt to set the behavior of the assistant.
account	A Snowflake account identifier , e.g. "testorg-test_account". Defaults to the value of the SNOWFLAKE_ACCOUNT environment variable.
credentials	A list of authentication headers to pass into <code>httr2::req_headers()</code> , a function that returns them when called, or NULL, the default, to use ambient credentials.
model	The model to use for the chat (defaults to "claude-3-7-sonnet"). We regularly update the default, so we strongly recommend explicitly specifying a model for anything other than casual use.
params	Common model parameters, usually created by <code>params()</code> .
api_args	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by ellmer with <code>modifyList()</code> .
echo	One of the following options: • none: don't emit any output (default when running in a function). • output: echo text and tool-calling output as it streams in (default when running at the console). • all: echo all input and output. Note this only affects the <code>chat()</code> method.
api_headers	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A `Chat` object.

Examples

```

chat <- chat_snowflake()
chat$chat("Tell me a joke in the form of a SQL query.")

```

chat_vllm

Chat with a model hosted by vLLM

Description

vLLM is an open source library that provides an efficient and convenient LLMs model server. You can use `chat_vllm()` to connect to endpoints powered by vLLM.

Usage

```
chat_vllm(
  base_url,
  system_prompt = NULL,
  model,
  seed = NULL,
  params = NULL,
  api_args = list(),
  api_key = vllm_key(),
  echo = NULL,
  api_headers = character()
)

models_vllm(base_url, api_key = vllm_key())
```

Arguments

<code>base_url</code>	The base URL to the endpoint; the default uses OpenAI.
<code>system_prompt</code>	A system prompt to set the behavior of the assistant.
<code>model</code>	The model to use for the chat. Use <code>models_vllm()</code> to see all options.
<code>seed</code>	Optional integer seed that ChatGPT uses to try and make output more reproducible.
<code>params</code>	Common model parameters, usually created by params() .
<code>api_args</code>	Named list of arbitrary extra arguments appended to the body of every chat API call. Combined with the body object generated by <code>ellmer</code> with modifyList() .
<code>api_key</code>	API key to use for authentication. You generally should not supply this directly, but instead set the <code>VLLM_API_KEY</code> environment variable. The best place to set this is in <code>.Renvi</code> , which you can easily edit by calling <code>usethis::edit_r_envi</code> .
<code>echo</code>	One of the following options: • <code>none</code>: don't emit any output (default when running in a function). • <code>output</code>: echo text and tool-calling output as it streams in (default when running at the console). • <code>all</code>: echo all input and output. Note this only affects the <code>chat()</code> method.
<code>api_headers</code>	Named character vector of arbitrary extra headers appended to every chat API call.

Value

A [Chat](#) object.

Examples

```
## Not run:
chat <- chat_vllm("http://my-vllm.com")
chat$chat("Tell me three jokes about statisticians")

## End(Not run)
```

Content	<i>Content types received from and sent to a chatbot</i>
---------	--

Description

Use these functions if you're writing a package that extends ellmer and need to customise methods for various types of content. For normal use, see [content_image_url\(\)](#) and friends.

ellmer abstracts away differences in the way that different [Providers](#) represent various types of content, allowing you to more easily write code that works with any chatbot. This set of classes represents types of content that can be either sent to and received from a provider:

- `ContentText`: simple text (often in markdown format). This is the only type of content that can be streamed live as it's received.
- `ContentImageRemote` and `ContentImageInline`: images, either as a pointer to a remote URL or included inline in the object. See [content_image_file\(\)](#) and friends for convenient ways to construct these objects.
- `ContentToolRequest`: a request to perform a tool call (sent by the assistant).
- `ContentToolResult`: the result of calling the tool (sent by the user). This object is automatically created from the value returned by calling the [tool\(\)](#) function. Alternatively, expert users can return a `ContentToolResult` from a [tool\(\)](#) function to include additional data or to customize the display of the result.

Usage

```
Content()
```

```
ContentText(text = stop("Required"))
```

```
ContentImage()
```

```
ContentImageRemote(url = stop("Required"), detail = "")
```

```
ContentImageInline(type = stop("Required"), data = NULL)
```

```
ContentToolRequest(
  id = stop("Required"),
  name = stop("Required"),
  arguments = list(),
  tool = NULL
```

```

)

ContentToolResult(value = NULL, error = NULL, extra = list(), request = NULL)

ContentThinking(thinking = stop("Required"), extra = list())

ContentPDF(
  type = stop("Required"),
  data = stop("Required"),
  filename = stop("Required")
)

```

Arguments

text	A single string.
url	URL to a remote image.
detail	Not currently used.
type	MIME type of the image.
data	Base64 encoded image data.
id	Tool call id (used to associate a request and a result). Automatically managed by ellmer .
name	Function name
arguments	Named list of arguments to call the function with.
tool	ellmer automatically matches a tool request to the tools defined for the chatbot. If NULL, the request did not match a defined tool.
value	The results of calling the tool function, if it succeeded.
error	The error message, as a string, or the error condition thrown as a result of a failure when calling the tool function. Must be NULL when the tool call is successful.
extra	Additional data.
request	The ContentToolRequest associated with the tool result, automatically added by ellmer when evaluating the tool call.
thinking	The text of the thinking output.
filename	File name, used to identify the PDF.

Value

S7 objects that all inherit from Content

Examples

```

Content()
ContentText("Tell me a joke")
ContentImageRemote("https://www.r-project.org/Rlogo.png")
ContentToolRequest(id = "abc", name = "mean", arguments = list(x = 1:5))

```

contents_text	<i>Format contents into a textual representation</i>
---------------	--

Description

[Experimental]

These generic functions can be use to convert [Turn](#) contents or [Content](#) objects into textual representations.

- `contents_text()` is the most minimal and only includes [ContentText](#) objects in the output.
- `contents_markdown()` returns the text content (which it assumes to be markdown and does not convert it) plus markdown representations of images and other content types.
- `contents_html()` returns the text content, converted from markdown to HTML with `commonmark::markdown_html()` plus HTML representations of images and other content types.

These content types will continue to grow and change as ellmer evolves to support more providers and as providers add more content types.

Usage

```
contents_text(content, ...)
```

```
contents_html(content, ...)
```

```
contents_markdown(content, ...)
```

Arguments

content	The Turn or Content object to be converted into text. <code>contents_markdown()</code> also accepts Chat instances to turn the entire conversation history into markdown text.
...	Additional arguments passed to methods.

Value

A string of text, markdown or HTML.

Examples

```
turns <- list(
  Turn("user", contents = list(
    ContentText("What's this image?"),
    content_image_url("https://placeholder.co/200x200")
  )),
  Turn("assistant", "It's a placeholder image.")
)
```

```
lapply(turns, contents_text)
lapply(turns, contents_markdown)
if (rlang::is_installed("commonmark")) {
  contents_html(turns[[1]])
}
```

content_image_url	<i>Encode images for chat input</i>
-------------------	-------------------------------------

Description

These functions are used to prepare image URLs and files for input to the chatbot. The `content_image_url()` function is used to provide a URL to an image, while `content_image_file()` is used to provide the image data itself.

Usage

```
content_image_url(url, detail = c("auto", "low", "high"))
```

```
content_image_file(path, content_type = "auto", resize = "low")
```

```
content_image_plot(width = 768, height = 768)
```

Arguments

<code>url</code>	The URL of the image to include in the chat input. Can be a data: URL or a regular URL. Valid image types are PNG, JPEG, WebP, and non-animated GIF.
<code>detail</code>	The detail setting for this image. Can be "auto", "low", or "high".
<code>path</code>	The path to the image file to include in the chat input. Valid file extensions are .png, .jpeg, .jpg, .webp, and (non-animated) .gif.
<code>content_type</code>	The content type of the image (e.g. image/png). If "auto", the content type is inferred from the file extension.
<code>resize</code>	If "low", resize images to fit within 512x512. If "high", resize to fit within 2000x768 or 768x2000. (See the OpenAI docs for more on why these specific sizes are used.) If "none", do not resize. You can also pass a custom string to resize the image to a specific size, e.g. "200x200" to resize to 200x200 pixels while preserving aspect ratio. Append > to resize only if the image is larger than the specified size, and ! to ignore aspect ratio (e.g. "300x200>!"). All values other than none require the magick package.
<code>width, height</code>	Width and height in pixels.

Value

An input object suitable for including in the ... parameter of the `chat()`, `stream()`, `chat_async()`, or `stream_async()` methods.

Examples

```
## Not run:
chat <- chat_openai()
chat$chat(
  "What do you see in these images?",
  content_image_url("https://www.r-project.org/Rlogo.png"),
  content_image_file(system.file("httr2.png", package = "ellmer"))
)

plot(waiting ~ eruptions, data = faithful)
chat <- chat_openai()
chat$chat(
  "Describe this plot in one paragraph, as suitable for inclusion in
  alt-text. You should briefly describe the plot type, the axes, and
  2-5 major visual patterns.",
  content_image_plot()
)

## End(Not run)
```

content_pdf_file

Encode PDFs content for chat input

Description

These functions are used to prepare PDFs as input to the chatbot. The `content_pdf_url()` function is used to provide a URL to an PDF file, while `content_pdf_file()` is used to for local PDF files.

Not all providers support PDF input, so check the documentation for the provider you are using.

Usage

```
content_pdf_file(path)
```

```
content_pdf_url(url)
```

Arguments

path, url Path or URL to a PDF file.

Value

A ContentPDF object

create_tool_def	Create metadata for a tool
-----------------	----------------------------

Description

In order to use a function as a tool in a chat, you need to craft the right call to `tool()`. This function helps you do that for documented functions by extracting the function's R documentation and using an LLM to generate the `tool()` call. It's meant to be used interactively while writing your code, not as part of your final code.

If the function has package documentation, that will be used. Otherwise, if the source code of the function can be automatically detected, then the comments immediately preceding the function are used (especially helpful if those are roxygen2 comments). If neither are available, then just the function signature is used.

Note that this function is inherently imperfect. It can't handle all possible R functions, because not all parameters are suitable for use in a tool call (for example, because they're not serializable to simple JSON objects). The documentation might not specify the expected shape of arguments to the level of detail that would allow an exact JSON schema to be generated. Please be sure to review the generated code before using it!

Usage

```
create_tool_def(
  topic,
  chat = NULL,
  model = deprecated(),
  echo = interactive(),
  verbose = FALSE
)
```

Arguments

topic	A symbol or string literal naming the function to create metadata for. Can also be an expression of the form <code>pkg::fun</code> .
chat	A Chat object used to generate the output. If NULL (the default) uses <code>chat_openai()</code> .
model	<code>lifecycle::badge("deprecated")</code> Formally used for defining the model used by the chat. Now supply chat instead.
echo	Emit the registration code to the console. Defaults to TRUE in interactive sessions.
verbose	If TRUE, print the input we send to the LLM, which may be useful for debugging unexpectedly poor results.

Value

A `register_tool` call that you can copy and paste into your code. Returned invisibly if echo is TRUE.

Examples

```
## Not run:
# These are all equivalent
create_tool_def(rnorm)
create_tool_def(stats::rnorm)
create_tool_def("rnorm")
create_tool_def("rnorm", chat = chat_azure_openai())

## End(Not run)
```

google_upload

Upload a file to gemini

Description

[Experimental]

This function uploads a file then waits for Gemini to finish processing it so that you can immediately use it in a prompt. It's experimental because it's currently Gemini specific, and we expect other providers to evolve similar feature in the future.

Uploaded files are automatically deleted after 2 days. Each file must be less than 2 GB and you can upload a total of 20 GB. ellmer doesn't currently provide a way to delete files early; please [file an issue](#) if this would be useful for you.

Usage

```
google_upload(
  path,
  base_url = "https://generativelanguage.googleapis.com/",
  api_key = NULL,
  mime_type = NULL
)
```

Arguments

path	Path to a file to upload.
base_url	The base URL to the endpoint; the default uses OpenAI.
api_key	API key to use for authentication. You generally should not supply this directly, but instead set the <code>GOOGLE_API_KEY</code> environment variable. The best place to set this is in <code>.Renvi</code> , which you can easily edit by calling <code>usethis::edit_r_envi</code> . For Gemini, you can alternatively set <code>GEMINI_API_KEY</code> .
mime_type	Optionally, specify the mime type of the file. If not specified, will be guesses from the file extension.

Value

A <ContentUploaded> object that can be passed to `$chat()`.

Examples

```
## Not run:
file <- google_upload("path/to/file.pdf")

chat <- chat_google_gemini()
chat$chat(file, "Give me a three paragraph summary of this PDF")

## End(Not run)
```

interpolate

Helpers for interpolating data into prompts

Description

These functions are lightweight wrappers around **glue** that make it easier to interpolate dynamic data into a static prompt:

- `interpolate()` works with a string.
- `interpolate_file()` works with a file.
- `interpolate_package()` works with a file in the `insts/prompt` directory of a package.

Compared to `glue`, dynamic values should be wrapped in `{ { }}`, making it easier to include R code and JSON in your prompt.

Usage

```
interpolate(prompt, ..., .envir = parent.frame())

interpolate_file(path, ..., .envir = parent.frame())

interpolate_package(package, path, ..., .envir = parent.frame())
```

Arguments

<code>prompt</code>	A prompt string. You should not generally expose this to the end user, since <code>glue</code> interpolation makes it easy to run arbitrary code.
<code>...</code>	Define additional temporary variables for substitution.
<code>.envir</code>	Environment to evaluate <code>...</code> expressions in. Used when wrapping in another function. See <code>vignette("wrappers", package = "glue")</code> for more details.
<code>path</code>	A path to a prompt file (often a <code>.md</code>).
<code>package</code>	Package name.

Value

A {glue} string.

Examples

```
joke <- "You're a cool dude who loves to make jokes. Tell me a joke about {{topic}}."

# You can supply values directly:
interpolate(joke, topic = "bananas")

# Or allow interpolate to find them in the current environment:
topic <- "applies"
interpolate(joke)
```

live_console

Open a live chat application

Description

- `live_console()` lets you chat interactively in the console.
- `live_browser()` lets you chat interactively in a browser.

Note that these functions will mutate the input chat object as you chat because your turns will be appended to the history.

Usage

```
live_console(chat, quiet = FALSE)
```

```
live_browser(chat, quiet = FALSE)
```

Arguments

chat	A chat object created by <code>chat_openai()</code> or friends.
quiet	If TRUE, suppresses the initial message that explains how to use the console.

Value

(Invisibly) The input chat.

Examples

```
## Not run:
chat <- chat_anthropic()
live_console(chat)
live_browser(chat)

## End(Not run)
```

parallel_chat

*Submit multiple chats in parallel***Description****[Experimental]**

If you have multiple prompts, you can submit them in parallel. This is typically considerably faster than submitting them in sequence, especially with Gemini and OpenAI.

If you're using `chat_openai()` or `chat_anthropic()` and you're willing to wait longer, you might want to use `batch_chat()` instead, as it comes with a 50% discount in return for taking up to 24 hours.

Usage

```
parallel_chat(chat, prompts, max_active = 10, rpm = 500)
```

```
parallel_chat_text(chat, prompts, max_active = 10, rpm = 500)
```

```
parallel_chat_structured(
  chat,
  prompts,
  type,
  convert = TRUE,
  include_tokens = FALSE,
  include_cost = FALSE,
  max_active = 10,
  rpm = 500
)
```

Arguments

chat	A base chat object.
prompts	A vector created by <code>interpolate()</code> or a list of character vectors.
max_active	The maximum number of simultaneous requests to send. For <code>chat_anthropic()</code> , note that the number of active connections is limited primarily by the output tokens per minute limit (OTPM) which is estimated from the <code>max_tokens</code> parameter, which defaults to 4096. That means if your usage tier limits you to 16,000 OTPM, you should either set <code>max_active = 4</code> (16,000 / 4096) to decrease the number of active connections or use <code>params()</code> in <code>chat_anthropic()</code> to decrease <code>max_tokens</code> .
rpm	Maximum number of requests per minute.
type	A type specification for the extracted data. Should be created with a <code>type_()</code> function.

convert	If TRUE, automatically convert from JSON lists to R data types using the schema. This typically works best when type is <code>type_object()</code> as this will give you a data frame with one column for each property. If FALSE, returns a list.
include_tokens	If TRUE, and the result is a data frame, will add <code>input_tokens</code> and <code>output_tokens</code> columns giving the total input and output tokens for each prompt.
include_cost	If TRUE, and the result is a data frame, will add <code>cost</code> column giving the cost of each prompt.

Value

For `parallel_chat()`, a list of [Chat](#) objects, one for each prompt. For `parallel_chat_text()`, a character vector of text responses. For `parallel_chat_structured()`, a single structured data object with one element for each prompt. Typically, when type is an object, this will be a data frame with one row for each prompt, and one column for each property.

Examples

```
chat <- chat_openai()

# Chat -----
country <- c("Canada", "New Zealand", "Jamaica", "United States")
prompts <- interpolate("What's the capital of {{country}}?")
parallel_chat(chat, prompts)

# Structured data -----
prompts <- list(
  "I go by Alex. 42 years on this planet and counting.",
  "Pleased to meet you! I'm Jamal, age 27.",
  "They call me Li Wei. Nineteen years young.",
  "Fatima here. Just celebrated my 35th birthday last week.",
  "The name's Robert - 51 years old and proud of it.",
  "Kwame here - just hit the big 5-0 this year."
)
type_person <- type_object(name = type_string(), age = type_number())
parallel_chat_structured(chat, prompts, type_person)
```

params

Standard model parameters

Description

This helper function makes it easier to create a list of parameters used across many models. The parameter names are automatically standardised and included in the correctly place in the API call.

Note that parameters that are not supported by a given provider will generate a warning, not an error. This allows you to use the same set of parameters across multiple providers.

Usage

```
params(  
  temperature = NULL,  
  top_p = NULL,  
  top_k = NULL,  
  frequency_penalty = NULL,  
  presence_penalty = NULL,  
  seed = NULL,  
  max_tokens = NULL,  
  log_probs = NULL,  
  stop_sequences = NULL,  
  ...  
)
```

Arguments

- temperature Temperature of the sampling distribution.
- top_p The cumulative probability for token selection.
- top_k The number of highest probability vocabulary tokens to keep.
- frequency_penalty Frequency penalty for generated tokens.
- presence_penalty Presence penalty for generated tokens.
- seed Seed for random number generator.
- max_tokens Maximum number of tokens to generate.
- log_probs Include the log probabilities in the output?
- stop_sequences A character vector of tokens to stop generation on.
- ... Additional named parameters to send to the provider.

Provider	<i>A chatbot provider</i>
----------	---------------------------

Description

A Provider captures the details of one chatbot service/API. This captures how the API works, not the details of the underlying large language model. Different providers might offer the same (open source) model behind a different API.

Usage

```
Provider(  
  name = stop("Required"),  
  model = stop("Required"),  
  base_url = stop("Required"),
```

```
    params = list(),
    extra_args = list(),
    extra_headers = character(0)
)
```

Arguments

name	Name of the provider.
model	Name of the model.
base_url	The base URL for the API.
params	A list of standard parameters created by <code>params()</code> .
extra_args	Arbitrary extra arguments to be included in the request body.
extra_headers	Arbitrary extra headers to be added to the request.

Details

To add support for a new backend, you will need to subclass `Provider` (adding any additional fields that your provider needs) and then implement the various generics that control the behavior of each provider.

Value

An S7 Provider object.

Examples

```
Provider(
  name = "CoolModels",
  model = "my_model",
  base_url = "https://cool-models.com"
)
```

token_usage	<i>Report on token usage in the current session</i>
-------------	---

Description

Call this function to find out the cumulative number of tokens that you have sent and recieved in the current session. The price will be shown if known.

Usage

```
token_usage()
```

Value

A data frame

Examples

```
token_usage()
```

tool	<i>Define a tool</i>
------	----------------------

Description

Annotate a function for use in tool calls, by providing a name, description, and type definition for the arguments.

Learn more in `vignette("tool-calling")`.

Usage

```
tool(  
  fun,  
  description,  
  ...,  
  arguments = list(),  
  name = NULL,  
  convert = TRUE,  
  annotations = list(),  
  .name = deprecated(),  
  .description = deprecated(),  
  .convert = deprecated(),  
  .annotations = deprecated()  
)
```

Arguments

fun	The function to be invoked when the tool is called. The return value of the function is sent back to the chatbot. Expert users can customize the tool result by returning a ContentToolResult object.
description	A detailed description of what the function does. Generally, the more information that you can provide here, the better.
...	[Deprecated] Use arguments instead.
arguments	A named list that defines the arguments accepted by the function. Each element should be created by a type_*() function (or NULL if you don't want the LLM to use that argument).
name	The name of the function. This can be omitted if fun is an existing function (i.e. not defined inline).
convert	Should JSON inputs be automatically convert to their R data type equivalents? Defaults to TRUE.

annotations Additional properties that describe the tool and its behavior. Usually created by `tool_annotations()`, where you can find a description of the annotation properties recommended by the **Model Context Protocol**.

`.name`, `.description`, `.convert`, `.annotations`

[Deprecated] Please switch to the non-prefixed equivalents.

Value

An S7 ToolDef object.

ellmer 0.3.0

In ellmer 0.3.0, the definition of the `tool()` function changed quite a bit. To make it easier to update old versions, you can use an LLM with the following system prompt

Help the user convert an ellmer 0.2.0 and earlier tool definition into a ellmer 0.3.0 tool definition. Here's what changed:

- * All arguments, apart from the first, should be named, and the argument names no longer use ``.`` prefixes. The argument order should be function, name (as a string), description, then arguments, then anything
- * Previously ``arguments`` was passed as ``...``, so all type specifications should now be moved into a named list and passed to the ``arguments`` argument. It can be omitted if the function has no arguments.

```
```R
old
tool(
 add,
 "Add two numbers together"
 x = type_number(),
 y = type_number()
)

new
tool(
 add,
 name = "add",
 description = "Add two numbers together",
 arguments = list(
 x = type_number(),
 y = type_number()
)
)
```
```

Don't respond; just let the user provide function calls to convert.

See Also

Other tool calling helpers: `tool_annotations()`, `tool_reject()`

Examples

```
# First define the metadata that the model uses to figure out when to
# call the tool
tool_rnorm <- tool(
  rnorm,
  description = "Draw numbers from a random normal distribution",
  arguments = list(
    n = type_integer("The number of observations. Must be a positive integer."),
    mean = type_number("The mean value of the distribution."),
    sd = type_number("The standard deviation of the distribution. Must be a non-negative number.")
  )
)
tool_rnorm(n = 5, mean = 0, sd = 1)

chat <- chat_openai()
# Then register it
chat$register_tool(tool_rnorm)

# Then ask a question that needs it.
chat$chat("Give me five numbers from a random normal distribution.")

# Look at the chat history to see how tool calling works:
chat
# Assistant sends a tool request which is evaluated locally and
# results are sent back in a tool result.
```

tool_annotations

Tool annotations

Description

Tool annotations are additional properties that, when passed to the `.annotations` argument of `tool()`, provide additional information about the tool and its behavior. This information can be used for display to users, for example in a Shiny app or another user interface.

The annotations in `tool_annotations()` are drawn from the **Model Context Protocol** and are considered *hints*. Tool authors should use these annotations to communicate tool properties, but users should note that these annotations are not guaranteed.

Usage

```
tool_annotations(
  title = NULL,
  read_only_hint = NULL,
  open_world_hint = NULL,
  idempotent_hint = NULL,
  destructive_hint = NULL,
  ...
)
```

Arguments

| | |
|-------------------------------|---|
| <code>title</code> | A human-readable title for the tool. |
| <code>read_only_hint</code> | If TRUE, the tool does not modify its environment. |
| <code>open_world_hint</code> | If TRUE, the tool may interact with an "open world" of external entities. If FALSE, the tool's domain of interaction is closed. For example, the world of a web search tool is open, but the world of a memory tool is not. |
| <code>idempotent_hint</code> | If TRUE, calling the tool repeatedly with the same arguments will have no additional effect on its environment. (Only meaningful when <code>read_only_hint</code> is FALSE.) |
| <code>destructive_hint</code> | If TRUE, the tool may perform destructive updates to its environment, otherwise it only performs additive updates. (Only meaningful when <code>read_only_hint</code> is FALSE.) |
| <code>...</code> | Additional named parameters to include in the tool annotations. |

Value

A list of tool annotations.

See Also

Other tool calling helpers: [tool\(\)](#), [tool_reject\(\)](#)

Examples

```
# See ?tool() for a full example using this function.
# We're creating a tool around R's `rnorm()` function to allow the chatbot to
# generate random numbers from a normal distribution.
tool_rnorm <- tool(
  rnorm,
  # Describe the tool function to the LLM
  .description = "Drawn numbers from a random normal distribution",
  # Describe the parameters used by the tool function
  n = type_integer("The number of observations. Must be a positive integer."),
  mean = type_number("The mean value of the distribution."),
  sd = type_number("The standard deviation of the distribution. Must be a non-negative number."),
```

```
# Tool annotations optionally provide additional context to the LLM
.annotations = tool_annotations(
  title = "Draw Random Normal Numbers",
  read_only_hint = TRUE, # the tool does not modify any state
  open_world_hint = FALSE # the tool does not interact with the outside world
)
)
```

tool_reject

Reject a tool call

Description

Throws an error to reject a tool call. `tool_reject()` can be used within the tool function to indicate that the tool call should not be processed. `tool_reject()` can also be called in an `Chat$on_tool_request()` callback. When used in the callback, the tool call is rejected before the tool function is invoked.

Here's an example where `utils::askYesNo()` is used to ask the user for permission before accessing their current working directory. This happens directly in the tool function and is appropriate when you write the tool definition and know exactly how it will be called.

```
chat <- chat_openai(model = "gpt-4.1-nano")

list_files <- function() {
  allow_read <- utils::askYesNo(
    "Would you like to allow access to your current directory?"
  )
  if (isTRUE(allow_read)) {
    dir(pattern = "[.](r|R|csv)$")
  } else {
    tool_reject()
  }
}

chat$register_tool(tool(
  list_files,
  "List files in the user's current directory"
))

chat$chat("What files are available in my current directory?")
#> [tool call] list_files()
#> Would you like to allow access to your current directory? (Yes/no/cancel) no
#> #> Error: Tool call rejected. The user has chosen to disallow the tool #' call.
#> It seems I am unable to access the files in your current directory right now.
#> If you can tell me what specific files you're looking for or if you can #' provide
#> the list, I can assist you further.
```

```

chat$chat("Try again.")
#> [tool call] list_files()
#> Would you like to allow access to your current directory? (Yes/no/cancel) yes
#> #> app.R
#> #> data.csv
#> The files available in your current directory are "app.R" and "data.csv".

```

You can achieve a similar experience with tools written by others by using a `tool_request` callback. In the next example, imagine the tool is provided by a third-party package. This example implements a simple menu to ask the user for consent before running *any* tool.

```

packaged_list_files_tool <- tool(
  function() dir(pattern = "[.](r|R|csv)$"),
  "List files in the user's current directory"
)

chat <- chat_openai(model = "gpt-4.1-nano")
chat$register_tool(packaged_list_files_tool)

always_allowed <- c()

# ContentToolRequest
chat$on_tool_request(function(request) {
  if (request@name %in% always_allowed) return()

  answer <- utils::menu(
    title = sprintf("Allow tool `%s()` to run?", request@name),
    choices = c("Always", "Once", "No"),
    graphics = FALSE
  )

  if (answer == 1) {
    always_allowed <- append(always_allowed, request@name)
  } else if (answer %in% c(0, 3)) {
    tool_reject()
  }
})

# Try choosing different answers to the menu each time
chat$chat("What files are available in my current directory?")
chat$chat("How about now?")
chat$chat("And again now?")

```

Usage

```

tool_reject(reason = "The user has chosen to disallow the tool call.")

```

Arguments

reason A character string describing the reason for rejecting the tool call.

Value

Throws an error of class `ellmer_tool_reject` with the provided reason.

See Also

Other tool calling helpers: [tool\(\)](#), [tool_annotations\(\)](#)

| | |
|------|---------------------------------|
| Turn | <i>A user or assistant turn</i> |
|------|---------------------------------|

Description

Every conversation with a chatbot consists of pairs of user and assistant turns, corresponding to an HTTP request and response. These turns are represented by the `Turn` object, which contains a list of [Contents](#) representing the individual messages within the turn. These might be text, images, tool requests (assistant only), or tool responses (user only).

Note that a call to `$chat()` and related functions may result in multiple user-assistant turn cycles. For example, if you have registered tools, `ellmer` will automatically handle the tool calling loop, which may result in any number of additional cycles. Learn more about tool calling in `vignette("tool-calling")`.

Usage

```
Turn(role, contents = list(), json = list(), tokens = c(0, 0, 0))
```

Arguments

role Either "user", "assistant", or "system".

contents A list of [Content](#) objects.

json The serialized JSON corresponding to the underlying data of the turns. Currently only provided for assistant.

 This is useful if there's information returned by the provider that `ellmer` doesn't otherwise expose.

tokens A numeric vector of length 2 representing the number of input and output tokens (respectively) used in this turn. Currently only recorded for assistant turns.

Value

An S7 `Turn` object

Examples

```
Turn(role = "user", contents = list(ContentText("Hello, world!")))
```

Type

Type definitions for function calling and structured data extraction.

Description

These S7 classes are provided for use by package developers who are extending ellmer. In every day use, use `type_boolean()` and friends.

Usage

```
TypeBasic(description = NULL, required = TRUE, type = stop("Required"))
```

```
TypeEnum(description = NULL, required = TRUE, values = character(0))
```

```
TypeArray(description = NULL, required = TRUE, items = Type())
```

```
TypeJsonSchema(description = NULL, required = TRUE, json = list())
```

```
TypeObject(
  description = NULL,
  required = TRUE,
  properties = list(),
  additional_properties = FALSE
)
```

Arguments

| | |
|-------------|--|
| description | The purpose of the component. This is used by the LLM to determine what values to pass to the tool or what values to extract in the structured data, so the more detail that you can provide here, the better. |
| required | Is the component or argument required?

In type descriptions for structured data, if <code>required = FALSE</code> and the component does not exist in the data, the LLM may hallucinate a value. Only applies when the element is nested inside of a <code>type_object()</code> .

In tool definitions, <code>required = TRUE</code> signals that the LLM should always provide a value. Arguments with <code>required = FALSE</code> should have a default value in the tool function's definition. If the LLM does not provide a value, the default value will be used. |
| type | Basic type name. Must be one of <code>boolean</code> , <code>integer</code> , <code>number</code> , or <code>string</code> . |
| values | Character vector of permitted values. |
| items | The type of the array items. Can be created by any of the <code>type_</code> function. |
| json | A JSON schema object as a list. |
| properties | Named list of properties stored inside the object. Each element should be an S7 Type object. |

additional_properties
Can the object have arbitrary additional properties that are not explicitly listed?
Only supported by Claude.

Value

S7 objects inheriting from Type

Examples

```
TypeBasic(type = "boolean")
TypeArray(items = TypeBasic(type = "boolean"))
```

| | |
|--------------|---------------------|
| type_boolean | Type specifications |
|--------------|---------------------|

Description

These functions specify object types in a way that chatbots understand and are used for tool calling and structured data extraction. Their names are based on the **JSON schema**, which is what the APIs expect behind the scenes. The translation from R concepts to these types is fairly straightforward.

- type_boolean(), type_integer(), type_number(), and type_string() each represent scalars. These are equivalent to length-1 logical, integer, double, and character vectors (respectively).
- type_enum() is equivalent to a length-1 factor; it is a string that can only take the specified values.
- type_array() is equivalent to a vector in R. You can use it to represent an atomic vector: e.g. type_array(type_boolean()) is equivalent to a logical vector and type_array(type_string()) is equivalent to a character vector). You can also use it to represent a list of more complicated types where every element is the same type (R has no base equivalent to this), e.g. type_array(type_array(type_string())) represents a list of character vectors.
- type_object() is equivalent to a named list in R, but where every element must have the specified type. For example, type_object(a = type_string(), b = type_array(type_integer())) is equivalent to a list with an element called a that is a string and an element called b that is an integer vector.
- type_from_schema() allows you to specify the full schema that you want to get back from the LLM as a JSON schema. This is useful if you have a pre-defined schema that you want to use directly without manually creating the type using the type_*() functions. You can point to a file with the path argument or provide a JSON string with text. The schema must be a valid JSON schema object.

Usage

```

type_boolean(description = NULL, required = TRUE)

type_integer(description = NULL, required = TRUE)

type_number(description = NULL, required = TRUE)

type_string(description = NULL, required = TRUE)

type_enum(values, description = NULL, required = TRUE)

type_array(items, description = NULL, required = TRUE)

type_object(
  .description = NULL,
  ...,
  .required = TRUE,
  .additional_properties = FALSE
)

type_from_schema(text, path)

```

Arguments

| | |
|---------------------------|--|
| description, .description | The purpose of the component. This is used by the LLM to determine what values to pass to the tool or what values to extract in the structured data, so the more detail that you can provide here, the better. |
| required, .required | Is the component or argument required?

In type descriptions for structured data, if <code>required = FALSE</code> and the component does not exist in the data, the LLM may hallucinate a value. Only applies when the element is nested inside of a <code>type_object()</code> .

In tool definitions, <code>required = TRUE</code> signals that the LLM should always provide a value. Arguments with <code>required = FALSE</code> should have a default value in the tool function's definition. If the LLM does not provide a value, the default value will be used. |
| values | Character vector of permitted values. |
| items | The type of the array items. Can be created by any of the <code>type_</code> function. |
| ... | Name-type pairs defining the components that the object must possess. |
| .additional_properties | Can the object have arbitrary additional properties that are not explicitly listed?
Only supported by Claude. |
| text | A JSON string. |
| path | A file path to a JSON file. |

Examples

```
# An integer vector
type_array(type_integer())

# The closest equivalent to a data frame is an array of objects
type_array(type_object(
  x = type_boolean(),
  y = type_string(),
  z = type_number()
))

# There's no specific type for dates, but you use a string with the
# requested format in the description (it's not guaranteed that you'll
# get this format back, but you should most of the time)
type_string("The creation date, in YYYY-MM-DD format.")
type_string("The update date, in dd/mm/yyyy format.")
```

Index

* chatbots

- chat_anthropic, 12
- chat_aws_bedrock, 14
- chat_azure_openai, 15
- chat_cloudflare, 17
- chat_databricks, 20
- chat_deepseek, 22
- chat_github, 23
- chat_google_gemini, 25
- chat_groq, 27
- chat_huggingface, 29
- chat_mistral, 31
- chat_ollama, 32
- chat_openai, 34
- chat_openrouter, 35
- chat_perplexity, 37
- chat_portkey, 38

* tool calling helpers

- tool, 56
- tool_annotations, 58
- tool_reject, 60

- batch_chat, 3
- batch_chat(), 52
- batch_chat_completed(batch_chat), 3
- batch_chat_structured(batch_chat), 3
- batch_chat_text(batch_chat), 3

- Chat, 4, 5, 13, 15, 17, 18, 20, 21, 23, 25, 27, 28, 30, 32, 33, 35, 36, 38, 40–42, 45, 53

- chat, 11
- chat_anthropic, 12, 15, 17, 18, 22, 23, 25, 27, 29, 30, 32, 33, 35, 36, 38, 40
- chat_anthropic(), 3, 52
- chat_aws_bedrock, 13, 14, 17, 18, 22, 23, 25, 27, 29, 30, 32, 33, 35, 36, 38, 40
- chat_azure_openai, 13, 15, 15, 18, 22, 23, 25, 27, 29, 30, 32, 33, 35, 36, 38, 40

- chat_cloudflare, 13, 15, 17, 17, 22, 23, 25, 27, 29, 30, 32, 33, 35, 36, 38, 40
- chat_cortex_analyst, 19
- chat_cortex_analyst(), 40
- chat_databricks, 13, 15, 17, 18, 20, 23, 25, 27, 29, 30, 32, 33, 35, 36, 38, 40
- chat_deepseek, 13, 15, 17, 18, 22, 22, 25, 27, 29, 30, 32, 33, 35, 36, 38, 40
- chat_github, 13, 15, 17, 18, 22, 23, 23, 27, 29, 30, 32, 33, 35, 36, 38, 40
- chat_google_gemini, 13, 15, 17, 18, 22, 23, 25, 25, 29, 30, 32, 33, 35, 36, 38, 40
- chat_google_vertex
(chat_google_gemini), 25
- chat_groq, 13, 15, 17, 18, 22, 23, 25, 27, 27, 30, 32, 33, 35, 36, 38, 40
- chat_huggingface, 13, 15, 17, 18, 22, 23, 25, 27, 29, 29, 32, 33, 35, 36, 38, 40
- chat_mistral, 13, 15, 17, 18, 22, 23, 25, 27, 29, 30, 31, 33, 35, 36, 38, 40
- chat_ollama, 13, 15, 17, 18, 22, 23, 25, 27, 29, 30, 32, 32, 35, 36, 38, 40
- chat_openai, 13, 15, 17, 18, 22, 23, 25, 27, 29, 30, 32, 33, 34, 36, 38, 40
- chat_openai(), 3, 5, 24, 27, 29, 32, 37, 48, 51, 52
- chat_openrouter, 13, 15, 17, 18, 22, 23, 25, 27, 29, 30, 32, 33, 35, 35, 38, 40
- chat_perplexity, 13, 15, 17, 18, 22, 23, 25, 27, 29, 30, 32, 33, 35, 36, 37, 40
- chat_portkey, 13, 15, 17, 18, 22, 23, 25, 27, 29, 30, 32, 33, 35, 36, 38, 38
- chat_snowflake, 40
- chat_snowflake(), 19
- chat_vllm, 41
- commonmark::markdown_html(), 45
- Content, 9, 43, 45, 62
- content_image_file(content_image_url), 46

content_image_file(), [8](#), [43](#)
 content_image_plot (content_image_url),
 [46](#)
 content_image_url, [46](#)
 content_image_url(), [8](#), [43](#)
 content_pdf_file, [47](#)
 content_pdf_url (content_pdf_file), [47](#)
 ContentImage (Content), [43](#)
 ContentImageInline (Content), [43](#)
 ContentImageRemote (Content), [43](#)
 ContentPDF (Content), [43](#)
 contents_html (contents_text), [45](#)
 contents_markdown (contents_text), [45](#)
 contents_text, [45](#)
 ContentText, [45](#)
 ContentText (Content), [43](#)
 ContentThinking (Content), [43](#)
 ContentToolRequest, [44](#)
 ContentToolRequest (Content), [43](#)
 ContentToolResult, [56](#)
 ContentToolResult (Content), [43](#)
 create_tool_def, [48](#)

 google_upload, [49](#)
 google_upload(), [25](#)

 http2::req_headers(), [16](#), [20](#), [41](#)

 interpolate, [50](#)
 interpolate(), [4](#), [52](#)
 interpolate_file (interpolate), [50](#)
 interpolate_package (interpolate), [50](#)

 live_browser (live_console), [51](#)
 live_console, [51](#)

 models_anthropic (chat_anthropic), [12](#)
 models_aws_bedrock (chat_aws_bedrock),
 [14](#)
 models_github (chat_github), [23](#)
 models_google_gemini
 (chat_google_gemini), [25](#)
 models_google_vertex
 (chat_google_gemini), [25](#)
 models_ollama (chat_ollama), [32](#)
 models_openai (chat_openai), [34](#)
 models_portkey (chat_portkey), [38](#)
 models_vllm (chat_vllm), [41](#)
 modifyList(), [13](#), [17](#), [18](#), [20](#), [21](#), [23](#), [24](#), [26](#),
 [28](#), [30](#), [31](#), [33](#), [35](#), [36](#), [38](#), [39](#), [41](#), [42](#)

 parallel_chat, [52](#)
 parallel_chat(), [3](#)
 parallel_chat_structured
 (parallel_chat), [52](#)
 parallel_chat_text (parallel_chat), [52](#)
 params, [53](#)
 params(), [12–14](#), [16](#), [18](#), [21](#), [23](#), [24](#), [26](#), [28](#),
 [30](#), [31](#), [33](#), [34](#), [36](#), [38](#), [39](#), [41](#), [42](#), [52](#),
 [55](#)
 Provider, [5](#), [43](#), [54](#)

 token_usage, [55](#)
 tool, [56](#), [59](#), [62](#)
 tool(), [10](#), [43](#), [48](#), [58](#)
 tool_annotations, [58](#), [58](#), [62](#)
 tool_annotations(), [57](#)
 tool_reject, [58](#), [59](#), [60](#)
 ToolDef (tool), [56](#)
 Turn, [5](#), [7](#), [45](#), [62](#)
 Type, [63](#)
 type_(), [4](#), [8](#), [52](#)
 type_*, [56](#)
 type_array (type_boolean), [64](#)
 type_boolean, [64](#)
 type_boolean(), [63](#)
 type_enum (type_boolean), [64](#)
 type_from_schema (type_boolean), [64](#)
 type_integer (type_boolean), [64](#)
 type_number (type_boolean), [64](#)
 type_object (type_boolean), [64](#)
 type_object(), [4](#), [53](#)
 type_string (type_boolean), [64](#)
 TypeArray (Type), [63](#)
 TypeBasic (Type), [63](#)
 TypeEnum (Type), [63](#)
 TypeJsonSchema (Type), [63](#)
 TypeObject (Type), [63](#)