

Package ‘hollr’

July 22, 2025

Type Package

Title Chat Completion and Text Annotation with Local and OpenAI Models

Version 1.0.0

Maintainer Jason Timm <JaTimm@salud.unm.edu>

Description Enables chat completion and text annotation with local and 'OpenAI' <<https://openai.com/>> language models, supporting batch processing, multiple annotators, and consistent output formats.

License MIT + file LICENSE

Encoding UTF-8

Depends R (>= 3.5)

Imports data.table, httr, jsonlite, pbapply, reticulate

RoxygenNote 7.3.1

URL <https://github.com/jaytimm/hollr>

BugReports <https://github.com/jaytimm/hollr/issues>

NeedsCompilation no

Author Jason Timm [aut, cre]

Repository CRAN

Date/Publication 2024-10-15 15:30:05 UTC

Contents

hollr	2
Index	4

hollr

LLM Completions

Description

This function generates text using the OpenAI API or a local model.

Usage

```
hollr(  
    id,  
    user_message = "",  
    annotators = 1,  
    model,  
    temperature = 1,  
    top_p = 1,  
    max_tokens = NULL,  
    max_length = 1024,  
    max_new_tokens = NULL,  
    system_message = "",  
    force_json = TRUE,  
    flatten_json = TRUE,  
    max_attempts = 10,  
    openai_api_key = Sys.getenv("OPENAI_API_KEY"),  
    openai_organization = NULL,  
    cores = 1,  
    batch_size = 1,  
    extract_json = TRUE,  
    verbose = TRUE  
)
```

Arguments

<code>id</code>	A unique identifier for the request.
<code>user_message</code>	The message provided by the user.
<code>annotators</code>	The number of annotators (default is 1).
<code>model</code>	The name of the model to use.
<code>temperature</code>	The temperature for the model's output (default is 1).
<code>top_p</code>	The top-p sampling value (default is 1).
<code>max_tokens</code>	The maximum number of tokens to generate (default is NULL).
<code>max_length</code>	The maximum length of the input prompt (default is 1024, for local models only).
<code>max_new_tokens</code>	The maximum number of new tokens to generate (default is NULL, for local models only).

system_message	The message provided by the system (default is "").
force_json	A logical indicating whether the output should be JSON (default is TRUE).
flatten_json	A logical indicating whether the output should be converted to DF (default is TRUE).
max_attempts	The maximum number of attempts to make for generating valid output (default is 10).
openai_api_key	The API key for the OpenAI API (default is retrieved from environment variables).
openai_organization	The organization ID for the OpenAI API (default is NULL).
cores	The number of cores to use for parallel processing (default is 1).
batch_size	The number of batch_size to process (default is 1, only for local models).
extract_json	A logical indicating whether to extract and clean JSON strings from the response (default is TRUE).
verbose	A logical indicating whether to show progress and other messages (default is TRUE).

Value

A data.table containing the generated text and metadata.

Examples

```
# Example using the OpenAI API
if (nzchar(Sys.getenv("OPENAI_API_KEY"))) {
  result <- hollr(id = "example_id",
                 user_message = "What is the capital of France?",
                 model = "gpt-3.5-turbo",
                 openai_api_key = Sys.getenv("OPENAI_API_KEY"))
  print(result)
} else {
  message("OpenAI API key is not available. Example not run.")
}
```

Index

hollr, [2](#)