

Package ‘chatRater’

August 18, 2025

Type Package
Title A Tool for Rating Text/Image Stimuli
Version 1.2.0
Date 2025-08-18
Maintainer Shiyang Zheng <Shiyang.Zheng@nottingham.ac.uk>
Description Evaluates stimuli using Large Language Models APIs with URL support.
License MIT + file LICENSE
Encoding UTF-8
Imports openai, tools
Suggests testthat (>= 3.0.0)
NeedsCompilation no
Author Shiyang Zheng [aut, cre]
RoxygenNote 7.3.2
Config/testthat/edition 3
Repository CRAN
Date/Publication 2025-08-18 19:50:26 UTC

Contents

generate_ratings	2
generate_ratings_for_all	3
Index	4

Description

Evaluates stimuli using Large Language Models APIs with URL support.

Usage

```
generate_ratings(  
    model = "gpt-4-turbo",  
    stim,  
    prompt = "You are an expert rater, limit your answer to numbers.",  
    question = "Please rate this:",  
    scale = "1-7",  
    top_p = 1,  
    temp = 0,  
    n_iterations = 1,  
    api_key = "",  
    debug = TRUE  
)
```

Arguments

model	LLM model name (default: 'gpt-4-turbo')
stim	Input stimulus (text string and image URL)
prompt	System instruction for the LLM
question	Specific rating question for the LLM
scale	Rating scale range (default: '1-7')
top_p	Top-p sampling parameter (default: 1)
temp	Temperature parameter (default: 0)
n_iterations	Number of rating iterations (default: 1)
api_key	OpenAI API key
debug	Debug mode flag (default: FALSE)

Value

A data frame containing ratings and metadata

generate_ratings_for_all
<i>Batch Rating Generator</i>

Description

Process multiple stimuli in sequence

Usage

```
generate_ratings_for_all(  
    model = "gpt-4-turbo",  
    stim_list,  
    prompt = "You are an expert rater, limit your answer to numbers.",  
    question = "Please rate this:",  
    scale = "1-7",  
    top_p = 1,  
    temp = 0,  
    n_iterations = 1,  
    api_key = "",  
    debug = TRUE  
)
```

Arguments

model	LLM model name (default: 'gpt-4-turbo')
stim_list	List of stimuli to process
prompt	System instruction for the LLM
question	Specific rating question for the LLM
scale	Rating scale range (default: '1-7')
top_p	Top-p sampling parameter (default: 1)
temp	Temperature parameter (default: 0)
n_iterations	Number of rating iterations (default: 1)
api_key	OpenAI API key
debug	Debug mode flag (default: FALSE)

Index

`generate_ratings`, [2](#)
`generate_ratings_for_all`, [3](#)