

```
search_rednote.py > ...
1 from selenium import webdriver
2 import subprocess
3 import time
4 import pandas as pd
5 from selenium.webdriver.chrome.options import Options
6 from selenium.webdriver.common.by import By
7 from selenium.webdriver.support.ui import WebDriverWait
8 from selenium.webdriver.support import expected_conditions as EC
9
10 def getElement(eles):
11     if eles and len(eles) != 0:
12         return eles[0].text
13     else:
14         return 'None'
15
16 def get_renote_posts(key, file_path):
17     browser_path = "C:\\Program Files (x86)\\Google\\Chrome\\Application\\chrome.exe"
18     subprocess.Popen([browser_path, '--remote-debugging-port=9222'])
19     time.sleep(5)
20
21     options = Options()
22     options.add_experimental_option("debuggerAddress", "localhost:9222")
23     driver = webdriver.Chrome(options=options)
24
25     driver.get('https://www.xiaohongshu.com/search_result?keyword='+key+'&source=web_explore_feed&type=51')
26     time.sleep(10) # Wait for the page to load
27
28     insert_Data = []
29     ids = []
30     y = 0
31     for i in range(20): # Loop for 20 pages
32         print(i)
33         try:
34             y += 600
35             driver.execute_script(f'window.scrollTo(0,{y})')
36             time.sleep(2.5)
37             video_list = driver.find_elements(By.CSS_SELECTOR, '.note-item')
38             if len(video_list) == 0:
39                 break
40             for video in video_list:
41                 if video.find_elements(By.CSS_SELECTOR, '.title'):
42                     url = video.find_element(By.CSS_SELECTOR, ".cover").get_attribute('href')
43                     title = video.find_element(By.CSS_SELECTOR, '.title').text
44                     name = video.find_element(By.CSS_SELECTOR, '.name').text
45                     count = video.find_element(By.CSS_SELECTOR, '.count').text
46                     shareurl = url.replace('pc_user', 'pc_share').replace('explore', 'discovery/item').replace('xsec_token=',
47                                                                                                     'source=webshare&xhsshare=pc_web&xsec_token=')
48
49                     if title in ids:
50                         continue
51                     print(title)
52                     ids.append(title)
53                     dict = {
54                         'Title': title,
55                         'Author': name,
56                         'Like Count': count,
57                         'URL': url,
58                         'Share URL': shareurl,
59                     }
60                     insert_Data.append(dict)
61         except Exception as e:
62             print(e)
63
64     df = pd.DataFrame(insert_Data)
65     df.to_excel(file_path, index=False, engine='openpyxl')
66
67 # Call the function with the keyword
68 key = 'Macau Street Market'
69 video_urls = get_renote_posts(key, 'Macau_Street_Market.xlsx')
70
```

File

Edit

Selection

View

Go

Macau Street Market

generate_wordcloud.py

wordcloud_english.py

translate.py

search_rednote.py

comments_collected.py

comments_collected.py > ...

1

import

csv

2

import

io

3

import

os

4

import

pandas

as

pd

5

import

random

6

import

re

7

import

subprocess

8

import

sys

9

import

time

10

from

datetime

import

datetime

,

timedelta

11

from

selenium

import

webdriver

12

from

selenium.common.exceptions

import

NoSuchElementException

13

from

selenium.webdriver.chrome.options

import

Options

14

from

selenium.webdriver.chrome.service

import

Service

15

from

selenium.webdriver.common.by

import

By

16

from

selenium.webdriver.support

import

expected_conditions

as

EC

17

from

selenium.webdriver.support.ui

import

WebDriverWait

18

19

def

getElement

(eles):

20

if

eles

and

len

(eles)

!=

0

:

21

return

eles

[0]

.text

22

else:

23

return

'N/A'

24

25

def

scrape

(url,

pid):

26

driver

.

get

(url)

27

time

.

sleep

(5)

28

WebDriverWait

(driver,

20)

.until

(

29

EC.visibility_of_element_located

((By.CSS_SELECTOR,

'

.list-container

'))

30

)

31

div_element

=

driver

.

find_element

(By.CSS_SELECTOR,

'

.note-scroller

'

32

y

=

0

33

for

i

in

range

(200):

34

try:

35

print

('Scrolling '

+

str

(i)

)

36

y

+=

2000

37

driver

.

execute_script

("arguments[0].scrollTop = " + str(y) + ";"

,

div_element

38

if

driver

.

find_elements

(By.CSS_SELECTOR,

'

.end-container

'

):

39

break

40

except:

41

print

('Jumping')

42

time

.

sleep

(3)

43

44

45

#

Expand

comments

46

clicklist

=

driver

.

find_elements

(By.CSS_SELECTOR,

'

.show-more

'

47

for

clickitem

in

clicklist:

48

try:

49

clickitem

.

click

()

50

time

.

sleep

(0.4)

51

except:

52

print

('Jumping')

53

54

#

Click

more

comments

55

for

j

in

range

(3000):

56

try:

57

if

driver

.

find_elements

(By.CSS_SELECTOR,

'

.show-more

'

):

58

driver

.

find_element

(By.CSS_SELECTOR,

'

.show-more

'

.

click

()

59

time

.

sleep

(0.2)

60

else:

61

break

62

except:

63

print

('Jumping')

64

65

video_list

=

driver

.

find_elements

(By.CSS_SELECTOR,

'

.comment-item

'

66

for

video

in

video_list:

67

author_name

=

getElement

(video

.

find_elements

(By.CSS_SELECTOR,

'

.author

'

a'))

68

content

=

getElement

(video

.

find_elements

(By.CSS_SELECTOR,

'

.note-text

'

)

69

comment_date

=

getElement

(

70

video

.

find_elements

(By.CSS_SELECTOR,

'

.right

>

div.info

>

div.date

>

span:nth-child(1)'

'

71

location

=

getElement

(video

.

find_elements

(By.CSS_SELECTOR,

'

.right

>

div.info

>

div.date

>

.location'

'

)

72

like_count

=

getElement

(video

.

find_elements

(By.CSS_SELECTOR,

'

.like

.

count'

'

)

73

reply_count

=

getElement

(video

.

find_elements

(By.CSS_SELECTOR,

'

.reply

.

count'

'

)

74

print

(str

(author_name)

)

75

dict

=

{

76

'Post ID':

str

(pid)

,

77

'Author Name':

str

(author_name)

,

78

'Comment Content':

str

(content)

,

79

'Comment Date':

str

(comment_date)

,

80

'Comment Location':

str

(location)

,

81

'Comment Likes':

str

(like_count)

,

82

'Comment Replies':

str

(reply_count)

,

83

}

84

#

writer.writerow

(dict)

85

data_list_insert.append

(dict)

86

print

(len

(data_list_insert)

)

87

88

if

__name__

==

"__main__"

:

89

#

Create

a

new

CSV

file

and

write

data

90

csv_filename

=

'data.csv'

91

browser_path

=

"C:\\Program Files (x86)\\Google\\Chrome\\Application\\chrome.exe"

92

subprocess.Popen

([browser_path,

'--remote-debugging-port=9222'])

93

time

.

sleep

(5)

94

print

('End')

95

options

=

Options

()

96

options.add_experimental_option

("debuggerAddress",

"localhost:9222")

97

driver

=

webdriver.Chrome

(options=options)

98

data_list_insert

=

[]

99

#

todo:

Replace

with

the

actual

URL

of

the

post

100

df

=

pd.read_excel

('Macau Street Market.xlsx',

sheet_name='Sheet1',

engine='openpyxl')

101

data_list

=

df.to_dict

(orient='records')

102

try:

103

for

pitem

in

data_list:

104

url

=

pitem

['url']

105

pid

=

pitem

['id']

106

try:

107

scrape

(url,

pid)

108

except

Exception

as

e:

109

print

(e)

110

except:

111

print

('Failed')

112

df

=

pd.DataFrame

(data_list_insert)

113

df.to_excel

('Macau Street Market Comments.xlsx',

index=False,

engine='xlsxwriter')

generate_wordcloud.py × wordcloud_english.py translate.py search_rednote.py comments_collected.py

generate_wordcloud.py > ...

```

1  # -*- coding: utf-8 -*-
2  import pandas as pd
3  from wordcloud import WordCloud
4  from PIL import Image
5  import os
6  import matplotlib
7  import jieba
8
9  matplotlib.rcParams['font.sans-serif'] = ['Microsoft YaHei']
10 matplotlib.rcParams['axes.unicode_minus'] = False
11 sheet_name = 'Sheet1'
12
13 column_name = 'Comment Content' # Replace with the actual column name
14 df = pd.read_excel('Comments.xlsx', sheet_name=sheet_name, usecols=[column_name], engine='openpyxl')
15
16 # Convert data to dictionary format
17 data_list = df.to_dict(orient='records')
18
19 # Set the folder to save the word cloud
20 wordcloud_save_folder = 'Macau Street Market'
21 if not os.path.exists(wordcloud_save_folder):
22     os.makedirs(wordcloud_save_folder)
23
24 # Combine all comments into a single text
25 text = ''
26 for item in data_list:
27     if str(item['Comment Content']) != 'nan' and len(item['Comment Content']) >= 3:
28         text += '\n' + str(item['Comment Content'])
29
30 # Use jieba for Chinese word segmentation
31 words = jieba.lcut(text)
32
33 # Load stopwords list
34 def add_stop_words(list):
35     stop_words = set()
36     for item in list:
37         stop_words.add(item.strip())
38     return stop_words
39
40 stopword_list = open('stopword.txt', encoding='utf-8').readlines()
41 stopwords = add_stop_words(stopword_list)
42
43 # Filter out stopwords and single character words
44 filtered_words = [word for word in words if word not in stopwords and len(word) > 1]
45 text = ' '.join(filtered_words)
46
47 # Generate word cloud
48 wordcloud = WordCloud(width=1000, height=700, background_color='white', font_path='C:\\Windows\\Fonts\\msyh.ttc').generate(text)
49
50 # Save the word cloud as an image file
51 image_path = os.path.join(wordcloud_save_folder, f'Word Frequency.png')
52 wordcloud_image = Image.fromarray(wordcloud.to_array())
53 wordcloud_image.save(image_path)
54
55 print(f"Word cloud has been saved to {image_path}.")
56

```

translate.py > ...

```

1  import pandas as pd
2  from deep_translator import GoogleTranslator
3
4  # Read Excel file, keep all columns
5  sheet_name = 'Sheet2' # Replace with the actual sheet name
6  column_name = 'Word' # Replace with the actual column name
7
8  # Read data
9  df = pd.read_excel('Word Frequency.xlsx', sheet_name=sheet_name, engine='openpyxl')
10
11 # Initialize Google Translator
12 translator = GoogleTranslator(source='zh-CN', target='en')
13
14 # Translate each comment
15 translated_comments = []
16 for comment in df[column_name]:
17     try:
18         if pd.isna(comment) or not comment.strip(): # If comment is empty or NaN
19             translated_comments.append('No comment') # Or you can choose another placeholder
20         else:
21             # Translate the comment content
22             translated_comment = translator.translate(comment)
23             translated_comments.append(translated_comment)
24     except Exception as e:
25         print(f"Error translating: {comment} - {e}")
26         translated_comments.append('Translation Error') # If translation fails, add error message
27
28 # Add translated results to a new column
29 df['Translated Comments'] = translated_comments
30
31 # Save the updated data back to the original Excel file's Sheet2, replacing the old data
32 with pd.ExcelWriter('Word Frequency.xlsx', engine='openpyxl', mode='a', if_sheet_exists='replace') as writer:
33     df.to_excel(writer, sheet_name=sheet_name, index=False)
34
35 print("Translation complete, results saved to 'Word Frequency.xlsx' Sheet2.")
36

```

I

wordcloud_english.py > ...

```

1  # -*- coding: utf-8 -*-
2  import pandas as pd
3  from wordcloud import WordCloud
4  from PIL import Image
5  import os
6  import matplotlib
7
8  matplotlib.rcParams['font.sans-serif'] = ['Microsoft YaHei']
9  matplotlib.rcParams['axes.unicode_minus'] = False
10 sheet_name = 'Sheet2'
11 column_name = 'Translated Comments'
12 count_column_name = 'Count'
13
14 df = pd.read_excel('Word Frequency.xlsx', sheet_name=sheet_name,
15                  usecols=[column_name, count_column_name], engine='openpyxl')
16
17 # Create a word frequency dictionary
18 word_freq = {}
19 for index, row in df.iterrows():
20     comment = str(row[column_name])
21     count = row[count_column_name]
22     if len(comment) >= 3 and comment != 'nan':
23         if comment in word_freq:
24             word_freq[comment] += count
25         else:
26             word_freq[comment] = count
27
28 wordcloud_save_folder = 'Macau Street Market'
29 if not os.path.exists(wordcloud_save_folder):
30     os.makedirs(wordcloud_save_folder)
31
32 # Generate the word cloud
33 wordcloud = WordCloud(width=1000, height=700, background_color='white',
34                      font_path='C:\\Windows\\Fonts\\msyh.ttc').generate_from_frequencies(word_freq)
35
36 image_path = os.path.join(wordcloud_save_folder, f'wordcloud_translated.png')
37 wordcloud_image = Image.fromarray(wordcloud.to_array())
38 wordcloud_image.save(image_path)
39 print(f"The word cloud has been saved to {image_path}.")
40

```

I