

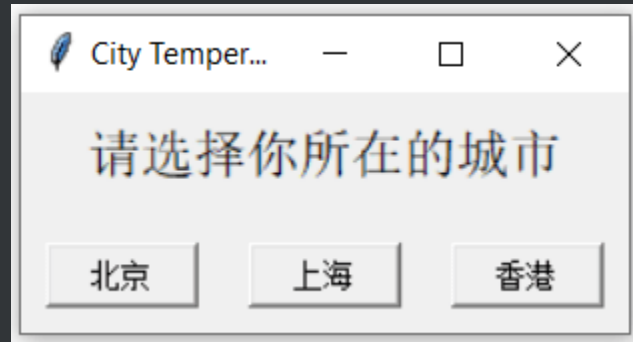
API



Demonstration and Vote!

This will be included as part of your in-class participation.

I have developed an APP. Let's try it.



In your opinion, how does it get the temperature data?

<https://www.youtube.com/embed/s7wmiS2mSXY?enablejsapi=1>

Let me show you the basics.

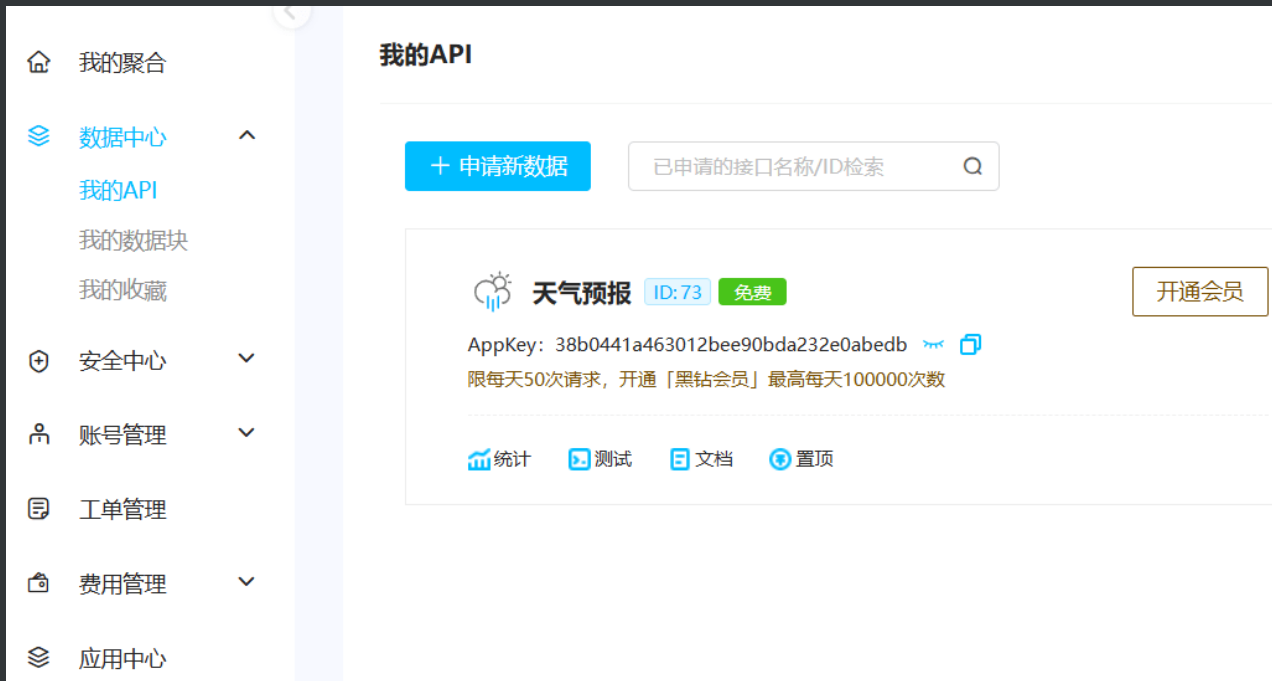
```
1 import requests
2 url = "http://apis.juhe.cn/simpleWeather/query"
3 params = {
4     "city": "北京",
5     "key": "your API key",
6 }
7 response = requests.get(url, params = params)
8 data = response.json()
9 temperature = data['result']['realtime']['temperature']
10 print(temperature)
```

You need to apply for your API key (password)

I am not sharing my key with you because a
free key can only be used 50 times a day...

How to get the key?

1. Sign up an account [here](#) and activate it.
2. Sign up for an API key.



我的API

+ 申请新数据

已申请的接口名称/ID检索

天气预报 ID: 73 免费 开通会员

AppKey: 38b0441a463012bee90bda232e0abedb

限每天50次请求, 开通「黑钻会员」最高每天100000次数

统计 测试 文档 置顶

Exercise

Apply for your own API key and develop your own weather APP!

```

1 import tkinter as tk
2 import requests
3 def get_temperature(city):
4     url = "http://apis.juhe.cn/simpleWeather/query"
5     params = { "city": city, "key": "your key" }
6     response = requests.get(url, params=params)
7     data = response.json()
8     temperature = data['result']['realtime']['temperature']
9     spaces_to_add = 3 - len(temperature)
10    output_string = " " * spaces_to_add + temperature
11    return output_string
12 def update_display(city):
13     temperature = get_temperature(city)
14     display_label.config(text="{}当前温度为: {}摄氏度".format(city, temperature))
15 root = tk.Tk()
16 root.title("City Temperature App")
17 display_label = tk.Label(root, text="    请选择你所在的城市    ", font=("Arial", 16))
18 display_label.grid(row=0, column=0, columnspan=3, pady=10)
19 cities = ["北京", "上海", "香港"]
20 for index, city in enumerate(cities):
21     city_button = tk.Button(root, text=city, command=lambda c=city: update_display(c))
22     city_button.grid(row=1, column=index, padx=10, pady=10, sticky="ew")
23 for i in range(3):
24     root.grid_columnconfigure(i, weight=1)
25
26 root.grid_rowconfigure(1, weight=1)
27 root.mainloop()

```

Let's try the following API for Chinese/English translation

Translation

```
1 import requests
2 url = "https://api.oicn.cn/api/fanyi"
3 params = {
4     "text": "欢迎来到香港大学",
5 }
6 response = requests.get(url, params = params)
7 data = response.json()
8 result = data['data']['result']
9 print(result)
```

Exercise: Design an APP so that users can input Chinese/English and the APP translates it into the other language

My APP created by GPT

```
1 import tkinter as tk
2 import requests
3 def translate(text):
4     url = "https://api.oicn.cn/api/fanyi"
5     params = {"text": text }
6     response = requests.get(url, params = params)
7     data = response.json()
8     result = data['data']['result']
9     return result
10
11 def translate_text():
12     chinese_text = input_text.get("1.0", "end-1c") # Get text from input text box
13     english_text = translate(chinese_text)
14     input_text.delete("1.0", "end") # Clear the input text
15     input_text.insert("1.0", english_text) # Display translated text in the same text box
16
17 root = tk.Tk()
18 root.title("Chinese to English Translator")
19 input_text = tk.Text(root, height=10, width=50)
20 input_text.pack(pady=10)
21 translate_button = tk.Button(root, text="Translate", command=translate_text)
22 translate_button.pack(pady=10)
23 root.mainloop()
```

AviationStack is a platform which provides an API to track the status of flights. Its website is here:

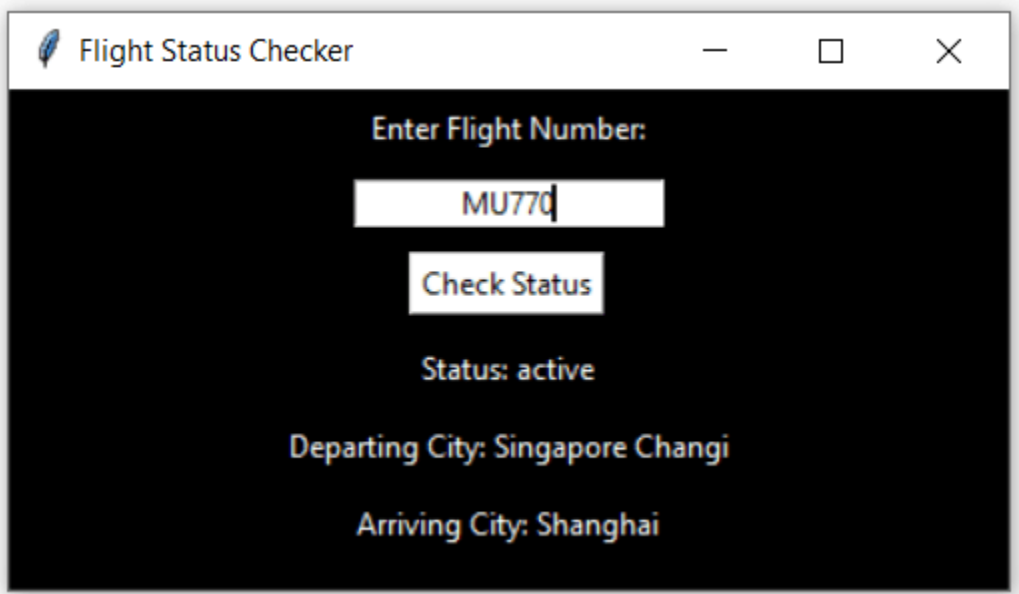
<https://aviationstack.com/>

You can find its API documentation [here](#).

Exercise

- Acquire a free API Key from the platform
- Develop an APP, the APP takes flight number as input from the user, and displays its flight status, departing airport and arriving airport on the screen.
- Try to do everything yourself (with help from GPT)!

This is my APP.



Flight Status Checker

Enter Flight Number:

MU770

Check Status

Status: active

Departing City: Singapore Changi

Arriving City: Shanghai

This is “my” AI (API)

 AI Question Answering

Enter your question:

Do you like HKU?

Ask

As an AI, I don't have personal feelings or preferences, so I don't "like" or "dislike" anything, including institutions like the University of Hong Kong (HKU). However, I can provide information about HKU! It is one of the most prestigious universities in Asia, known for its strong academic programs, research contributions, and vibrant campus life. If you have specific questions about HKU, feel free to ask!

Exercise

- Acquire an API Key from [DeepSeek](https://platform.deepseek.com/) (<https://platform.deepseek.com/>)
- This is not free, but very cheap. You need to top up \$2 or RMB 10 to use the API.
- Then, develop an APP which takes a question from user, sends the question to the API, and displays the answer from API to the user.

Top up

USD

CNY

Amount

\$2

\$5

\$10

\$20

\$50

\$100

\$500

Custom

Pricing

Total

\$2.12

Total excluding tax \$2.00 + VAT(6%) \$0.12 ?

Off-Peak Discounts: DeepSeek-V3 with 50% off and DeepSeek-R1 with 75% off at off-peak hours (16:30-00:30 UTC daily). Optimize your workflow while enjoying these exclusive savings. [View Details](#)

API keys

Your API keys are listed below. The API key is only visible and can be copied once at creation. Save it securely. Do not share your API key with others, or expose it in the browser or other client-side code. In order to protect the security of your account, DeepSeek may also automatically disable any API key that we've found has leaked publicly. Usage of API keys created before April 25, 2024, was not tracked.

Name	Key	Created	Last used	
ximarketing	sk-16926*****2e3e	2025-02-26	2025-02-26	 

If you need help on how to use the API, check the document [here: https://api-docs.deepseek.com/](https://api-docs.deepseek.com/)

What are the main topics in modern economics?

Scraping made Simple

Scraping

Consider a simple URL:

<https://www.aeaweb.org/issues/785>

American Economic Review is an academic journal in economics (we don't care anyway). Let's try to write a scraper to download the article titles from the page:

Decisions under Risk Are Decisions under Complexity

by Ryan Oprea

(pp. 3789–3811)

Curbing Leakage in Public Programs: Evidence from India's Direct Benefit Transfer Policy

by Prabhat Barnwal

(pp. 3812–46)

Aiming for the Goal: Contribution Dynamics of Crowdfunding

by Joyee Deb, Aniko Öry, and Kevin R. Williams

(pp. 3847–76)

Scraping

Xi Li

I want to scrape article titles from <https://www.aeaweb.org/issues/785> A title is like this: Decisions under Risk Are Decisions under Complexity Show me python code

23/02/2025, 10:08:28



GPT-4o (Tokens used 527)

To scrape article titles from the provided webpage, you can use Python with libraries such as `requests` and `BeautifulSoup` . Here's a sample code to get you started:

The code is a bit problematic...

```
1 import requests
2 from bs4 import BeautifulSoup
3 # URL of the page to scrape
4 url = "https://www.aeaweb.org/issues/785"
5 # Send a GET request to fetch the page content
6 response = requests.get(url)
7 # Check if the request was successful (status code 200)
8 if response.status_code == 200:
9     # Parse the page content with BeautifulSoup
10    soup = BeautifulSoup(response.content, 'html.parser')
11
12    # Find all article titles
13    # Note: The exact HTML structure of the page needs to be inspected to
14    find the correct tags and classes
15    titles = soup.find_all('h2', class_='title') # Example: Change these
16    tags and classes based on actual HTML structure
17
18    # Print the titles
19    for title in titles:
20        print(title.get_text().strip())
21 else:
22    print(f"Failed to retrieve the page. Status code:
23    {response.status_code}")
```

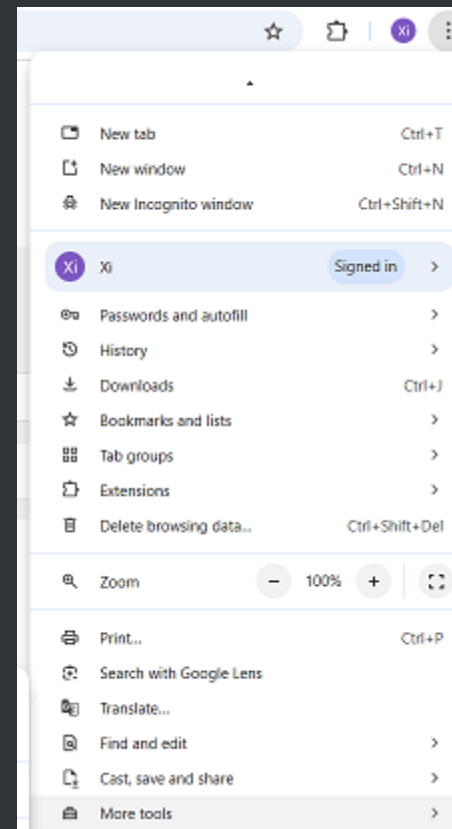
What's wrong with it?

```
1 import requests
2 from bs4 import BeautifulSoup
3 # URL of the page to scrape
4 url = "https://www.aeaweb.org/issues/785"
5 # Send a GET request to fetch the page content
6 response = requests.get(url)
7 # Check if the request was successful (status code 200)
8 if response.status_code == 200:
9     # Parse the page content with BeautifulSoup
10    soup = BeautifulSoup(response.content, 'html.parser')
11
12    # Find all article titles
13    # Note: The exact HTML structure of the page needs to be inspected to
14    find the correct tags and classes
15    titles = soup.find_all('h2', class_='title') # Example: Change these
16    tags and classes based on actual HTML structure
17
18    # Print the titles
19    for title in titles:
20        print(title.get_text().strip())
21 else:
22    print(f"Failed to retrieve the page. Status code:
23    {response.status_code}")
```

Check HTML with Chrome

Open the webpage in your Chrome browser.

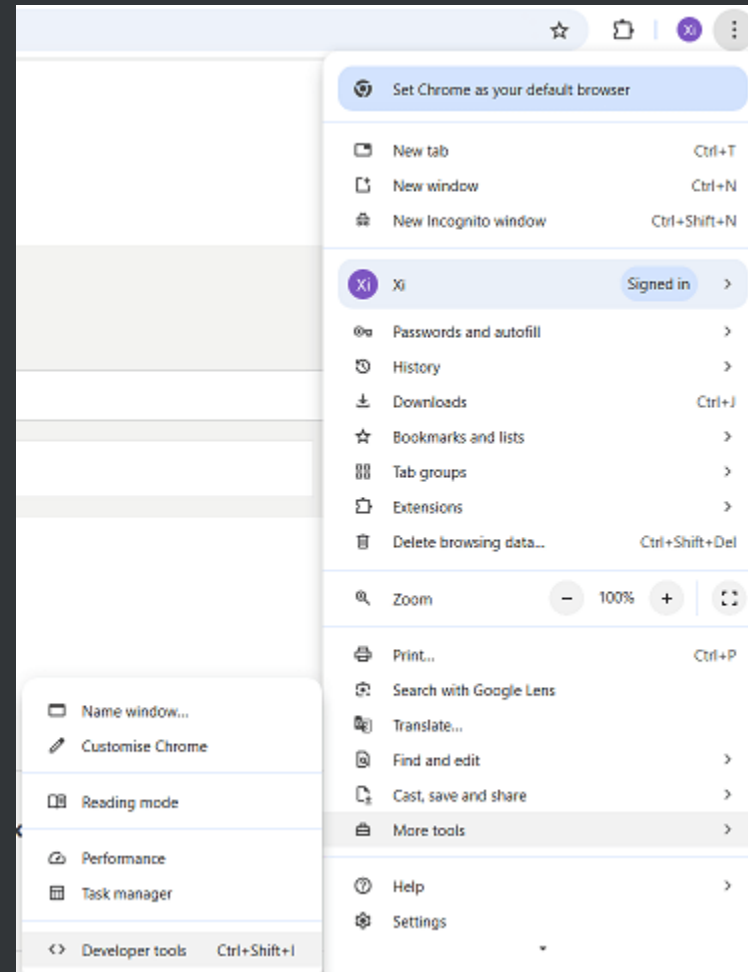
Click the upper right Chrome setting button of your browser and you will be directed here.



Check HTML with Chrome

Choose “More tools” ...

Choose “Developer tools” ...

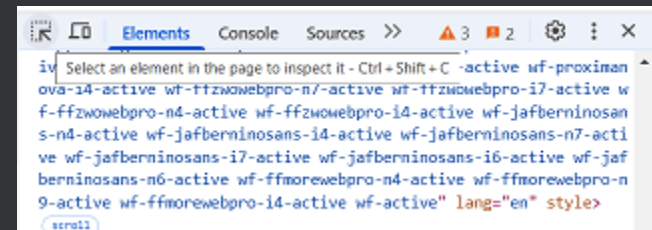


Check HTML with Chrome

Click the  button and you will get to “select an element in the page to inspect it”.

Alternatively, use “Ctrl + Shift + C.”

If needed, also click the  button to switch between desktop and mobile version.



Check HTML with Chrome

```
▼ <article class="journal-article art_18081" id="10.1257/aer.20221227">
  ▼ <h3 class="title">
    ...    <a href="/articles?id=10.1257/aer.20221227">Decisions under Risk Are
           Decisions under Complexity</a> == $0
          </h3>
          ▶ <div class="article-item-authors">...</div>
            <div class="page-range">(pp. 3789–3811)</div>
          </article>
          ▶ <article class="journal-article art_18082" id="10.1257/aer.20161864">...
            </article>
          ▶ <article class="journal-article art_18083" id="10.1257/aer.20181851">...
            </article>
```

The title information is here.

Upload your screenshot to ChatGPT!

Xi Li

I want to download article titles from <https://www.aeaweb.org/issues/785> A title is like this: "Decisions under Risk Are Decisions under Complexity". Show me python code.

 Capture.JPG
JPG 44.23 KB



23/02/2025, 10:22:52



GPT-4o (Tokens used 714)

Certainly! Below is a Python code snippet that uses the `requests` and `beautifulsoup4` libraries to scrape the article titles from the specified webpage.

The code I got

```
1 import requests
2 from bs4 import BeautifulSoup
3
4 # URL of the page to scrape
5 url = 'https://www.aeaweb.org/issues/785'
6
7 # Send a GET request to the page
8 response = requests.get(url)
9
10 # Parse the page content with BeautifulSoup
11 soup = BeautifulSoup(response.content,
12                        'html.parser')
13
14 # Find all article titles
15 titles = soup.find_all('h3', class_='title')
16
17 # Print each title text
18 for title in titles:
19     print(title.get_text(strip=True))
```

Scraping

Consider a simple URL:

<https://www.aeaweb.org/issues/785>

Try to download the links of the articles now, with the help of GPT. Here is an example of the link:

<https://www.aeaweb.org/articles?id=10.1257/aer.20221227>

Scraping

Xi Li

I want to scrape article links from this page: <https://www.aeaweb.org/issues/785>. A link is like this: <https://www.aeaweb.org/articles?id=10.1257/aer.20221227> Show me python code

 Capture2.JPG
JPG 35.32 KB 

23/02/2025, 17:21:07

GPT-4o (Tokens used 764)

```
<article class="journal-article symposia-title" style="border:none;"> Articles
</article>
▼ <article class="journal-article art_18081" id="10.1257/aer.20221227">
  ▼ <h3 class="title">
    <a href="/articles?id=10.1257/aer.20221227">Decisions under Risk Are
    Decisions under Complexity</a> == $0
  </h3>
  ▶ <div class="article-item-authors">...</div>
  <div class="page-range">(pp. 3789–3811)</div>
```

My Code

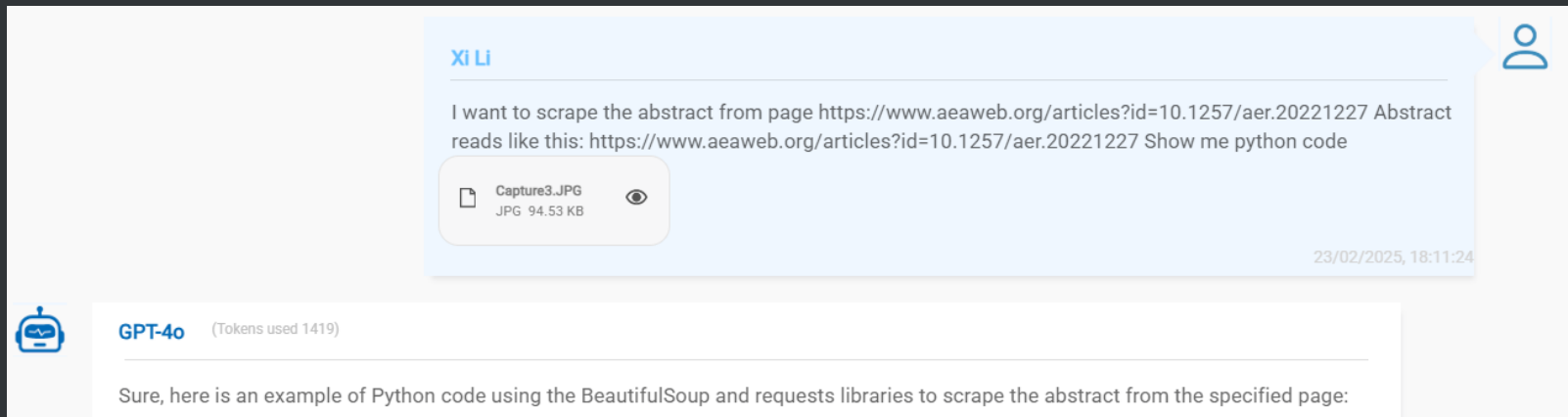
```
1 import requests
2 from bs4 import BeautifulSoup
3 # URL of the page to scrape
4 url = 'https://www.aeaweb.org/issues/785'
5 # Send a GET request to the page
6 response = requests.get(url)
7 response.raise_for_status() # Check that the request was
    successful
8 # Parse the page content with BeautifulSoup
9 soup = BeautifulSoup(response.content, 'html.parser')
10 # Find all article links
11 article_links = []
12 for a_tag in soup.find_all('a', href=True):
13     href = a_tag['href']
14     if href.startswith('/articles?id='):
15         full_url = 'https://www.aeaweb.org' + href
16         article_links.append(full_url)
17 # Print the article links
18 for link in article_links:
19     print(link)
```

Exercise

- Go to <https://www.aeaweb.org/issues/785> to download the links to each article.
- Go to each of the above links to download the abstract of the article.
- Print the all abstracts on the screen.

Exercise

First, I asked GPT to scrape the abstract for a given URL.



Exercise

I got the following code.

```
import requests
from bs4 import BeautifulSoup

# URL of the page to scrape
url = 'https://www.aeaweb.org/articles?id=10.1257/aer.20221227'

# Send a GET request to the page
response = requests.get(url)
response.raise_for_status() # Check that the request was successful

# Parse the page content with BeautifulSoup
soup = BeautifulSoup(response.content, 'html.parser')

# Find the abstract section
abstract_section = soup.find('section', class_='article-information abstract')
if abstract_section:
    # Extract and print the abstract text
    abstract_text = abstract_section.get_text(strip=True, separator=' ')
    print('Abstract:', abstract_text)
else:
    print('Abstract section not found')
```

Exercise

I change it to a function “getabstract”:

```
import requests
from bs4 import BeautifulSoup

def getabstract(url):
    response = requests.get(url)
    response.raise_for_status()
    soup = BeautifulSoup(response.content, 'html.parser')
    abstract_section = soup.find('section', class_='article-information
abstract')
    if abstract_section:
        # Extract and print the abstract text
        abstract_text = abstract_section.get_text(strip=True, separator=' ')
        print('Abstract:', abstract_text)
    else:
        print('Abstract section not found')
```

Exercise

Combine the previous codes, you got it!

```
1 import requests
2 from bs4 import BeautifulSoup
3
4 def getabstract(url):
5     response = requests.get(url)
6     response.raise_for_status()
7     soup = BeautifulSoup(response.content, 'html.parser')
8     abstract_section = soup.find('section', class_='article-information abstract')
9     if abstract_section:
10         # Extract and print the abstract text
11         abstract_text = abstract_section.get_text(strip=True, separator=' ')
12         print('Abstract:', abstract_text)
13     else:
14         print('Abstract section not found')
15
16
17 url = 'https://www.aeaweb.org/issues/785'
18 response = requests.get(url)
19 response.raise_for_status() # Check that the request was successful
20 soup = BeautifulSoup(response.content, 'html.parser')
21 article_links = []
22 for a_tag in soup.find_all('a', href=True):
23     href = a_tag['href']
24     if href.startswith('/articles?id='):
25         full_url = 'https://www.aeaweb.org' + href
26         article_links.append(full_url)
27
28 for link in article_links:
29     getabstract(link)
```

We can do more!

- On page <https://www.aeaweb.org/journals/aer/issues>, we can scrape the URL of each issue.
- On the issue page, we can scrape the URL of each article.
- On the article page, we can scrape the abstract of each article (**it took me two hours**).
- Then, we can save all abstracts in one single file.
- After that, we can perform topic modeling!

We can do more!

I have already uploaded the txt file to the course website, which can be found [here](#):

http://ximarketing.github.io/data/AER_abstracts.txt

Please try the top words and LDA algorithm yourself and demonstrate the main topics in modern economics.

Exercise: Download the Photos of All HKU Business School
Professors from

[https://www.hkubs.hku.hk/people/faculty?
pg=1&staff_type=faculty&subject_area=all&track=professor
iate®ion=hku](https://www.hkubs.hku.hk/people/faculty?pg=1&staff_type=faculty&subject_area=all&track=professoriate®ion=hku)

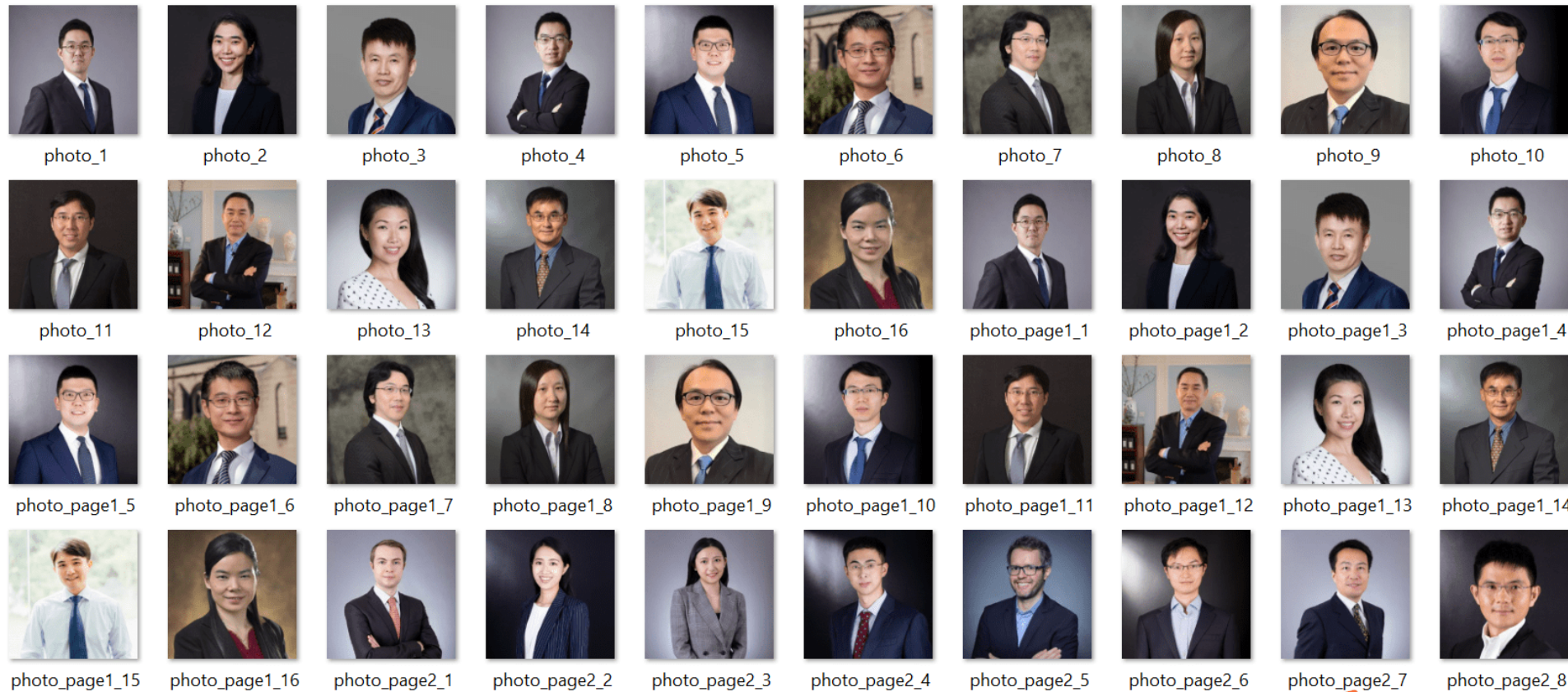
Downloading from a single page:

```
1 import requests
2 from bs4 import BeautifulSoup
3 import os
4 url = "https://www.hkubs.hku.hk/people/faculty?
pg=1&staff_type=faculty&subject_area=all&track=professoriate&region=hku"
5 response = requests.get(url)
6 if response.status_code == 200:
7     soup = BeautifulSoup(response.text, 'html.parser')
8
9     photo_divs = soup.find_all('div', class_='people-item')
10    photo_urls = []
11
12    for div in photo_divs:
13        img_tag = div.find('img')
14        if img_tag and 'src' in img_tag.attrs:
15            photo_url = img_tag['src']
16            photo_urls.append(photo_url)
17
18    os.makedirs('profile_photos', exist_ok=True)
19
20    for idx, photo_url in enumerate(photo_urls):
21        img_response = requests.get(photo_url)
22
23        if img_response.status_code == 200:
24            # Save the photo to the directory
25            with open(f'profile_photos/photo_{idx + 1}.jpg', 'wb') as f:
26                f.write(img_response.content)
27                print(f'Downloaded: photo_{idx + 1}.jpg')
28        else:
29            print(f'Failed to download photo from {photo_url}')
30    else:
31        print(f'Failed to retrieve the page. Status code: {response.status_code}')
```

I further asked the following question...

I have 10 URLs each containing 16 photos, and I want to download all these photos. The URLs are https://www.hkubs.hku.hk/people/faculty?pg=1&staff_type=faculty&subject_area=all&track=professoriate®ion=hku to https://www.hkubs.hku.hk/people/faculty?pg=10&staff_type=faculty&subject_area=all&track=professoriate®ion=hku

I am done!



The code I used

```
1 import requests
2 from bs4 import BeautifulSoup
3 import os
4 base_url = "https://www.hkubs.hku.hk/people/faculty?pg=
   {}&staff_type=faculty&subject_area=all&track=professoriate&region=hku"
5 os.makedirs('profile_photos', exist_ok=True)
6 for page in range(1, 11):
7     url = base_url.format(page)
8     print(f"Fetching page: {url}")
9     response = requests.get(url)
10    if response.status_code == 200:
11        soup = BeautifulSoup(response.text, 'html.parser')
12        photo_divs = soup.find_all('div', class_='people-item')
13        for idx, div in enumerate(photo_divs):
14            img_tag = div.find('img')
15            if img_tag and 'src' in img_tag.attrs:
16                photo_url = img_tag['src']
17                img_response = requests.get(photo_url)
18                if img_response.status_code == 200:
19                    photo_filename = f'profile_photos/photo_page{page}_{idx + 1}.jpg'
20                    with open(photo_filename, 'wb') as f:
21                        f.write(img_response.content)
22                    print(f'Downloaded: {photo_filename}')
23                else:
24                    print(f'Failed to download photo from {photo_url}')
25    else:
26        print(f'Failed to retrieve the page. Status code: {response.status_code}')
```

Try to scrape article titles from the following page:
<https://pubsonline.informs.org/toc/mnsc/71/1>

What else can we do with scraping?

What else can we do with scraping?

- Scraping a website with selenium.
- Web scraping with concurrency (parallel scraping).
- Changing IP addresses with VPN (using proxies / VPNs).

<https://www.youtube.com/embed/cc26zFE8X1k?enablejsapi=1>

Scraping with Selenium

You need to install selenium and download Chrome driver
Moreover, add Chrome driver to your system path

```
1 from selenium import webdriver
2 from selenium.webdriver.common.by import By
3 url = "https://pubsonline.informs.org/toc/mnsc/71/1"
4 driver = webdriver.Chrome()
5 driver.get(url)
6
7 titles = driver.find_elements(By.CLASS_NAME, 'issue-item__title')
8 # Extract and print the text of each title
9 for title in titles:
10     article_title = title.find_element(By.TAG_NAME, 'a').text
11     print(article_title)
12 # Close the WebDriver
13 driver.implicitly_wait(5)
14 driver.quit()
```

Class Review

Basic Operations

Linear Regression

Logistic Regression

Multinomial Logit Model

Artificial Neural Network

Data Visualization

Statistical Tests

Advanced Algorithms

Clustering (and Expectation-Maximization)

Monte Carlo Algorithm

Dynamic Programming

Divide and Conquer

Unstructured Data

Text Data: Word Frequency, Topic Modeling,
Summarization

Image: Color, Brightness, Convolution, etc.

You can also find APIs for text and image analysis.

API and Scraper

How to apply for and use an API.

How to rely on AI to generate your own data scraper.