

Thonny - D:\Python Μαθημα 3\python\scraper.py @ 14 : 1

File Edit View Run Tools Help



scraper.py x

```
1 import requests
2 from bs4 import BeautifulSoup
3
4 url = "https://news.google.com"
5
6 response = requests.get(url)
7 soup = BeautifulSoup(response.text, "html.parser")
8
9 titles = soup.find_all("h3")
10
11 print("\nΒρέθηκαν τίτλοι:\n")
12 for i, t in enumerate(titles, start=1):
13     print(f"{i}. {t.get_text(strip=True)}")
14
```

Shell x

```
>>> %Run scraper.py
```

```
Βρέθηκαν τίτλοι:
```

1. Top stories
2. Local news
3. Picks for you

```
...
```

Web Image Downloade.py ×

```
1 import os
2 import requests
3 from bs4 import BeautifulSoup
4 from urllib.parse import urljoin
5
6 url = "https://www.python.org" # άλλαξέ το σε όποιο site θέλεις
7
8 r = requests.get(url)
9 soup = BeautifulSoup(r.text, "html.parser")
10
11 os.makedirs("images", exist_ok=True)
12
13 for img in soup.find_all("img"):
14     src = img.get("src")
15     if not src:
16         continue
17     img_url = urljoin(url, src)
18     filename = img_url.split("/")[-1]
19
20     try:
21         data = requests.get(img_url).content
22         with open(os.path.join("images", filename), "wb") as f:
23             f.write(data)
24         print("✓ Κατέβηκε:", filename)
25     except:
26         pass
27
```

Shell ×

```
>>> %Run 'Web Image Downloade.py'
```

```
✓ Κατέβηκε: python-logo.png
```

```
>>>
```

Web Title Extractor.py ×

```
1 import requests
2 from bs4 import BeautifulSoup
3
4 url = "https://news.google.com" # άλλαξέ το σε ό
5
6 r = requests.get(url)
7 soup = BeautifulSoup(r.text, "html.parser")
8
9 titles = soup.find_all(["h1", "h2", "h3"])
10
11 for t in titles:
12     print("•", t.get_text(strip=True))
13 |
```

Shell ×

```
>>> %Run 'Web Title Extractor.py'
```

```
• Your
briefing
• Top stories
• Local news
• Picks for you
```

```
>>>
```

Copy

```
1 import requests
2 from bs4 import BeautifulSoup
3
4 url = "https://www.python.org"    # βάλε όποια σελίδα θέλεις
5 r = requests.get(url)
6
7 soup = BeautifulSoup(r.text, "html.parser")
8 text = soup.get_text(separator="\n", strip=True)
9
10 with open("page.txt", "w", encoding="utf-8") as f:
11     f.write(text)
12
13 print("✓ Το καθαρό κείμενο αποθηκεύτηκε στο page.txt")
14 |
```

%Run Web2TXT.py

Το καθαρό κείμενο αποθηκεύτηκε στο page.txt

WebWatch.py ×

```
1 import requests
2 import time
3 import hashlib
4
5 url = "https://www.python.org" # όποιο site θες
6 print("Παρακολουθώ την σελίδα:", url)
7
8 old_hash = ""
9
10 while True:
11     r = requests.get(url)
12     new_hash = hashlib.md5(r.text.encode("utf-8")).hexdigest()
13
14     if old_hash and new_hash != old_hash:
15         print("⚠️ Η ιστοσελίδα ΑΛΛΑΞΕ!")
16     else:
17         print("✅ Καμία αλλαγή...")
18
19     old_hash = new_hash
20     time.sleep(10) # έλεγχος κάθε 10 δευτ.
21
```

Shell ×

>>> %Run WebWatch.py

Παρακολουθώ την σελίδα: <https://www.python.org>

✅ Καμία αλλαγή...