

D:\myprojects\nlp\rag\rag_wiki.py

```
1 # conda create -n rag transformers
2 # conda activate rag
3 # conda install pytorch torchvision torchaudio pytorch-cuda=11.8 -c pytorch -c nvidia
4 # conda install -c conda-forge faiss-gpu
5 # pip install langchain pdfplumber datasets
6 # pip install wikipedia sentence-transformers tiktoken
7 # pip install --upgrade huggingface_hub
8 # pip install --upgrade openai
9
10 import numpy as np
11
12 from transformers import AutoTokenizer
13 from langchain_community.document_loaders import WikipediaLoader
14 from langchain.text_splitter import RecursiveCharacterTextSplitter
15 from langchain_community.embeddings import HuggingFaceEmbeddings
16 from langchain_community.vectorstores import FAISS
17 from langchain_community.vectorstores.utils import DistanceStrategy
18 from sentence_transformers import CrossEncoder
19
20 import os, openai
21 from openai import OpenAI
22
23 SHOW_INSIDE = True
24 # [1] harrypotter 키워드에 대한 wikipedia page 3개 읽어오기
25 pages = WikipediaLoader(query="Harry Potter", load_max_docs=3, lang="en").load()
26
27 # [2.1] text splitter 정의
28 text_splitter = RecursiveCharacterTextSplitter.from_huggingface_tokenizer(
29     tokenizer=AutoTokenizer.from_pretrained("sentence-transformers/all-MiniLM-L12-v2"),
30     chunk_size=256, # 256 chars / chunk
31     chunk_overlap=16, # 이웃하는 chunk사이 16 chars overlap
32     strip_whitespace=True, # chunk 생성시에 공백 제거
33 )
34 # [2.2] 불러온 페이지에 대하여 text_splitting하여 chunks 구하기
35 chunks = text_splitter.split_documents(pages)
36
37 if SHOW_INSIDE :
38     print ('\nsplit된 chunk의 갯수: ',len(chunks))
39     print ('\nchunk[0]의 전체내용: \n',chunks[0])
40     print ('\nchunk[0]의 page_content: \n',chunks[0].page_content)
41     print ('\nchunk[0]의 metadata: \n',chunks[0].metadata)
42     print ('\nchunk[0]의 metadata title: \n',chunks[0].metadata['title'])
43     print ('\nchunk[0]의 metadta summary: \n',chunks[0].metadata['summary'])
44     print ('\nchunk[0]의 metadata source: \n',chunks[0].metadata['source'])
45
46 # [3] bidirectional encoder 모델을 이용한 Embedding 정의
47 bi_encoder = HuggingFaceEmbeddings(
48     #model_name="sentence-transformers/all-MiniLM-L12-v2", model_kwargs={"device": "cpu"}
49     model_name="sentence-transformers/all-MiniLM-L12-v2", model_kwargs={"device": "cuda"}
50 )
51
52 # question과 첫번째 chunk의 내용 embedding 해보기
53 '''
54 embedding1 = bi_encoder.embed_query("Harry Potter")
55 embedding2 = bi_encoder.embed_query(chunks[0].page_content)
56 '''
```

```

57 # [4] chunks들로부터 bi_encoder를 이용한 임베딩을 사용하고, dot-product거리를 이용하여 faiss
index 생성.
58 # FAISS (Facebook AI Similarity Search)
59 faiss_db = FAISS.from_documents(
60     chunks, bi_encoder, distance_strategy=DistanceStrategy.DOT_PRODUCT
61 )
62
63 # [5] cross encoder 정의
64 cross_encoder = CrossEncoder(
65     "cross-encoder/ms-marco-TinyBERT-L-2-v2", max_length=512, device="cpu"
66 )
67 '''
68 cross_encoder.rank(
69     query="Do I like cats?",
70     documents=["I like cats", "I like horses", "I like dogs"],
71     return_documents=True,
72 )
73 '''
74 client = OpenAI(
75     #api_key=os.environ.get("OPENAI_API_KEY"),
76     api_key = "sk-proj-P*****kFJFytKn0x7vYAeqXrL18AR"
77 )
78 #=====
79 default_question = "On what date was the first book in the Harry Potter series originally
published?"
80
81 while True :
82     question = input("\nType 'q' to quit: ")
83
84     if question == 'q' : break
85     elif question == '' :
86         question = default_question
87
88     # [5] question과 유사도가 가장 가까운 상위 25의 chunk를 추출
89     retrieved_chunks = faiss_db.similarity_search(question, k=25)
90
91     reranked_chunks = cross_encoder.rank(
92         (question),
93         [doc.page_content for doc in retrieved_chunks],
94         top_k=3,
95         return_documents=True,
96     )
97
98     context = "".join(doc["text"] + "\n-----\n" for doc in reranked_chunks)
99
100     print("\n\n***최종 상위 답변 3개 통합문장 : \n",context)
101
102     if SHOW_INSIDE :
103         print(f"\n\n*** RAG에 의한 연관된 상위 {len(reranked_chunks)} 개 chunks의 제목과 출
처(URL)")
104         for doc in reranked_chunks:
105             corpus_id = doc["corpus_id"]
106             print(
107                 f"Chunk {corpus_id} from '{retrieved_chunks[corpus_id].metadata['title']}'
\n\t(URL: {retrieved_chunks[corpus_id].metadata['source']})"
108             )
109
110     #-----
111     # Use OpenAI's LLM to generate a response based on the context

```

```

112
113     prompt=f"You are a helpful assistant. Based on the following details about Harry
114     Potter, {context} Answer the question.",
115     try :
116         response = client.chat.completions.create(
117             messages=[
118                 {
119                     "role": "system", "content": prompt,
120                     "role": "user", "content": question,
121                 }
122             ],
123             model="gpt-3.5-turbo-instruct", # or "text-davinci-002" "gpt-3.5-turbo-
instruct"
124             max_tokens=150,
125             temperature=0.5,
126             stop=["\n"]
127         )
128         print("\nOpenAI LLM의 답변:\n", response.choices[0].text.strip())
129     except Exception as e:
130         print("Error : ",e)

```