

D:\myprojects\nlp\rag\rag_pdf.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 import pdfplumber
24
25 SHOW_INSIDE = True
26 '''
27 # [1] harrypotter 키워드에 대한 wikipedia page 3개 읽어오기
28 pages = WikipediaLoader(query="Harry Potter", load_max_docs=3, lang="en").load()
29 '''
30 # [1] Load PDF documents
31 class Document:
32     def __init__(self, page_content, metadata=None):
33         self.page_content = page_content
34         self.metadata = metadata or {} # text_splitter가 metadata를 기본 필요로 함
35
36 pdf_paths = ['example1.pdf', 'example2.pdf', 'glocal1.pdf', 'glocal2.pdf', 'glocal3.pdf']
37 # Specify your PDF file paths here
38 pages_text = []
39 for pdf_path in pdf_paths:
40     with pdfplumber.open(pdf_path) as pdf:
41         for page in pdf.pages:
42             page_text = page.extract_text()
43             if page_text: # Ensure there is text on the page
44                 doc = Document(page_text)
45                 pages_text.append(doc)
46
47 # [2.1] text splitter 정의
48 text_splitter = RecursiveCharacterTextSplitter.from_huggingface_tokenizer(
49     tokenizer=AutoTokenizer.from_pretrained("sentence-transformers/all-MiniLM-L12-v2"),
50     chunk_size=256, # 256 chars / chunk
51     chunk_overlap=16, # 이웃하는 chunk사이 16 chars overlap
52     strip_whitespace=True, # chunk 생성시에 공백 제거
53 )
54 # [2.2] 불러온 페이지에 대하여 text_splitting하여 chunks 구하기
55 #chunks = text_splitter.split_documents(pages)
56 chunks = text_splitter.split_documents(pages_text)
```

```

57
58 if SHOW_INSIDE :
59     print ('\nsplit된 chunk의 갯수: ',len(chunks))
60     print('\nchunk[0]의 전체내용: \n',chunks[0])
61     print('\nchunk[0]의 page_content: \n',chunks[0].page_content)
62
63 # [3] bidirectional encoder 모델을 이용한 Embedding 정의
64 bi_encoder = HuggingFaceEmbeddings(
65     #model_name="sentence-transformers/all-MiniLM-L12-v2", model_kwargs={"device": "cpu"}
66     model_name="sentence-transformers/all-MiniLM-L12-v2", model_kwargs={"device": "cuda"}
67 )
68
69 # question과 첫번째 chunk의 내용 embedding 해보기
70 '''
71 embedding1 = bi_encoder.embed_query("글로컬대학")
72 embedding2 = bi_encoder.embed_query(chunks[0].page_content)
73 '''
74 # [4] chunks들로부터 bi_encoder를 이용한 임베딩을 사용하고, dot-product거리를 이용하여 faiss
index 생성.
75 # FAISS (Facebook AI Similarity Search)
76 faiss_db = FAISS.from_documents(
77     chunks, bi_encoder, distance_strategy=DistanceStrategy.DOT_PRODUCT
78 )
79
80 # [5] cross encoder 정의
81 cross_encoder = CrossEncoder(
82     "cross-encoder/ms-marco-TinyBERT-L-2-v2", max_length=768, device="cuda"
83 )
84
85 client = OpenAI(
86     #api_key=os.environ.get("OPENAI_API_KEY"),
87     api_key = "sk-proj-*****AR"
88 )
89 #=====
==
90 default_question = "글로컬대학 30 예비지정대학으로 선정된 대학은?"
91
92 while True :
93     question = input("\nType 'q' to quit: ")
94
95     if question == 'q' : break
96     elif question == '' :
97         question = default_question
98
99 # [5] question과 유사도가 가장 가까운 상위 25의 chunk를 추출
100 retrieved_chunks = faiss_db.similarity_search(question, k=25)
101
102 reranked_chunks = cross_encoder.rank(
103     (question),
104     [doc.page_content for doc in retrieved_chunks],
105     top_k=3,
106     return_documents=True,
107 )
108
109 context = "".join(doc["text"] + "\n-----\n" for doc in reranked_chunks)
110
111 print("\n\n*** 최종 상위 답변 3개 통합문장 : \n",context)
112
113 if SHOW_INSIDE :
114     print(f"\n\n*** RAG에 의한 연관된 상위 {len(reranked_chunks)} 개 chunks")

```

```

115         for doc in reranked_chunks:
116             corpus_id = doc["corpus_id"]
117             print(
118                 f"Chunk {corpus_id} "
119             )
120
121     #-----
122     # Use OpenAI's LLM to generate a response based on the context
123
124     prompt=f"You are a helpful assistant. Based on the following details about Harry
125     Potter, {context} Answer the question.",
126     try :
127         response = client.chat.completions.create(
128             messages=[
129                 {
130                     "role": "system", "content": prompt,
131                     "role": "user", "content": question,
132                 }
133             ],
134             model="gpt-3.5-turbo-instruct", # or "text-davinci-002" "gpt-3.5-turbo-
instruct"
135             max_tokens=150,
136             temperature=0.5,
137             stop=["\n"]
138         )
139         print("\nOpenAI LLM의 답변:\n", response.choices[0].text.strip())
140     except Exception as e:
141         print("Error : ",e)

```