

```

1 #*****
2 # etri+google 한글,영어 음성인식 및 STT 테스트conda create etri
3 #*****
4 # pydub API 문서 - https://github.com/jiaaro/pydub/blob/master/API.markdown
5 # 설치할 패키지들
6 # pip install SpeechRecognition gTTS
7 # pip install playsound==1.2.2          # version 1.3.0은 오류 발생
8 # pip install pydub librosa
9 # conda install pyAudio matplotlib ffmpeg
10 # python etri ko
11
12 # STT관련
13 import speech_recognition as sr
14 from pydub import AudioSegment, silence
15 import librosa, librosa.display
16 import matplotlib.pyplot as plt
17 import numpy as np
18
19 # TTS 관련
20 import playsound
21 from gtts import gTTS
22
23 # web browser 구동 관련
24 import webbrowser
25
26 import os, sys, random
27 from time import ctime
28
29 import urllib3
30 import json
31 import base64
32
33 # ETRI 서버 및 개인 accesskey 관련 정보
34 accessKey = "3ab9d53f-cce1-4d82-b207-5593e15407e0"
35 audioFilePath = "KOR_M_RM0276MKDH0135.pcm"
36
37 #*****
38 # ETRI 한글 음성인식과 발음 평가
39 #*****
40 def sentence_from_etri(wavFile) :
41     make_melspectrogram(wavFile, 'original_voice')
42
43     # 16khz mono 오디오 데이터로 읽기
44     wavAudio = AudioSegment.from_file(wavFile, format="wav").set_frame_rate(16000)
45
46     print('앞쪽 조용한 구간 시작점: ',silence.detect_silence(wavAudio.silent(10000),
47                                                             silence_thresh=-50))
48
49     print('frame_rate,channels,sample_width:',wavAudio.frame_rate,wavAudio.channels,wavAudio.sample_width)
50     print('loudness(rms), peak값:',wavAudio.rms,wavAudio.max)
51
52     # 헤더를 제외한 실제 음성 데이터 부분만 추출
53     rawData = wavAudio.raw_data
54     audioContents = base64.b64encode(rawData).decode("utf8")
55     # .wav 파일로 부터 실행시 사용
56     '''
57     audioFile = open(audioFilePath, "rb")
58     audioContents = base64.b64encode(audioFile.read()).decode("utf8")
59     audioFile.close()
60     '''
61
62     # 음성인식
63     http = urllib3.PoolManager()
64     openApiURL = "http://aiopen.etri.re.kr:8000/WiseASR/Recognition"
65     requestJson = {

```

```

66     "access_key": accessKey,
67     "argument": {
68         "language_code": 'korean',      # languageCode
69         "audio": audioContents
70     }
71 }
72
73 response = http.request(
74     "POST",
75     openApiURL,
76     headers={"Content-Type": "application/json; charset=UTF-8"},
77     body=json.dumps(requestJson)
78 )
79
80 print("[responseCode_음성인식] " + str(response.status))
81 print("[responseData_음성인식] " + str(response.data, "utf-8"))
82
83 sentence_str = str(response.data, "utf-8")
84 sentence_start_pos = sentence_str.find('recognized')+len('recognized: ')+2
85 sentence_end_pos = sentence_str.find('"', sentence_start_pos+1)
86 sentence = sentence_str[sentence_start_pos:sentence_end_pos]
87
88 # 발음 평가
89 openApiURL = "http://aiopen.etri.re.kr:8000/WiseASR/PronunciationKor" # 한국어
90 #openApiURL = "http://aiopen.etri.re.kr:8000/WiseASR/Pronunciation" # 영어
91 script = "이름이 뭐예요?"
92 requestJson = {
93     "access_key": accessKey,
94     "argument": {
95         "language_code": 'korean',      # languageCode
96         "script": sentence,             # script,
97         "audio": audioContents
98     }
99 }
100
101 response2 = http.request(
102     "POST",
103     openApiURL,
104     headers={"Content-Type": "application/json; charset=UTF-8"},
105     body=json.dumps(requestJson)
106 )
107
108 print("[responseCode_발음평가] " + str(response2.status))
109 print("[responseData_발음평가] " + str(response2.data, "utf-8"), "\n")
110
111 score_str = str(response2.data, "utf-8")
112 score_pos = score_str.find('score')+len('score:')+1
113 score = score_str[score_pos:score_pos+6]
114
115 return sentence, score
116
117 #*****
118 # 한글,영어음성 문장을 음성으로 입력받아 Text로 출력하는 음성인식
119 #*****
120 def sentence_from_audio( lang, direction='How can I help You? [say "goodbye" to finish]',
121     soundFrom='MIC') :
122     COMPARE_RECOGNIZERS = False
123
124     print("언어: ", lang)
125
126     recog = sr.Recognizer()                # Recognizer 인스턴스 recog 생성
127     sentence = []                          # 인식기별로 인식결과 문장들 저장
128
129     with sr.Microphone() as source:
130         recog.adjust_for_ambient_noise(source, duration=1)
131         while True:

```

```

131 print("Wn",direction)
132 audio = recog.listen(source) # 오디오 데이터 추출
133
134 # 사용자가 말한 음성을 wave파일로 저장
135 with open("spokenWav.wav","wb") as fpWave:
136     fpWave.write(audio.get_wav_data())
137
138 #print("audio type: ",type(audio)," data: ",audio)
139 try :
140     if soundFrom == 'MIC': # 마이크로 입력한 음성 인식
141         if lang == 'en' :
142             result = recog.recognize_google(audio) # 구글 WebSpeech 음성인식 이용, 인식
143         elif lang == 'ko' :
144             result, score = sentence_from_etri("spokenWav.wav")
145
146             sentence.append( result )
147             print("==> you spoke: ", sentence[0], "["+score+" 점"]")
148     else : # 음성 파일로 부터 구글 Webspeech 이용 인식
149         audio_file = sr.AudioFile('spokenWav.wav')
150         with audio_file as source:
151             audio = recog.record(source)
152             sentence.append( recog.recognize_google(audio) ) # 구글 WebSpeech 음성인식 이용, 인
153             print("The FILE says: ", sentence[0])
154
155     ...
156     if COMPARE_RECOGNIZERS == True:
157         sentence.append( recog.recognize_bing(audio, key) ) # Microsoft 음성인식 이용
158         sentence.append( recog.recognize_ibm(audio, username, password) ) # IBM 음성인식 이용
159         sentence.append( recog.recognize_houndify(audio client_id, client_key) ) # Houndify 음
160         sentence.append( recog.recognize_wit(audio, key) ) # Wit.ai 음성인식 이용
161         # 아래 인식엔진들은 SpeechRecognition외에 별도 패키지 추가가 필요
162         sentence.append( recog.recognize_google_cloud(audio, ...) ) # Google Cloud 이용
163         sentence.append( recog.recognize_sphinx(audio, ...) ) # CMYU Sphinx 이용, Offline 이용
164
165         print('* 다른 인식기 인식 결과 *')
166         for ndx in range(1,len(sentence)):
167             print(ndx, "번째 인식기 결과: ", sentence[ndx])
168
169         ...
170         return sentence[0]
171
172     except sr.UnknownValueError: # 인식결과 오류 발생시
173         print('Please say again')
174
175     except sr.RequestError as e: # 서버접속, 사용한도 초과, 인터넷연결 등의 문제시 발생
176         print('오류'.format(e))
177
178 *****
179 # 구글 TTS를 이용하여, TEXT를 음성파일로 저장후 소리 출력
180 *****
181 def googleTTS(sentence, language='en') :
182     tts = gTTS(text=sentence, lang=language) # language로 설정한 언어와 문장으로 tts 인스턴스 생성
183
184     r = random.randint(1,1000000) # 파일명 난수로 임의로 설정
185     audioFile = './audio_' +str(r) + '.mp3'
186
187     #print('audio_file_name: ',audioFile)
188     tts.save(audioFile) # 문장을 음성파일로 저장
189     playsound.playsound(audioFile) # 음성파일을 불러서 음성 출력
190
191     print(sentence) # 문장 내용 출력 후 음성 파일은 지움
192     #os.remove(audioFile)
193
194 *****
195 # 사용자가 입력한 문장에 따라 몇가지만 반응하는 함수

```

```

194 #*****
195 def respond(sentence=''):
196     searchWords = ['search', 'find', 'i want to know', 'open web browser']
197     nameWords = ['your name', 'should i call']
198     positiveAnswerWords = ['you free', 'can i call you', 'is it okay']
199
200     if 'time' in sentence: # time 단어가 들어있는 경우
201         print(ctime())
202         googleTTS(ctime(), 'en')
203         return 'ok'
204
205     for phrase in nameWords :
206         if phrase in sentence: # 이름을 묻는 질문인 경우
207             txt = 'My name is AI Kim'
208             txtKor = '저는 AI 로봇입니다'
209             print(txt)
210             googleTTS(txtKor, 'ko')
211             return 'ok'
212
213     for phrase in positiveAnswerWords :
214         if phrase in sentence: # 긍정적 답을 하는게 좋은 질문인 경우
215             googleTTS('Sure, Feel free to say anything', 'en')
216             return 'ok'
217
218     for phrase in searchWords :
219         if phrase in sentence: # search 단어가 들어있는 경우, 웹브라우저로 검색결과 실행
220             search_data = sentence_from_audio('en', 'What do you want to search for?', 'MIC')
221             url = 'https://google.com/search?q=' + search_data
222             webbrowser.get().open(url)
223             str = 'Search Result for ' + search_data
224             print(str)
225             return 'ok'
226
227     return 'NO_MATCH'
228
229 #*****
230 # ETRI 한글 질의 응답
231 #*****
232 def answer_from_etri(question) :
233     # ETRI 서버로 전송할 데이터를 json형태로 준비
234     openApiURL_QnA = "http://aiopen.etri.re.kr:8000/WiseQAnal"
235
236     requestJson = {
237         "access_key": accessKey,
238         "argument": {
239             "text": question
240         }
241     }
242
243     http = urllib3.PoolManager()
244     response = http.request(
245         "POST",
246         openApiURL_QnA,
247         headers={"Content-Type": "application/json; charset=UTF-8"},
248         body=json.dumps(requestJson)
249     )
250
251     #print("[responseCode] " + str(response.status))
252
253     return response.status, str(response.data, "utf-8")
254
255 #*****
256 # Make MelSpectrogram
257 #*****
258 def make_melspectrogram(wavFilename, filename):
259     spokenWaveData, srVal = librosa.load(wavFilename, sr=None) # 음성파일을 읽어서 spokenWaveData에 저장

```

```

260 spokenWaveData *= 1.0
261 plt.interactive(False)
262 plt.figure()
263 S = librosa.feature.melspectrogram(y=spokenWaveData, sr=srVal, n_mels=64)
264 S_db = librosa.power_to_db(S, ref=np.max)
265 fig = plt.figure(figsize = [0.8, 0.8])
266 ax = fig.add_subplot(1,1,1)
267 ax.axes.get_xaxis().set_visible(False)
268 ax.axes.get_yaxis().set_visible(False)
269 ax.set_frame_on(False)
270 librosa.display.specshow(S_db)
271 melSavePath = './'
272 melSavePath += filename+'.jpg'
273 print(melSavePath)
274 plt.savefig(melSavePath, dpi=400, bbox_inches='tight',pad_inches=0)    # 이미지파일로 출력
275 plt.close('all')
276
277 #*****
278 # 직접 실행시 프로그램 시작 부분
279 #*****
280 if __name__ == '__main__':
281
282     lang = 'en'
283     question = ''
284     if len(sys.argv) > 1:
285         if sys.argv[1] == 'ko' or sys.argv[1] == 'en':
286             lang = sys.argv[1]
287             googleTTS('문장을 말하세요', 'ko')
288         else :
289             question = sys.argv[1]
290     else :
291         print('[1]한영 음성인식 사용법: Python etri.py ko/en')
292         print('[2]한글 질의응답 사용법: Python etri.py 질문문장')
293
294     while True:
295
296         if question == '' : # 음성인식
297             result_sentence = sentence_from_audio(lang) # 음성인식하여 TEXT(문자) 값을 가져옴
298
299             if 'goodbye' in result_sentence:
300                 googleTTS('bye bye', 'en')
301                 break
302             else : # 종료명령이 아니면
303                 if respond(result_sentence) == 'NO_MATCH' and sys.argv[1]=='en':
304                     print("Sorry, I didn't get it. Would you say again?")
305             else : # 질의 응답
306                 result_status, result_sentence = answer_from_etri(question)
307                 if result_status == 200 : # 답 문장을 받은 경우
308                     print('ETRI 리턴 값: Wn',result_sentence, 'Wn')
309
310                     start_pos = result_sentence.find('vTitles')    # 답 위치 항목 찾기
311                     last_pos = result_sentence.find('vQTopic',start_pos)    # 답 위치 항목 찾기
312
313                     if start_pos == -1 : # 문장을 받았으나 답이 포함되지 않은 경우
314                         print('받은 답이 없습니다 !')
315                     else :
316                         answer_no = 1;
317                         while True :
318                             start_pos = result_sentence.find('strExplain',start_pos+1)
319                             if start_pos == -1 or start_pos > last_pos:
320                                 break
321                             start_pos += len('strExplain: ')
322                             end_pos = result_sentence.find('dWeightEn',start_pos)
323                             answer =result_sentence[start_pos:end_pos-2]
324                             print('ETRI 답 [',answer_no,']Wn',answer, 'Wn')
325                             answer_no += 1

```

```
326
327 question = input('*** 질문을 입력하세요[종료시엔 end 라고 입력하세요] : ')
328 if question == 'end' :
329     break
```