

D:\myprojects\nlp\transformer\self_attention.py

```
1 # Vanilla Self Attention
2
3 # conda activate word2vec
4 # python self_attention.py
5
6 import torch
7 import torch.nn.functional as F
8 import numpy as np
9 import math
10
11 #=====
12 # Input Target batch 데이터 생성 함수
13 # #=====
14 def make_batch(sentences, forTrain=True):
15     input_batch = []
16
17     if forTrain :
18         target_batch = []
19
20         for sen in sentences:
21             words = sen.split()
22             inputNdx = [wordDictionary[n] for n in words[:-1]]
23             target = wordDictionary[words[-1]]
24
25             # input번째 자리만 1이 되게 One-Hot Encoding
26             # input list 가 2개 값만 가지면 아래는 2개의 one-hot vectors 생성
27             input_batch.append(np.eye(noOfClasses)[inputNdx])
28
29             target_batch.append(target)
30
31         return input_batch, target_batch
32
33     else : # for Prediction
34         sentence = sentences
35
36         words = sentence.split()
37         #print(f'words: {words}')
38         inputNdx = [wordDictionary[n] for n in words[:-1]]
39
40         input_batch.append(np.eye(noOfClasses)[inputNdx]) # One-Hot Encoding
41
42         return input_batch
43
44 #=====
45 # Self-Attention Class
46 # #=====
47 class SelfAttention(torch.nn.Module):
48     def __init__(self, embeddingSize, headSize):
49         super(SelfAttention, self).__init__()
50         self.embeddingSize = embeddingSize
51         self.headSize = headSize
52         self.headDimension = embeddingSize // headSize
53
54         assert (
55             self.headDimension * headSize == embeddingSize
56         ), "Embedding size needs to be divisible by headSize"
57
```

```

58 # query, key, value를 scalar값이 아닌 2차원 행렬 형태로 처리
59 # 그래서 이 샘플 코드에서는 2차원 weighr표현에 좋은 nn.Linear 형식을 사용
60 # 만약, outer product 가 아닌 dot product를 할 경우엔 vector가 필요함.
61 self.values = torch.nn.Linear(self.headDimension, self.headDimension, bias=False)
62 print(f'\nvalues 벡터 입출력 크기: {self.values}')
63 self.keys = torch.nn.Linear(self.headDimension, self.headDimension, bias=False)
64 self.queries = torch.nn.Linear(self.headDimension, self.headDimension, bias=False)
65 self.fc_out = torch.nn.Linear(headSize * self.headDimension, embeddingSize)
66
67 def forward(self, valuesBatch, keysBatch, queriesBatch):
68     batchSize = queriesBatch.shape[0] # sentence 갯수와 동일 => 12
69     valueLength, keyLength, queryLength = valuesBatch.shape[1], keysBatch.shape[1],
queriesBatch.shape[1]
70     print(f'\nbatchSize: {batchSize}\n [query, key, value] 의 시퀀스 길이는: \
71         [{valueLength}, {keyLength}, {queryLength}] \n')
72
73     # Split the embedding into self.headSize different pieces
74     queriesBatchMultiHead = queriesBatch.reshape(batchSize, queryLength,
self.headSize, self.headDimension)
75     keysBatchMultiHead = keysBatch.reshape(batchSize, keyLength, self.headSize,
self.headDimension)
76     valuesBatchMultiHead = valuesBatch.reshape(batchSize, valueLength, self.headSize,
self.headDimension)
77
78     queriesBatch = self.queries(queriesBatchMultiHead)
79     keysBatch = self.keys(keysBatchMultiHead)
80     valuesBatch = self.values(valuesBatchMultiHead)
81
82     # q * t 를 하여 similarity(attention 값)을 구함,
83     # similarity 계산법은 정하기 나름인데, 여기서는 outer product(행렬곱)를 수행하므로
84     # attention 결과는 4차원 tensor [batch, head, queryDimension, keyDimension]
85     attention = torch.einsum("nqhd,nkhd->nhqk", [queriesBatch, keysBatch])
86     # queries shape: (N, queryLength, headSize, head_dim),
87     # keysBatch shape: (N, keyLength, headSize, head_dim),
88     # attention: (N, headSize, queryLength, keyLength)
89
90     # Scale attention scores
91     attention = attention / (self.embeddingSize ** (1/2))
92     attention = F.softmax(attention, dim=-1)
93
94     # 방금 계산한 attention과 values를 곱하여 출력값 계산,
95     # 그 후 embeddingSize 사이즈를 갖도록 한 차원 축소된 shape 변경
96     out = torch.einsum("nhql,nlhd->nqhd", [attention, valuesBatch]).reshape(
97         batchSize, queryLength, self.headSize * self.headDimension
98     )
99     # attention shape: (N, headSize, queryLength, keyLength)
100    # valuesBatch shape: (N, valueLength, headSize, head_dim)
101    # after einsum: (N, queryLength, headSize, head_dim) then we flatten the last two
dimensions
102
103    out = self.fc_out(out) # (N, queryLength, embeddingSize) 크기로 리턴
104    print(f'attention 후 출력 y: {out}')
105    return out
106
107 #=====
108 sentences = ["i like my dogs", "i love hot coffee", "i hate cold milk", "i enjoy hot tea",
109     "you like cute cats", "you love hot milk", "you hate cold coffee", "you want rainy days"
110     ,
111     "he expects big promotion", "he wants nice gifts", "he wants my dogs", "i expect rainy
days"]
112 dtype = torch.float

```

```

113
114 wordList = list(set(" ".join(sentences).split()))
115 wordList.sort()      # 매 실행시마다 동일한 word index 번호를 갖게
116
117 wordDictionary = {w: i for i, w in enumerate(wordList)}
118 print('\n','='*80)
119 print("word dictionary:",wordDictionary,'\n')
120 number_dict = {i: w for i, w in enumerate(wordList)}
121
122 windowSize = 3
123 headSize = 4
124 noOfClasses = int(math.ceil(len(wordDictionary) / headSize) * headSize)
125 embeddingSize = noOfClasses
126
127 print(f'\nheadSize: {headSize}, noOfClasses: {noOfClasses}')
128
129 # make one-hot encoded input batch
130 input_batch, target_batch = make_batch(sentences,True)
131 valueLength, keyLength, queryLength = windowSize, windowSize, windowSize # sequence
lengths => 3
132
133 queries = torch.tensor(np.array(input_batch), dtype=torch.float32)
134 keys = torch.tensor(np.array(input_batch), dtype=torch.float32)
135 values = torch.tensor(np.array(input_batch), dtype=torch.float32)
136 print('\n','='*80)
137 print(f"Self-Attention용 [{len(queries)}]개 input batch:",input_batch)
138
139 attention = SelfAttention(embeddingSize, headSize)
140
141 # Self-attention(Forward Pass only) 실행
142 out = attention(values, keys, queries)
143
144 print(f'\noutput의 shape[batchSize, sequenceLength, VectorDimension]: {out.shape}\n') #
Expected shape: (batchSize, sequenceLength,embeddingSize)
145 print('첫 번째 batch의 앞 쪽 2개의 시퀀스에 대한 attention 결과값 출력해보기')
146 for i in range(2):
147     print('-'*80)
148     for j in range(out.shape[1]):
149         print(f'sequence #{i} [{j}] {sentences[i].split()[j]}\n')
150         for k in range(out.shape[2]):
151             print(out[i, j, k].item(), end=' ')
152         print('\n')
153     print('\n')
154
155

```