

ECE 60146 Assignment 10

Adhitha Dias

Wednesday 24th April, 2024; 18:14

Contents

1	Introduction	1
2	Output during training	1
3	Qualitative Evaluation	2
4	Quantitative Evaluation with DistilBERT	4
5	Code Snippets	5

1 Introduction

This document contains the results of the training and evaluation of the model on the SQuAD dataset. The model is trained on the SQuAD dataset using the DistilBERT model (for comparison) on question and answers, and locally fine-tuning a bert pretrained model. The model is trained for 100 epochs but the model overfits and the accuracy goes down at the 100th epoch. Therefore, I am reporting the results using the checkpoint obtained at the 35th epoch. All code snippets are borrowed from the class notes, Piazza posts, and the assignment description.

2 Output during training

```
1 {'loss': 2.484, 'grad_norm': 17.250259399414062, 'learning_rate': 4.714285714285714e-05, 'epoch':  
2 0.57}  
3 {'loss': 1.442, 'grad_norm': 17.010774612426758, 'learning_rate': 4.428571428571428e-05, 'epoch':  
4 1.14}  
5 {'loss': 0.9571, 'grad_norm': 39.44040298461914, 'learning_rate': 4.1428571428571437e-05, 'epoch':  
6 1.71}  
7 {'loss': 0.7126, 'grad_norm': 42.816986083984375, 'learning_rate': 3.857142857142858e-05, 'epoch':  
8 2.29}  
9 {'loss': 0.5209, 'grad_norm': 21.191295623779297, 'learning_rate': 3.571428571428572e-05, 'epoch':  
10 2.86}  
11 {'loss': 0.3601, 'grad_norm': 16.623889923095703, 'learning_rate': 3.285714285714286e-05, 'epoch':  
12 3.43}  
13 {'loss': 0.3125, 'grad_norm': 0.35003775358200073, 'learning_rate': 3e-05, 'epoch': 4.0}  
14 {'loss': 0.2026, 'grad_norm': 0.2589515149593353, 'learning_rate': 2.714285714285714e-05, 'epoch':  
15 4.57}  
16 {'loss': 0.1959, 'grad_norm': 8.793052673339844, 'learning_rate': 2.4285714285714288e-05, 'epoch':  
17 5.14}  
18 {'loss': 0.1411, 'grad_norm': 52.41377639770508, 'learning_rate': 2.1428571428571428e-05, 'epoch':  
19 5.71}
```

```

13 {'loss': 0.1168, 'grad_norm': 0.037031445652246475, 'learning_rate': 1.8571428571428572e-05, 'epoch':
    6.29}
14 {'loss': 0.0986, 'grad_norm': 1.3157042264938354, 'learning_rate': 1.5714285714285715e-05, 'epoch':
    6.86}
15 {'loss': 0.0565, 'grad_norm': 43.251949310302734, 'learning_rate': 1.2857142857142857e-05, 'epoch':
    7.43}
16 {'loss': 0.062, 'grad_norm': 3.89984393119812, 'learning_rate': 1e-05, 'epoch': 8.0}
17 {'loss': 0.04, 'grad_norm': 0.0038578195963054895, 'learning_rate': 7.142857142857143e-06, 'epoch':
    8.57}
18 {'loss': 0.0326, 'grad_norm': 0.0037382750306278467, 'learning_rate': 4.285714285714286e-06, 'epoch':
    9.14}
19 {'loss': 0.0162, 'grad_norm': 0.0022480105981230736, 'learning_rate': 1.4285714285714286e-06, 'epoch':
    9.71}
20 9.71}
21 {'train_runtime': 12699.0025, 'train_samples_per_second': 5.512, 'train_steps_per_second': 0.689, 'train_loss': 0.4435298219953264, 'epoch': 10.0}

```

Listing 1: Output during the first 8 epochs of training

3 Qualitative Evaluation

Since the model returns uncased strings, I had to convert the expected answers to lowercase before comparing them. The evaluation script calculates the exact match and F1 score for each answer. The F1 score and exact match are calculated as shown in Listing8.

Because token join operation does not properly stitch the predicted answers, I have used several regular expression replacements to fix the predicted answers. The regular expressions used are shown in Listing6 and below.

```

1 prediction = re.sub(r'##', '', prediction)
2 prediction = re.sub(r',', ', ', prediction)
3 prediction = re.sub(r'\.', '.', prediction)
4 prediction = re.sub(r'\$', '$', prediction)
5 prediction = re.sub(r'%', '%', prediction)
6 prediction = re.sub(r' - ', ' - ', prediction)
7 prediction = re.sub(r'\(', '(', prediction)
8 prediction = re.sub(r'\)', ')', prediction)
9 prediction = re.sub(r'\'', '\'', prediction)
10 prediction = re.sub(r'. com', '.com', prediction)
11 prediction = re.sub(ur' 0xE2 ', '0xE2', prediction)
12 prediction = re.sub(r' 0', '0', prediction)

```

Listing 2: Regular Expressions used to fix the predicted answers

Note that these regular expressions do not fix the irregularities in the answers completely. I had to eye the output and select the regular expressions that I thought would fix the predicted answers.

Some of the examples of questions and answers are shown in Listing3. I have listed why the model did not get a perfect score in the comments.

```

1 — the model is capable of selecting the correct answers with multiple words
2 Question: What are three games, in addition to Brick, which have been included with the iPod?
3 Predicted Answer: parachute, solitaire, and music quiz
4 Expected: parachute, solitaire, and music quiz
5 Exact match: 1
6 F1 score: 1.0
7 — model is able to answer questions with numbers
8 Question: How many Chinese troops and medics were involved in the relief efforts?
9 Predicted Answer: 135,000
10 Expected: 135,000
11 Exact match: 1
12 F1 score: 1.0
13 — model sometimes returns empty answers
14 Question: How tall is One World Trade Center in meters?
15 Predicted Answer:
16 Expected: 541.3

```

```

17 Exact match: 0
18 F1 score: 0
19 — model return partial answers
20 Question: What was the first commercial solar concentrating system?
21 Predicted Answer: the solar total energy project
22 Expected: solar total energy project (step) in shenandoah, georgia, usa
23 Exact match: 0
24 F1 score: 0.5714285714285714
25 — model answer is correct but not in the expected format, and I could not figure out how to fix
    it because xx, xx is a valid format for a comma separated list of values
26 Question: How many hectares does the island have in total?
27 Predicted Answer: 2, 500
28 Expected: 2,500
29 Exact match: 0
30 F1 score: 0
31 — slight mismatch due to punctuation. There are several similar examples in the output that
    reduces the F1 score and exact match score
32 Question: What website showed video of an altercation between Mariah Carey and Nicki Minaj?
33 Predicted Answer: tmz.
34 Expected: tmz
35 Exact match: 0
36 F1 score: 0
37 — right answer but slightly different letter in names, the tokenizer does not have this name in
    its vocabulary
38 Question: Who designed Chopin's tombstone?
39 Predicted Answer: clesinger
40 Expected: clésinger.
41 Exact match: 0
42 F1 score: 0
43 —
44 Question: What organization was Broca in the process of disentangling himself from?
45 Predicted Answer: societe de biologie
46 Expected: société de biologie
47 Exact match: 0
48 F1 score: 0.6666666666666666
49 — sometimes does not include the article in the answer
50 Question: In what historical era does the book take place?
51 Predicted Answer: great depression
52 Expected: the great depression
53 Exact match: 0
54 F1 score: 0.8
55 — model is able to answer questions with dates
56 Question: When was the Montana Territory formed?
57 Predicted Answer: april 26, 1864
58 Expected: april 26, 1864
59 Exact match: 1
60 F1 score: 1.0
61 — just adding this here because I was born a Buddhist
62 Question: Who is the Buddha of this Buddha era?
63 Predicted Answer: gautama buddha
64 Expected: gautama buddha
65 Exact match: 1
66 F1 score: 1.0
67 — model is able to answer questions with URLs or website names
68 Question: Where was the message posted?
69 Predicted Answer: cyberctm.com
70 Expected: cyberctm.com
71 Exact match: 1
72 F1 score: 1.0
73 —

```

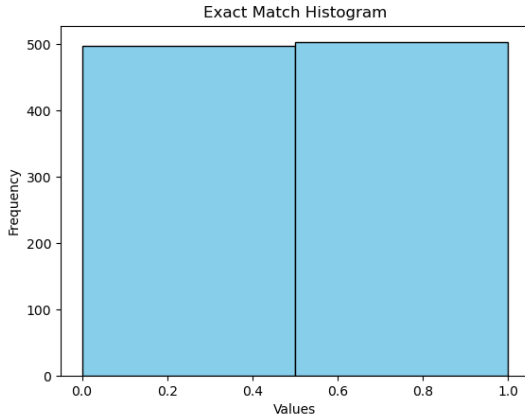
Listing 3: Example Question and Answer Outputs

4 Quantitative Evaluation with DistilBERT

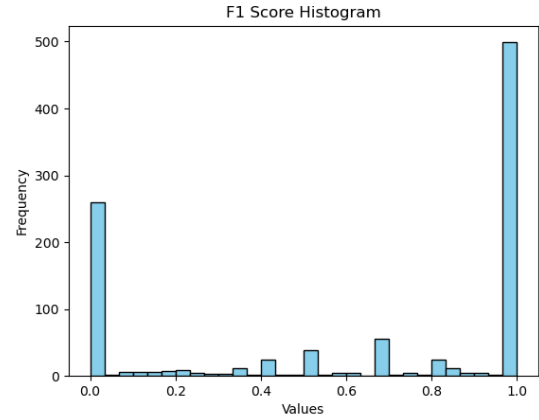
The model selects mostly correct answers as evident from the histograms in Figure1. The model has a mean exact match score of 0.769 and a median exact match score of 1.0. The mean F1 score is 0.891 and the median F1 score is 1.0. The comparison of the performance of the local model and DistilBERT is shown in Table1.

	Mean Exact Match	Median Exact Match	Mean F1 Score	Median F1 Score
Local	0.509	1.0	0.636	1.0
DistilBERT	0.769	1.0	0.891	1.0

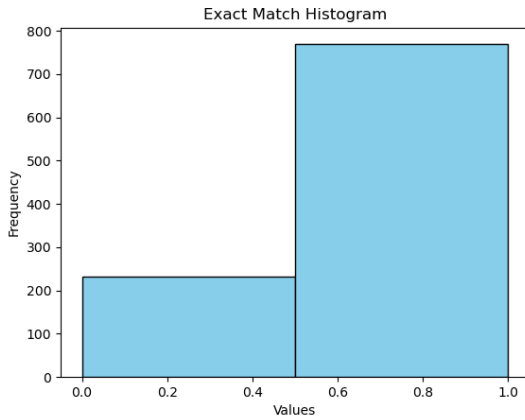
Table 1: Comparison of the performance of the local model and DistilBERT



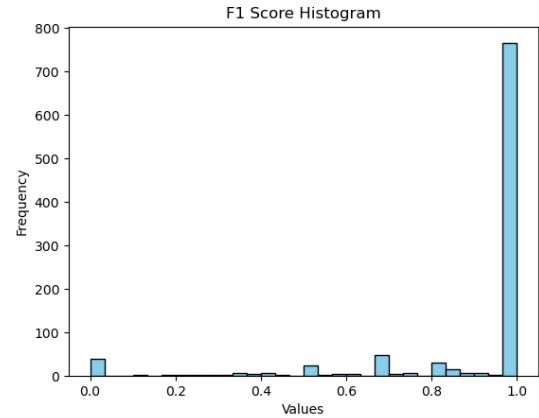
(a) Histogram of Exact Match Scores



(b) Histogram of F1 Scores



(c) Histogram of Exact Match Scores for DistilBERT



(d) Histogram of F1 Scores for DistilBERT

Figure 1: Histograms of Exact Match and F1 Scores

5 Code Snippets

```
1 import pickle
2 import pandas as pd
3 from datasets import Dataset
4
5 with open('dataset/train_dict.pkl', 'rb') as f:
6     train_dict = pickle.load(f)
7
8 with open('dataset/test_dict.pkl', 'rb') as f:
9     test_dict = pickle.load(f)
10
11 with open('dataset/eval_dict.pkl', 'rb') as f:
12     eval_dict = pickle.load(f)
13
14 with open('dataset/train_data_processed.pkl', 'rb') as f:
15     train_processed = pickle.load(f)
16
17 with open('dataset/test_data_processed.pkl', 'rb') as f:
18     test_processed = pickle.load(f)
19
20 with open('dataset/eval_data_processed.pkl', 'rb') as f:
21     eval_processed = pickle.load(f)
22
23 print(train_dict.keys())
24 print(test_dict.keys())
25 print(eval_dict.keys())
26
27 print(train_processed.keys())
28 print(test_processed.keys())
29 print(eval_processed.keys())
30
31 # print one element of the train_dict
32 for k,v in train_dict.items():
33     print('key', k)
34     print('value', len(v))
35     # break
36
37 train_dataset = Dataset.from_pandas(pd.DataFrame(train_processed))
38 test_dataset = Dataset.from_pandas(pd.DataFrame(test_processed))
39 eval_dataset = Dataset.from_pandas(pd.DataFrame(eval_processed))
```

Listing 4: Dataset Definition

```
1 from src.dataset import \
2     train_dict, test_dict, eval_dict, \
3     train_processed, test_processed, eval_processed, \
4     train_dataset, test_dataset, eval_dataset
5
6 from transformers import BertForQuestionAnswering
7
8 model_name = 'bert-base-uncased'
9 model = BertForQuestionAnswering.from_pretrained(model_name)
10
11 # load the model from the saved checkpoint
12 # model = BertForQuestionAnswering.from_pretrained('./results/checkpoint-2000')
13
14 print(model._modules)
15
16 from transformers import TrainingArguments
17
18 training_args = TrainingArguments(
19     output_dir='./results',           # output directory
20     use_mps_device=False,             # use 16-bit (mixed precision) training
21     num_train_epochs=100,             # number of training epochs
22     per_device_train_batch_size=8,    # batch size for training
23     per_device_eval_batch_size=8,     # batch size for evaluation
24     warmup_steps=0,                   # number of warmup steps for learning rate scheduler
25     weight_decay=0.01,                # strength of weight decay
26     logging_dir='./logs',             # directory for storing logs
```

```

27     # save_steps=1_000,                # frequency of model checkpoints
28 )
29
30 from transformers import Trainer
31
32 trainer = Trainer(
33     model=model,                        # the instantiated Transformers model to be trained
34     args=training_args,                # training arguments, defined above
35     train_dataset=train_dataset,       # training dataset
36     eval_dataset=eval_dataset         # evaluation dataset
37 )
38
39 print('training the network')
40 trainer.train()

```

Listing 5: Training Script

```

1 from src.dataset import \
2     train_dict, test_dict, eval_dict, \
3     train_processed, test_processed, eval_processed, \
4     train_dataset, test_dataset, eval_dataset
5
6 from transformers import BertForQuestionAnswering
7
8 model_name = 'bert-base-uncased'
9 # model = BertForQuestionAnswering.from_pretrained(model_name)
10
11 # load the model from the saved checkpoint
12 model = BertForQuestionAnswering.from_pretrained('./results/checkpoint-31500')
13
14 print(model._modules)
15
16 from transformers import TrainingArguments
17
18 training_args = TrainingArguments(
19     output_dir='./results',            # output directory
20     use_mps_device=False,              # use 16-bit (mixed precision) training
21     num_train_epochs=10,               # number of training epochs
22     per_device_train_batch_size=8,     # batch size for training
23     per_device_eval_batch_size=8,     # batch size for evaluation
24     warmup_steps=0,                   # number of warmup steps for learning rate scheduler
25     weight_decay=0.01,                # strength of weight decay
26     logging_dir='./logs',              # directory for storing logs
27     # save_steps=1_000,                # frequency of model checkpoints
28 )
29
30 from transformers import Trainer
31
32 trainer = Trainer(
33     model=model,                        # the instantiated Transformers model to be trained
34     args=training_args,                # training arguments, defined above
35     train_dataset=train_dataset,       # training dataset
36     eval_dataset=eval_dataset         # evaluation dataset
37 )
38
39 print('NOT training the network')
40 # trainer.train()
41
42 import numpy as np
43 from transformers import BertTokenizer
44
45 # Initialize the tokenizer
46 tokenizer = BertTokenizer.from_pretrained(model_name)
47 # tokenizer = BertTokenizer.from_pretrained('./results/checkpoint-2000')
48
49
50 x = trainer.predict(test_dataset)
51 start_pos, end_pos = x.predictions
52 start_pos = np.argmax(start_pos, axis=1)
53 end_pos = np.argmax(end_pos, axis=1)
54

```

```

55 from src.util import compute_exact_match, f1_score
56
57 lexact_match = []
58 lf1_score = []
59
60 import re
61
62 for k, (i, j) in enumerate(zip(start_pos, end_pos)):
63     tokens = tokenizer.convert_ids_to_tokens(test_processed['input_ids'][k])
64     print('Question:', test_dict['question'][k])
65
66     prediction = ' '.join(tokens[i:j+1])
67     prediction = re.sub(r'##', '', prediction)
68     prediction = re.sub(r',', '', prediction)
69     prediction = re.sub(r'\.', '', prediction)
70     prediction = re.sub(r'\$', '$', prediction)
71     prediction = re.sub(r'%', '%', prediction)
72     prediction = re.sub(r'—', '-', prediction)
73     prediction = re.sub(r'\(', '(', prediction)
74     prediction = re.sub(r'\)', ')', prediction)
75     prediction = re.sub(r'\'', '\'', prediction)
76     prediction = re.sub(r'.com', '.com', prediction)
77     prediction = re.sub(r'\ ', '', prediction)
78     prediction = re.sub(r'0', '0', prediction)
79     print('Predicted Answer:', prediction)
80     ground_truth = test_dict['answers'][k]['text'][0]
81     ground_truth = ground_truth.lower()
82     print('Expected:', ground_truth)
83
84     exmatch = compute_exact_match(prediction, ground_truth)
85     f1 = f1_score(prediction, ground_truth)
86     print('Exact match:', exmatch)
87     print('F1 score:', f1)
88
89     lexact_match.append(exmatch)
90     lf1_score.append(f1)
91
92     print('——')
93
94 from src.util import print_stats, plot_histogram
95
96 print_stats(lexact_match, lf1_score)
97 plot_histogram(lexact_match, 2, 'Values', 'Frequency', 'Exact Match Histogram', 'images/
    histogram_exact_match.png')
98 plot_histogram(lf1_score, 30, 'Values', 'Frequency', 'F1 Score Histogram', 'images/
    histogram_f1_score.png')

```

Listing 6: Evaluation Script

```

1 from src.util import compute_exact_match, f1_score
2 from src.dataset import \
3     train_dict, test_dict, eval_dict, \
4     train_processed, test_processed, eval_processed, \
5     train_dataset, test_dataset, eval_dataset
6
7 import torch
8 # device = "cuda" if torch.cuda.is_available() else "cpu"
9 device = "cpu"
10
11 from transformers import pipeline
12
13 question_answerer = pipeline('question-answering', model='distilbert-base-cased-distilled-squad')
14
15 lexact_match = []
16 lf1_score = []
17
18 for i in range(len(test_dict['question'])):
19     result = question_answerer(context=test_dict['context'][i], question=test_dict['question'][i])
20
21     question = test_dict['question'][i]
22     print('Question:', question)

```

```

23 answer = result['answer']
24 print('Answer:', answer)
25 correct_answer = test_dict['answers'][i]['text'][0]
26 print('Correct Answer:', correct_answer)
27 exmatch = compute_exact_match(answer, correct_answer)
28 print('Exact Match:', compute_exact_match(answer, correct_answer))
29 f1 = f1_score(answer, correct_answer)
30 print('F1 Score:', f1)
31 lexact_match.append(exmatch)
32 lf1_score.append(f1)
33 print('—')
34
35 from src.util import print_stats
36 from src.util import plot_histogram
37
38 print_stats(lexact_match, lf1_score)
39 plot_histogram(lexact_match, 2, 'Values', 'Frequency', 'Exact Match Histogram', 'images/
    histogram_comp_exact_match.png')
40 plot_histogram(lf1_score, 30, 'Values', 'Frequency', 'F1 Score Histogram', 'images/
    histogram_comp_f1_score.png')

```

Listing 7: Script to run distilbert-base-cased-distilled-squad

```

1 def compute_exact_match(prediction, truth):
2     return int(prediction == truth)
3
4 def f1_score(prediction, truth):
5     prediction_tokens = prediction.split()
6     truth_tokens = truth.split()
7
8     # if either the prediction or the truth is no-answer then f1 = 1 if they agree, 0 otherwise
9     if len(prediction_tokens) == 0 or len(truth_tokens) == 0:
10         return int(prediction_tokens == truth_tokens)
11
12     common_tokens = set(prediction_tokens) & set(truth_tokens)
13
14     # if there are no common tokens then f1 = 0
15     if len(common_tokens) == 0:
16         return 0
17
18     precision = len(common_tokens) / len(prediction_tokens)
19     recall = len(common_tokens) / len(truth_tokens)
20     if precision + recall == 0:
21         return 0
22
23     return 2 * (precision * recall) / (precision + recall)
24
25 def plot_histogram(l, bins, xlabel, ylabel, title, filename):
26     import matplotlib.pyplot as plt
27     plt.hist(l, bins=bins, color='skyblue', edgecolor='black')
28     plt.xlabel(xlabel)
29     plt.ylabel(ylabel)
30     plt.title(title)
31     plt.savefig(filename)
32     plt.clf()
33
34 def print_stats(lexact_match, lf1_score):
35     import numpy as np
36     print(lexact_match)
37     print(lf1_score)
38     mean_exact_match = np.mean(lexact_match)
39     mean_f1_score = np.mean(lf1_score)
40
41     median_exact_match = np.median(lexact_match)
42     median_f1_score = np.median(lf1_score)
43
44     print('Mean Exact Match:', mean_exact_match)
45     print('Mean F1 Score:', mean_f1_score)
46
47     print('Median Exact Match:', median_exact_match)

```



```
48 print('Median F1 Score:', median_f1_score)
```

Listing 8: Utility Functions

References

- [1] Class Notes
- [2] DLStudio, <https://engineering.purdue.edu/kak/distDLS/>