

Distributions

September 2, 2025

1 Distributions of Words & Sentences

This assignment is comprised of three required tasks and an additional task for bonus points:

1. The first task is to compute the frequency vs. rank distribution of the words in Moby Dick. For this, you will need to tokenize the document and create a vocabulary mapping word types to their document frequency.
2. The second task is to segment the document into sentences and compute the sentence length distribution. Here you will experiment with spaCy's default sentence segmenter as well as the simple rule-based Sentencizer.
3. The third task is the same as the first except that we use subword tokenization.
4. [Bonus] Use spaCy's NE recognizer to find all named entities in the first 2,500 paragraphs. Count how many times they appear in the document and consolidate them based on their most frequent type.

1.1 Write Your Name Here:

2 Submission instructions

1. Click the Save button at the top of the Jupyter Notebook.
2. Please make sure to have entered your name above.
3. Select Edit -> Clear Output of All Cells. This will clear all the outputs from all cells (but will keep their content).
4. Select Run -> Run All Cells. This will run all the cells in order, and may take several seconds.
5. Once you've rerun everything, select File -> Save and Export Notebook As -> PDF and download a PDF version showing the code and the output of all cells, and save it in the same folder that contains the notebook file.
6. Look at the PDF file and make sure all your solutions and outputs are there, displayed correctly.
7. Submit **both** the PDF and your notebook file .ipynb on Canvas.
8. Make sure your Canvas submission contains the correct files by downloading it after posting it on Canvas.

2.1 Word distributions using the SpaCy tokenizer (40 + 10 points)

First, create the spaCy tokenizer.

```
[1]: from spacy.lang.en import English
      nlp = English()
```

```
tokenizer = nlp.tokenizer
```

Create a *vocab* dictionary. This dictionary will map tokens to their counts in the input text file.

```
[ ]: vocab = {}
```

Read the input file line by line.

1. Tokenize each line.
2. For each token in the line that contains only letters, convert it to lower case and increment the corresponding count in the dictionary.
 - If the token does not exist in the dictionary yet, insert it with a count of 1. For example, the first time the token 'water' is encountered, the code should evaluate $\text{vocab}[\text{'water'}] = 1$.

At the end of this code segment, *vocab* should map each word type to the number of times it appeared in the entire document. There should be 16830 word types and 214287 words in Moby Dick.

```
[ ]: with open('../data/melville-moby_dick.txt', 'r') as f:
      for line in f:
          # YOUR CODE GOES HERE

      print('There are', len(vocab), 'word types in Moby Dick.')
      print('There are', sum(vocab.values()), 'words in Moby Dick.')
```

Create a list *ranked* of tuples (*word*, *freq*) that contains all the words in the vocabulary *vocab* sorted by frequency. For example, if $\text{vocab} = \{\text{'duck':}2, \text{'goose':}5, \text{'turkey':}3\}$, then $\text{ranked} = [(\text{'goose'}, 5), (\text{'turkey'}, 3), (\text{'duck'}, 2)]$.

```
[ ]: ranked = [] # YOUR CODE GOES HERE
```

Print the top 10 words in the sorted list.

```
[ ]: print('Size of vocabulary:', len(ranked))
      for word, freq in ranked[:10]:
          print(word, freq)
```

Plot the frequency vs. rank of the top ranked words in Moby Dick.

```
[ ]: import matplotlib.pyplot as plt
      ranks = range(1, 50 + 1)
      freqs = [t[1] for t in ranked[:50]]
      plt.scatter(ranks, freqs, c='#1f77b4', alpha=0.5)
      plt.show()
```

```
[ ]: import math
      ranks = [1 + math.log(r) for r in range(1, len(ranked) + 1)]
      freqs = [math.log(t[1]) for t in ranked]
      plt.scatter(ranks, freqs, c='#1f77b4', alpha=0.5)
```

```
plt.show()
```

2.2 Sentence distributions (40 + 10 points)

First, try to create the spaCy nlp object from the entire text of Moby Dick. This will likely not work, it is not a good idea to read all the text.

```
[ ]: import spacy

nlp = spacy.load("en_core_web_sm")
text = open('../data/melville-moby_dick.txt', 'r').read()
doc = nlp(text)
```

Instead, read the document paragraph by paragraph, i.e. in chunks of text separated by empty lines. Before using spaCy to segment a paragraph into sentences, replace each end of line character with a whitespace, to allow a sentence to span multiple lines. After sentence segmentation, for each sentence in the paragraph append its length (in tokens) to *lengths*. Use the default *nlp* class to process each paragraph and split it into sentences. Stop after processing 1000 paragraphs. This will be slow, so be patient.

```
[ ]: import spacy

nlp = spacy.load("en_core_web_sm")

# the number of paragraphs read so far.
count = 0
# stores the length of each sentence processed so far.
lengths = []
# make sure the file is read and processed line by line.
with open('../data/melville-moby_dick.txt', 'r') as f:
    # YOUR CODE GOES HERE

len150 = [l for l in lengths if l <= 150]
plt.hist(len150, bins = 20)
plt.show()
```

Next, do the same processing as above, but use the more robust Sentencizer to split paragraphs into sentences. Note the speedup.

```
[ ]: from spacy.lang.en import English

nlp = English()
```

```

nlp.add_pipe("sentencizer")

# the number of paragraphs read so far.
count = 0
# stores the length of each sentence processed so far.
lengths = []
with open('../data/melville-moby_dick.txt', 'r') as f:
    # YOUR CODE GOES HERE

len150 = [1 for l in lengths if l <= 150]
plt.hist(len150, bins = 20)
plt.show()

```

Note the difference between the two histograms. Identify at least 5 examples of sentences in Moby Dick that are segmented differently by the two approaches. Copy them below and explain the differences. Which method seems to be more accurate?

2.3 Word distribution using OpenAI's subword tokenization (30 points)

In this part, we will compute the frequency vs. rank based on the the BPE subword tokenization created by the [tiktoken module from OpenAI](#).

Read the input file line by line.

1. Tokenize each line using `tiktoken` encoder and decoder for GPT-3.5.
2. For each token in the line that contains only letters, convert it to lower case and increment the corresponding count in the dictionary.
 - If the token does not exist in the dictionary yet, insert it with a count of 1. For example, the first time the token 'water' is encountered, the code should evaluate `vocab['water'] = 1`.

At the end of this code segment, `vocab` should map each word type to the number of times it appeared in the entire document. There should be 12659 unique types and 248615 total tokens in Moby Dick.

```

[ ]: import tiktoken

# To get the tokeniser corresponding to a specific model in the OpenAI API:
enc = tiktoken.encoding_for_model("gpt-3.5-turbo")

vocab = {}
with open('../data/melville-moby_dick.txt', 'r') as f:
    for line in f:
        # YOUR CODE HERE

```

```
print('There are', len(vocab), 'unique tokens in Moby Dick.')
print('There are', sum(vocab.values()), 'tokens in Moby Dick.')
```

Rank the tokens based on their frequency, then plot frequency vs. rank.

```
[ ]: ranked = [] # YOUR CODE GOES HERE

print('Size of vocabulary:', len(ranked))
for word, freq in ranked[:10]:
    print(word, freq)

[ ]: import matplotlib.pyplot as plt
ranks = range(1, 50 + 1)
freqs = [t[1] for t in ranked[:50]]
plt.scatter(ranks, freqs, c='#1f77b4', alpha=0.5)
plt.show()
```

2.4 [Bonus] Named Entities (10 + 10 + 10 + 10 + 10 points)

Useful documentation is at: - <https://spacy.io/usage/linguistic-features#named-entities> - <https://spacy.io/api/entityrecognizer>

```
[2]: import spacy

nlp = spacy.load("en_core_web_sm")

# These are all the entity types covered by spaCy's NE recognizer.
nlp.pipe_labels['ner']
```

```
[2]: ['CARDINAL',
      'DATE',
      'EVENT',
      'FAC',
      'GPE',
      'LANGUAGE',
      'LAW',
      'LOC',
      'MONEY',
```

```
'NORP',
'ORDINAL',
'ORG',
'PERCENT',
'PERSON',
'PRODUCT',
'QUANTITY',
'TIME',
'WORK_OF_ART']
```

Read the first 2,500 paragraphs in Moby Dick and extract all named entities into a dictionary `ne_counts` that maps each *named entity* to its frequency. By *named entity* we mean a tuple (*name*, *type*) where *name* is the entity name as a string, and *type* is its entity type. For example, if the name 'Ahab' appears with the NE type 'PERSON' 50 times, then the dictionary should map the key ('Ahab', 'PERSON') to the value 50.

```
[ ]: # The number of paragraphs read so far.
count = 0
# Stores the dictionary of named entites and their counts.
ne_counts = {}

# Make sure the file is read line by line.
with open('../data/melville-moby_dick.txt', 'r') as f:
    # YOUR CODE GOES HERE
```

Create a list `ranked_ne` containing all the items in the `ne_counts` dictionary that is sorted in descending order by their frequency.

```
[ ]: ranked_ne = [] # YOUR CODE GOES HERE

# This should display 2974 unique named entities, with the top two being
# ('one', 'CARDINAL') 354 and ('Ahab', 'PERSON') 351
print('Unique named entities:', len(ranked_ne))
for ne, count in ranked_ne[:50]:
    print(ne, count)
```

2.4.1 Consolidate named entities

Some names appear with more than one type, most often due to errors in named entity recognition. One way to fix such errors is to use the fact that typically a name occurs with just one meaning in a document, as such it has just one type. In this part of the assignment, we will consolidate the extracted names such that the counts for the same name appearing with multiple types are added together, and by associating the name with the type that it appears with most often.

Create a dictionary `ne_types` that maps each name to a dictionary that contains all the types the name appears with, where each type is mapped to the corresponding count. Use information from the dictionary `ne_counts` above.

```
[ ]: ne_types = {}

# YOUR CODE HERE


print(ne_types['Queequeg']) # this should print {'PERSON': 15, 'GPE': 23, 'LOC':
↪ 3, 'NORP': 10, 'ORG': 1}

print(ne_types['Gabriel']) # this should print {'PERSON': 20}
```

Create the consolidated dictionary `ne_cons` that maps each name to a tuple that contains its most frequent type and the total count over all types. Use information from the dictionary `ne_types` above.

```
[ ]: ne_cons = {}

# YOUR CODE HERE


print(ne_cons['Queequeg']) # this should print ('GPE', 52)

print(ne_cons['Gabriel']) # this should print ('PERSON', 20)
```

Create a list `ranked_nec` that contains only the consolidated entries from `ne_cons` whose type is among the types listed in the list `types` below, sorted in descending order based on their total counts.

```
[ ]: types = ['PERSON', 'GPE', 'ORG', 'LOC', 'FAC']

# YOUR CODE HERE


ranked_nec =

# This should display 1632 consolidated named entities, with the top two
↪ entries being
# Ahab ('PERSON', 351) and Pequod ('GPE', 150)
print('Consolidated named entities:', len(ranked_nec))
for ne, count in ranked_nec[:30]:
    print(ne, count)
```

[Extra Bonus points 1] (5 points) Select one name from the dictionary `ne_counts` that appears frequently with 2 types and explain why you think spaCy's named entity recognizer associated the name with those 2 types.

[Extra Bonus points 2] (10 points) Find all the syntactic dependency paths connecting the subject Ahab with a direct object, e.g. ‘Ahab’ $\rightarrow$ nsubj $\rightarrow$ <verb> $\rightarrow$ dobj $\rightarrow$ <object>. Rank all the object words based on how frequently they appear connected to ‘Ahab’ through this syntactic pattern, and for the top 10 objects display the list of verbs that are used with each object.

Useful documentation is at: - <https://spacy.io/usage/linguistic-features#dependency-parse>

[4]: `# YOUR CODE HERE`

2.5 Bonus points

Anything extra goes here. For example:

- Write code Li (1992) showing that just random typing of letters including a space will generate “words” with a Zipfian distribution. Generate at least 1 million characters before you compute word frequencies.
 - Show mathematically that random typing results in a Zipf’s distribution by computing probabilities for all words that contain just 1 letter, 2 letters, ...
- Implement the BPE algorithm, where you break ties by selecting to merge in lexicographic order. Train the BPE algorithm on a large corpus and then use it to do subword tokenization on the Moby Dick corpus. What are the top 10 most frequent tokens and how does it compare with what you got from `tiktokenizer`.

[]: