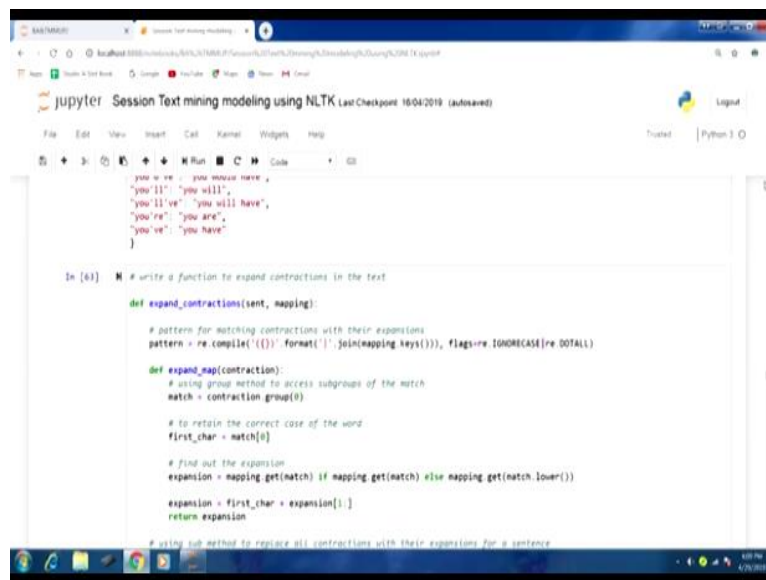


Business Analytics & Text Mining Modeling Using python
Prof. Gaurav Dixit
Department of Management Studies
Indian Institute of Technology Roorkee

Lecture-39
Text Mining and Modeling-Part II

Welcome to the course business analytics and text mining modeling using python. So, in previous few lectures we have been discussing some of the steps that are required to process the text convert you know transform the unstructured text into a form where we can you know do our text mining modeling. So, specifically we were talking about expanding contractions, so we talked about that aspect, we talked about the mapping you know contraction mapping that we discussed in the previous lecture.

(Refer Slide Time: 00:52)



```
you're : you're have ,
you'll : you will ,
you'll've : you will have ,
you're : you are ,
you've : you have
}

In [43]: # write a function to expand contractions in the text

def expand_contractions(text, mapping):
    # pattern for matching contractions with their expansions
    pattern = re.compile('({})'.format(' '.join(mapping.keys()))).flags=re.IGNORECASE|re.DOTALL)

    def expand_map(contraction):
        # using group method to access subgroups of the match
        match = contraction.group(0)

        # to retain the correct case of the word
        first_char = match[0]

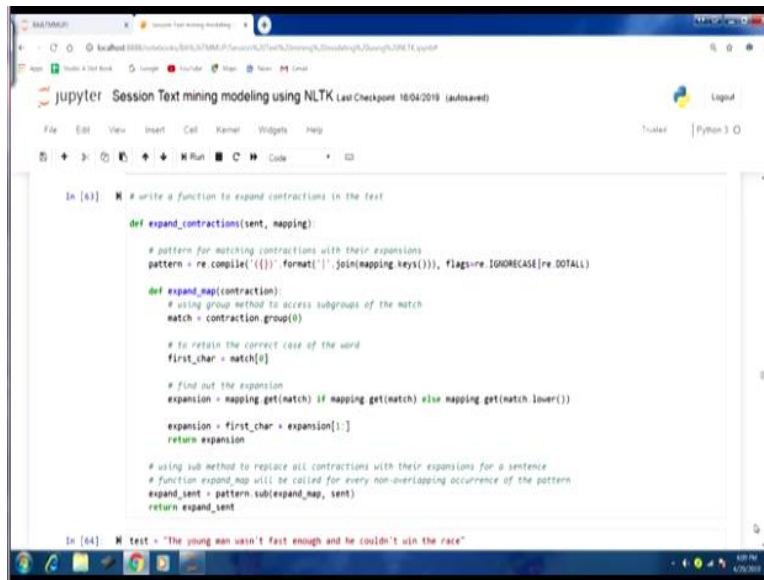
        # find out the expansion
        expansion = mapping.get(match) if mapping.get(match) else mapping.get(match.lower())

        expansion = first_char + expansion[1:]
        return expansion

    # using sub method to replace all contractions with their expansions for a sentence
```

So, let us go through quickly briefly this and then we will continue.

(Refer Slide Time: 01:00)



```
In [43]: # write a function to expand contractions in the text

def expand_contractions(sent, mapping):

    # pattern for matching contractions with their expansions
    pattern = re.compile('({})'.format(' '.join(mapping.keys()))).flags=re.IGNORECASE|re.DOTALL

    def expand_map(contraction):
        # using group method to access subgroups of the match
        match = contraction.group(0)

        # to retain the correct case of the word
        first_char = match[0]

        # find out the expansion
        expansion = mapping.get(match) if mapping.get(match) else mapping.get(match.lower())

        expansion = first_char + expansion[1:]
        return expansion

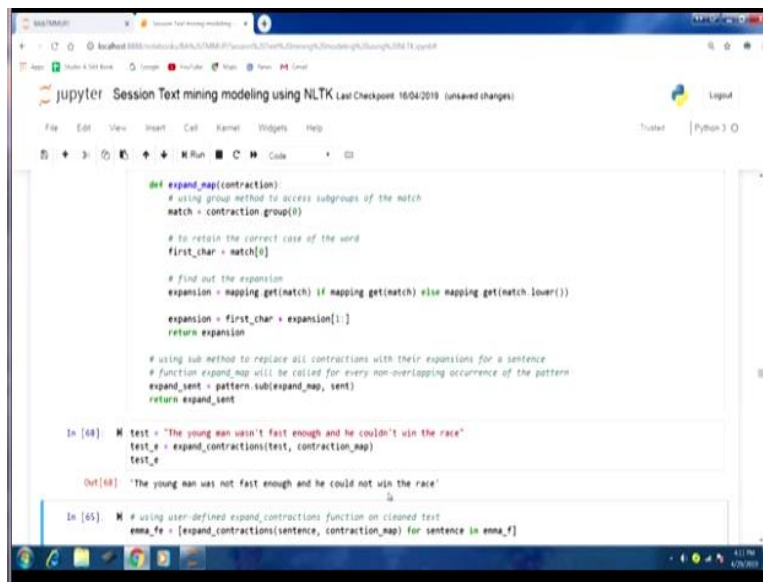
    # using sub method to replace all contractions with their expansions for a sentence
    # function expand_map will be called for every non-overlapping occurrence of the pattern
    expand_sent = pattern.sub(expand_map, sent)
    return expand_sent

In [44]: # test = "The young man wasn't fast enough and he couldn't win the race"
```

So, here we are defining a function that can be use to expand contractions in the text, so you can see define expand underscore contraction and any sentence you know text and then mapping. So, mapping that we discussed before we have a dict object there, so we can use that to perform of our mapping here. So, here we are again using pattern for matching contractions with their expansion, so you can see that you know this pattern here.

And we will be using this pattern to actually find out all the contractions with their expansions for a sentence, you can look at this part of the code. So, here I am calling you know this using the pattern dot sub I am finding out all the occurrences here you know expanded underscore map and this particular this function is being called which will actually sense essentially it will you know do the mapping, so this part we have discussed in the previous lecture, so this is the function that we define to perform this.

(Refer Slide Time: 02:09)



```
def expand_map(contraction):
    # using group method to access subgroups of the match
    match = contraction.group(0)

    # to retain the correct case of the word
    first_char = match[0]

    # find out the expansion
    expansion = mapping.get(match) if mapping.get(match) else mapping.get(match.lower())

    expansion = first_char + expansion[1:]
    return expansion

# using sub method to replace all contractions with their expansions for a sentence
# function expand_map will be called for every non-overlapping occurrence of the pattern
expand_sent = pattern.sub(expand_map, sent)
return expand_sent

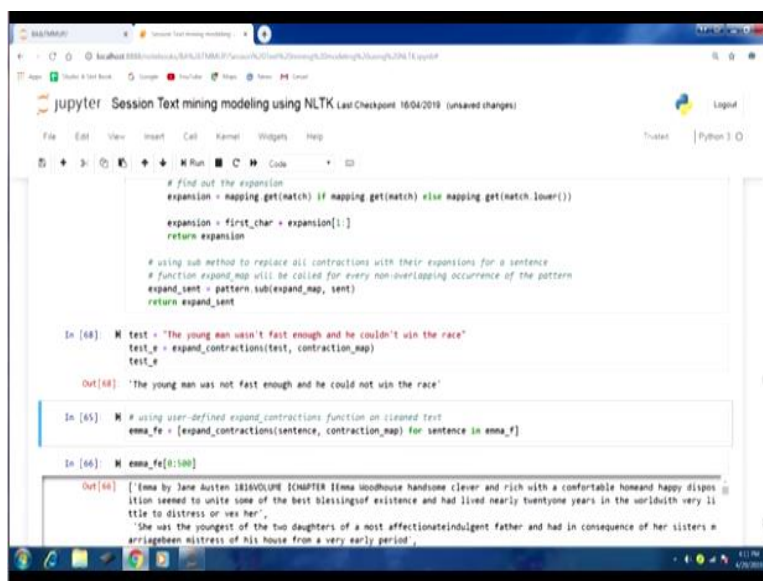
In [48]: test = "The young man wasn't fast enough and he couldn't win the race"
test_e = expand_contractions(test, contraction_map)
test_e

Out[48]: 'The young man was not fast enough and he could not win the race'

In [45]: # using user-defined expand_contractions function on cleaned text
emma_fe = [expand_contractions(sentence, contraction_map) for sentence in emma_f]
```

Now once we through this then we can apply this expand underscore contraction in one of the text. So, I have taken this sample text here you can see the young man was not fast enough and he could not win the race. So, was not and could not are the 2 contractions that are there and let us use our function and let us see whether they are you know converted into their expansion form, so let me run this and you can see in the output that you know was not has been converted in to was not and could not has you know has been converted into could not, so you can see the function is working quite well.

(Refer Slide Time: 02:42)



```
# find out the expansion
expansion = mapping.get(match) if mapping.get(match) else mapping.get(match.lower())

expansion = first_char + expansion[1:]
return expansion

# using sub method to replace all contractions with their expansions for a sentence
# function expand_map will be called for every non-overlapping occurrence of the pattern
expand_sent = pattern.sub(expand_map, sent)
return expand_sent

In [48]: test = "The young man wasn't fast enough and he couldn't win the race"
test_e = expand_contractions(test, contraction_map)
test_e

Out[48]: 'The young man was not fast enough and he could not win the race'

In [45]: # using user-defined expand_contractions function on cleaned text
emma_fe = [expand_contractions(sentence, contraction_map) for sentence in emma_f]

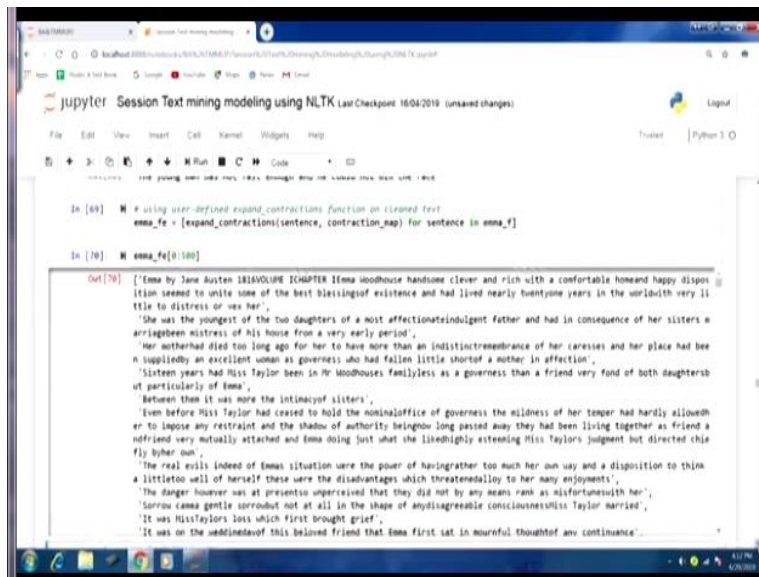
In [46]: emma_fe[0:100]

Out[46]: ['Emma by Jane Austen 1816VOLUME ICHAPTER IEmma Woodhouse handsome clever and rich with a comfortable homeand happy dispositio
tion seemed to unite some of the best blessingsof existence and had lived nearly twentyone years in the worldwith very li
ttle to distress or vex her',
'She was the youngest of a most affectionateindulgent father and had in consequence of her sisters m
arriagebeen mistress of his house from a very early period;']
```

Now the same thing we can apply in our this you know this text document that is Emma file that we have. So, we can use this user define expand underscore contraction function on clean text M

underscore F that we had cleaned in previous lectures.

(Refer Slide Time: 03:00)



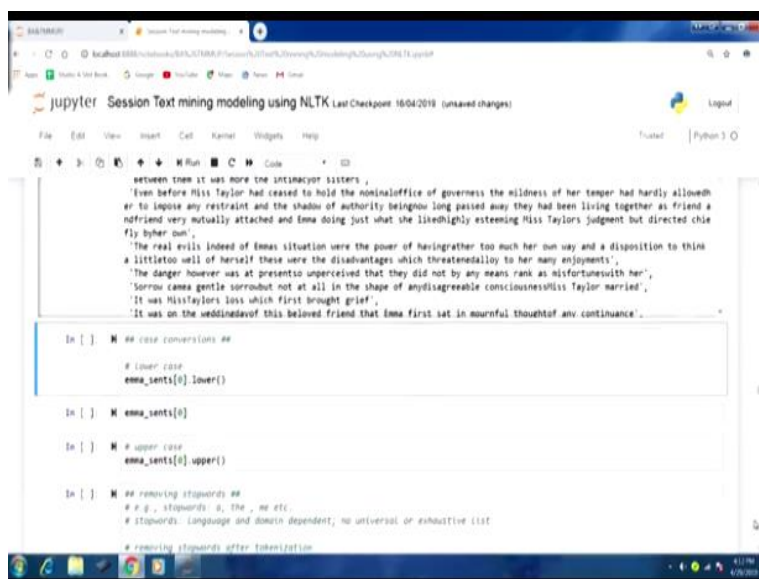
```
In [69]: # using user-defined expand_contractions function on cleaned text
emma_fe = [expand_contractions(sentence, contraction_map) for sentence in emma_f]

In [70]: # emma_fe[0:100]

Out[70]: ['Emma by Jane Austen 1814VOLUME ICHAPTER IEmma Woodhouse handsome clever and rich with a comfortable home and happy dispositi
tion seemed to unite some of the best blessings of existence and had lived nearly twentyone years in the world with very li
ttle to distress or vex her',
'She was the youngest of the two daughters of a most affectionate indulgent father and had in consequence of her sisters m
arriage been mistress of his house from a very early period',
'Her mother had died too long ago for her to have more than an indistinct remembrance of her caresses and her place had bee
n supplied by an excellent woman as governess who had fallen little short of a mother in affection',
'Sixteen years had Miss Taylor been in Mr Woodhouses family less as a governess than a friend very fond of both daughters b
ut particularly of Emma',
'Between them it was more the intimacy of sisters',
'Even before Miss Taylor had ceased to hold the nominal office of governess the mildness of her temper had hardly allowed h
er to impose any restraint and the shadow of authority being long passed away they had been living together as friend a
nd friend very mutually attached and Emma doing just what she liked highly esteeming Miss Taylors judgment but directed chie
fly by her own',
'The real evils indeed of Emmas situation were the power of having rather too much her own way and a disposition to think
a little too well of herself these were the disadvantages which threatened alloy to her many enjoyments',
'The danger however was at present unperceived that they did not by any means rank as misfortunes with her',
'Sorrow came gentle surround not at all in the shape of any disagreeable consciousness Miss Taylor married',
'It was Miss Taylors loss which first brought grief',
'It was on the wedding day of this beloved friend that Emma first sat in mournful thought of any continuance']
```

So, if I run this and you can see now we have Emma underscore fe after this expansion process. And if you will have a look at the first you know 500 you in the sentences here, so many places you would we you know if you compare this output with the am Emma sentence output.

(Refer Slide Time: 03:19)



```
In [ ]: # case conversions ##
# lower case
emma_sents[0].lower()

In [ ]: # emma_sents[0]

In [ ]: # upper case
emma_sents[0].upper()

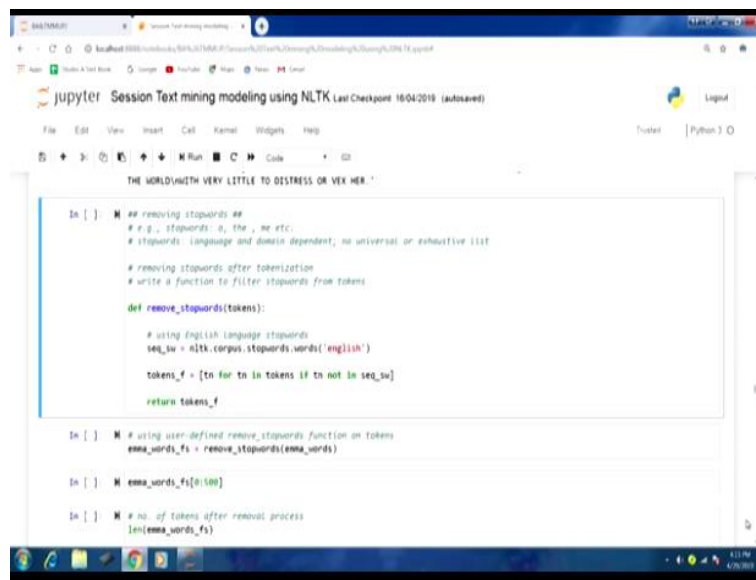
In [ ]: # removing stopwords ##
# e.g., stopwords: a, the, we etc.
# stopwords: language and domain dependent; no universal or exhaustive list
# removing stopwords after tokenization
```

Then you would be able to identify that many places where you know contractions were present, now they have been map to expansions. So, in this fashion in one go we can actually perform this, now let us move to the next aspect that is case convergence. So, sometimes if you might require you know certain case convergence to be performed. So, for that we can just take the you

know we have taken this example if we want lower cases.

So, we can type like this M underscore sense and 0 just for an example you know in the first sentence and dot lower, so we are calling this lower method and it will convert all the you know text into lower case.

(Refer Slide Time: 04:01)

A screenshot of a Jupyter Notebook interface. The title bar says "Jupyter Session Text mining modeling using NLTK (Last Checkpoint: 16/04/2019 (auto-saved))". The notebook contains several code cells. The first cell defines a function to remove stopwords from a list of tokens. The second cell applies this function to a list of tokens. The third cell prints the first 100 tokens of the resulting list. The fourth cell prints the length of the resulting list.

```
In [ ]: ## removing stopwords ##
# e.g., stopwords: a, the, we etc.
# stopwords: language and domain dependent; no universal or exhaustive list

# removing stopwords after tokenization
# write a function to filter stopwords from tokens

def remove_stopwords(tokens):
    # using English language stopwords
    seq_sw = nltk.corpus.stopwords.words('english')
    tokens_f = [tn for tn in tokens if tn not in seq_sw]
    return tokens_f

In [ ]: # using user-defined remove_stopwords function on tokens
emma_words_fs = remove_stopwords(emma_words)

In [ ]: # emma_words_fs[0:100]

In [ ]: # no. of tokens after removal process
len(emma_words_fs)
```

If I run this will have this you know convert into lower case, this is we can compare, similarly for upper also similarly you can see all the text can be converted into the upper case as well. Let us move forward now I will talk about another aspect that is removing stop words, so what we mean by stop word essentially it is about the terms like a dha, dhe, me. So, these words which are not you know really meaningful in the text analytics context in the test you know classification you know context.

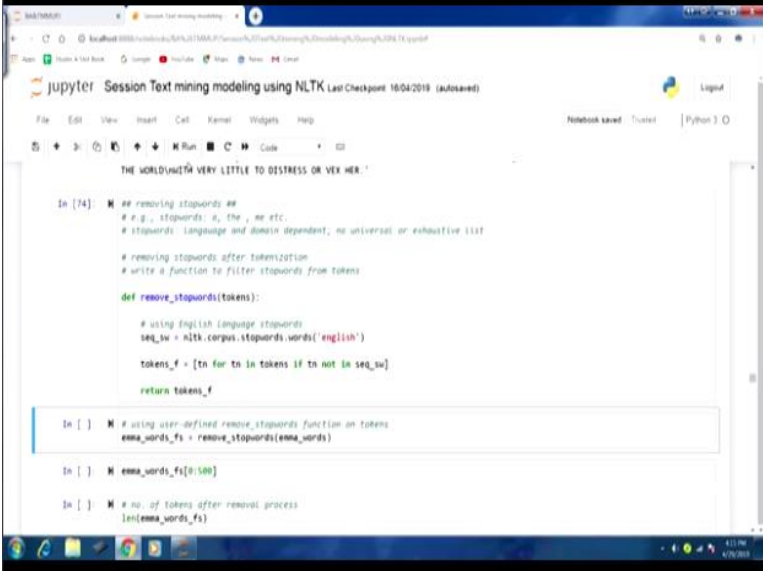
So, we might prefer to remove these terms and you know because we would like to later on I will discuss will like to identify relevant features, relevant vectors which will help us in our text classification problems. So, therefore some of these stop words which are not going to add much in that context we would like to get rid of these terms. So, stop words again the regional language and domain dependent, so depending on the you know language English typically you know these stop words differ from language to language and also domain you know.

So, certain domains, certain words might be meaningful and certain you know and another domain they might lose you know importance there. So, however we do not have any universal or exhaustive list in this case still we have for English language will use this NLTK you know module. Here we have a list of stop words, so that will be using to get rid of these words from our corpus.

So, typically removing stop words after tokenization because it makes the process slightly easier for us. So, we will talk about writing a function to filter stop words from tokens you can see now we are defining the function `remove` and of course stop words tokens will be passing tokens here and then we will use English language stop words here. So, you can see we are creating this sequence underscore SW this list of the stop words and NLTK dig it dot corpus dot stop words dot words and English.

So, we will have the list here and then we are running this you know list comprehension tokens underscore F. So, we are filtering these tough words, so that is why underscore F, so tn, so here you can see tn for tn in tokens if tn not in sequence underscore sw. So, if a particular token is not already present in the list of stop words then you know will have that and our filtered tokens list, so this is the functions.

(Refer Slide Time: 06:53)



```
THE WORLD WITH VERY LITTLE TO DISTRESS OR VEX HER.'

In [74]: ## removing stopwords ##
# e.g., stopwords: a, the, we etc.
# stopwords: language and domain dependent, no universal or exhaustive list
# removing stopwords after tokenization
# write a function to filter stopwords from tokens

def remove_stopwords(tokens):
    # using english language stopwords
    seq_sw = nltk.corpus.stopwords.words('english')
    tokens_f = [tn for tn in tokens if tn not in seq_sw]
    return tokens_f

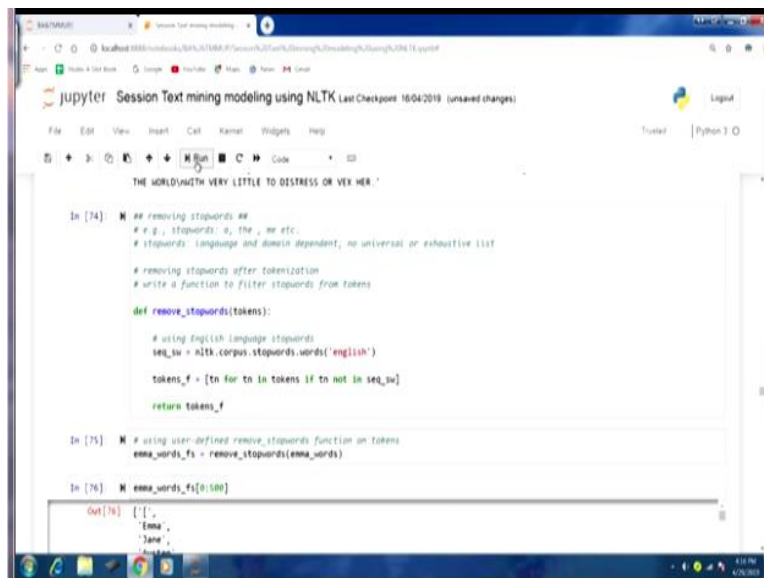
In [ ]: # using user-defined remove_stopwords function on tokens
emma_words_f = remove_stopwords(emma_words)

In [ ]: emma_words_f[0:500]

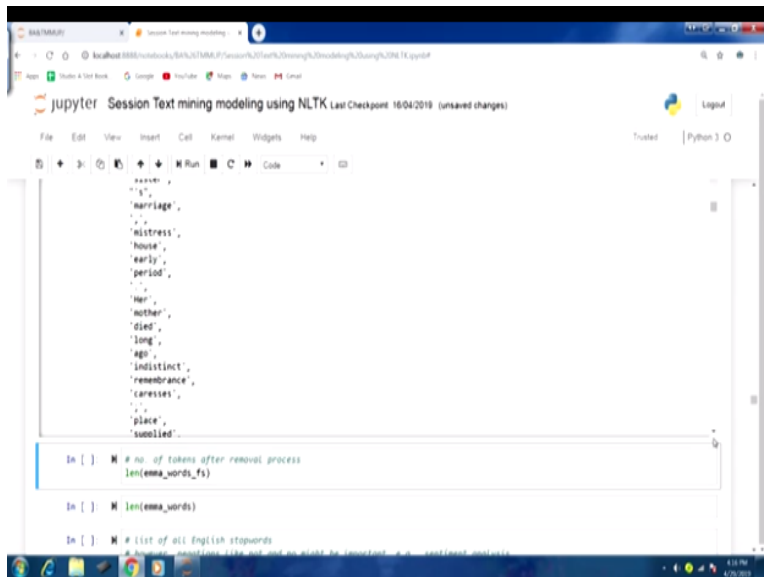
In [ ]: # no. of tokens after removal process
len(emma_words_f)
```

So, if I run this and now we can apply this particular function in our Emma underscore words, so

(Refer Slide Time: 07:10)

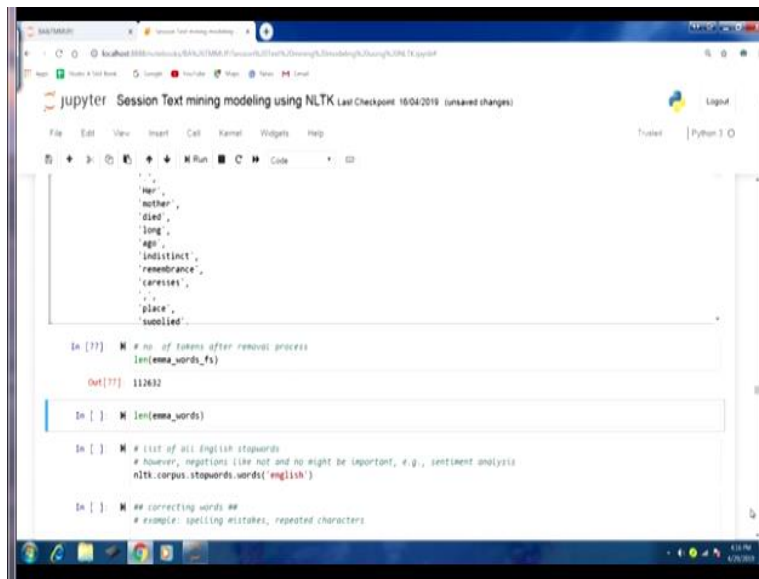


(Refer Slide Time: 07:16)



And in the output if I scroll through this you can see if you look at Emma, Jane, Austen 1816 and you would see that all the you know stop words a and these kind of terms are gone. If I scroll through this you would see only meaningful words have remained and you know so all others are gone. So, let us have a look at the number of tokens after removal process, so we can use this learn function.

(Refer Slide Time: 07:51)

A screenshot of a Jupyter Notebook interface. The title bar says "jupyter Session Text mining modeling using NLTK Last checkpoint: 16/04/2019 (unsaved changes)". The code cell contains a list of words: 'her', 'mother', 'died', 'long', 'ago', 'indistinct', 'remembrance', 'careless', 'place', 'suicided'. Below the list, the code cell shows:

```
In [77]: # no. of tokens after removal process  
len(emma_words_fs)
```

 The output cell shows:

```
Out[77]: 112632
```

 The next code cell shows:

```
In [ ]: len(emma_words)
```

 The next code cell shows:

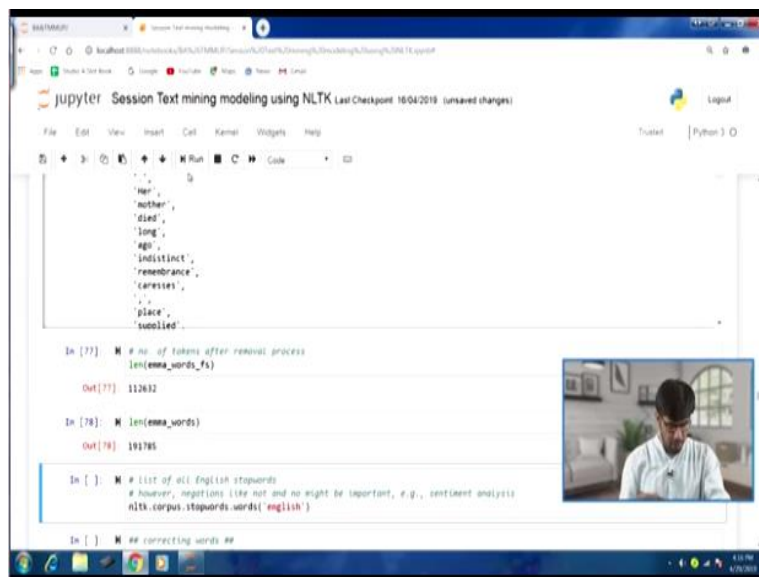
```
In [ ]: # list of all English stopwords  
# however, negations like not and no might be important, e.g., sentiment analysis  
nltk.corpus.stopwords.words('english')
```

 The next code cell shows:

```
In [ ]: # correcting words ##  
# example: spelling mistakes, repeated characters
```

And you can see the number of words after this stop word removal process we have one lakh twelve thousand six hundred thirty-two.

(Refer Slide Time: 07:58)

A screenshot of a Jupyter Notebook interface, similar to the previous one. The code cell shows the same list of words. Below the list, the code cell shows:

```
In [77]: # no. of tokens after removal process  
len(emma_words_fs)
```

 The output cell shows:

```
Out[77]: 112632
```

 The next code cell shows:

```
In [78]: len(emma_words)
```

 The output cell shows:

```
Out[78]: 191785
```

 The next code cell shows:

```
In [ ]: # list of all English stopwords  
# however, negations like not and no might be important, e.g., sentiment analysis  
nltk.corpus.stopwords.words('english')
```

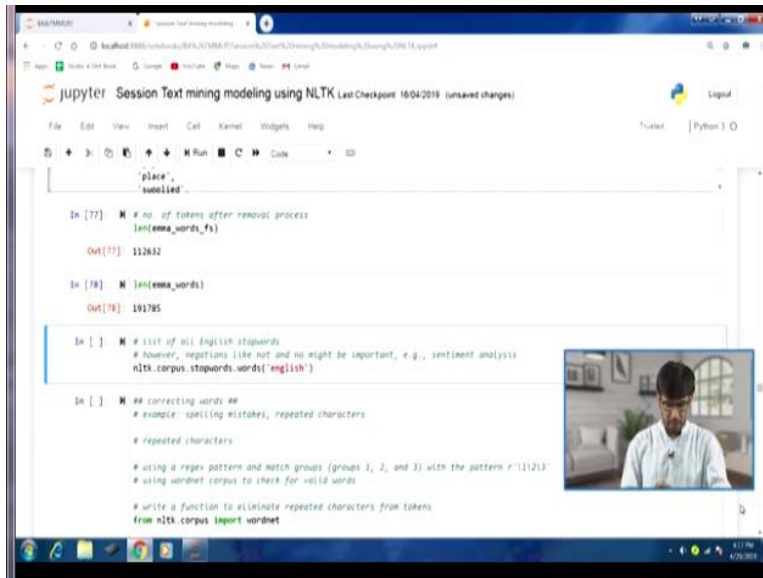
 The next code cell shows:

```
In [ ]: # correcting words ##  
# example: spelling mistakes, repeated characters
```

 A small video inset in the bottom right corner shows a man speaking.

And if we look at the original list here, so you can see my one lakh ninety one thousand, so approximately about you know eighty thousand words we have lost just by removing stop words, this is the reduction in size that we have been able to perform.

(Refer Slide Time: 08:17)



A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK". The notebook shows the following code and output:

```
place',
'sweepled',

In [77]: # # no. of tokens after removal process
         len(emma_words_fs)

Out[77]: 112632

In [78]: # len(emma_words)

Out[78]: 191785

In [ ]: # # list of all English stopwords
         # however, negations like not and no might be important, e.g., sentiment analysis
         nltk.corpus.stopwords.words('english')

In [ ]: # # correcting words ##
         # example: spelling mistakes, repeated characters

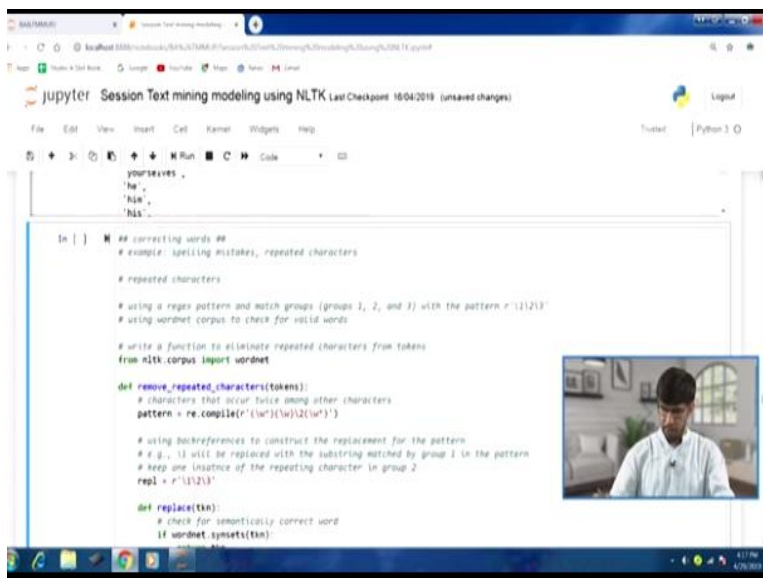
         # repeated characters

         # using a regex pattern and match groups (groups 1, 2, and 3) with the pattern r'(\1|2|3)'
         # using wordnet corpus to check for valid words

         # write a function to eliminate repeated characters from tokens
         from nltk.corpus import wordnet
```

Now if you are interested in having a look at the list of all English stop words, so you know you can again print this list `NLTK dot corpus dot stop words dot words English`.

(Refer Slide Time: 08:30)



A screenshot of a Jupyter Notebook showing the following code:

```
# # correcting words ##
# example: spelling mistakes, repeated characters

# repeated characters

# using a regex pattern and match groups (groups 1, 2, and 3) with the pattern r'(\1|2|3)'
# using wordnet corpus to check for valid words

# write a function to eliminate repeated characters from tokens
from nltk.corpus import wordnet

def remove_repeated_characters(tokens):
    # characters that occur twice among other characters
    pattern = re.compile(r'(\w)(\1|2|3)')

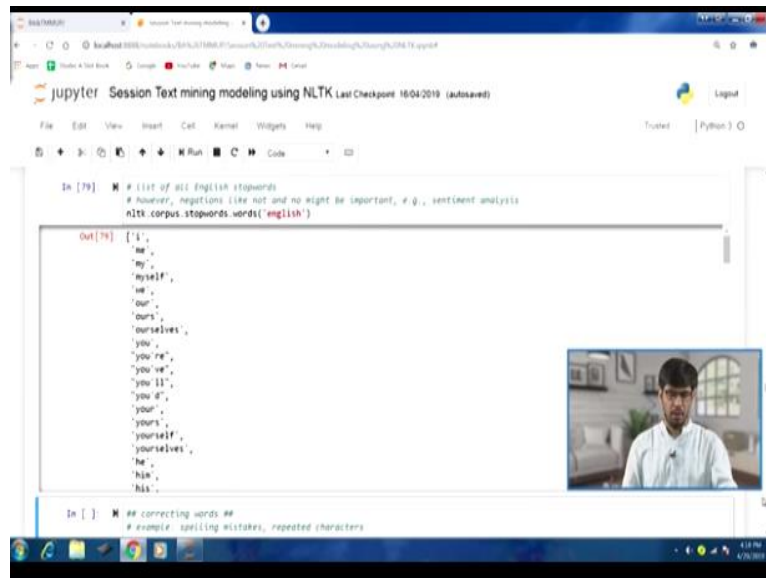
    # using backreferences to construct the replacement for the pattern
    # e.g., \1 will be replaced with the substring matched by group 1 in the pattern
    # keep one instance of the repeating character in group 2
    repl = r'\1\2\3'

    def replace(tkn):
        # check for semantically correct word
        if wordnet.synsets(tkn):
```

If I run this yeah and if I go back to the output you can see these are these could be some of these stop words these are some of the listed stop words for English language you can see I, me, my, myself, we, our, hours. So, these are some of the words which are not which might not be meaningful for our text and analytics context or text classification problems. So, these are the stopper that we typically would like to remove, so the same thing I mentioned in the comment part also.

So, however there might be certain exceptions also for example negations like not and no they might be important sometimes, specifically we are applying if we are you know doing sentiment analysis. So, there sentiment analysis would also involve you know we will have negation aspect.

(Refer Slide Time: 09:23)

A screenshot of a Jupyter Notebook interface. The browser address bar shows a local file path. The notebook title is "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (autosaved)". The code cell contains a comment about negations and a call to `nltk.corpus.stopwords.words('english')`. The output cell shows a list of 25 common English stop words. A small video inset in the bottom right corner shows a man speaking.

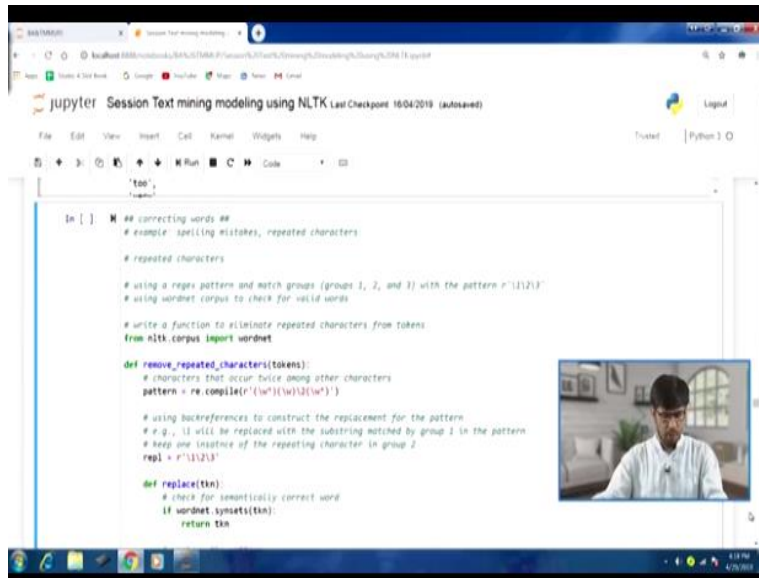
```
In [79]: # list of all English stopwords
# however, negations like not and no might be important, e.g., sentiment analysis
nltk.corpus.stopwords.words('english')

Out[79]: ['I',
'ae',
'ay',
'myself',
'ed',
'our',
'ours',
'ourselves',
'you',
'you\'re',
'you\'re',
'you\'ll',
'you\'d',
'your',
'yours',
'yourself',
'yourselves',
'he',
'his',
'his']

In [ ]: ## correcting words ##
# example: spelling mistakes, repeated characters
```

So, therefore these words might become important there, so we will have to really see you know whether you know we can really remove you know these you know stop words whether we can afford no and not in these stop words .So, right scroll through we might encounter however if that is the case you can see here no, nor, not here. So, they are right now in the stop words but if we are performing sentiment analysis probably will be better off you know not removing stop words. Because he might lose no were not or we can prepare our own list, own stop word list.

(Refer Slide Time: 10:03)



```

In [ ]: ## correcting words ##
# example: spelling mistakes, repeated characters

# repeated characters

# using a regex pattern and match groups (groups 1, 2, and 3) with the pattern r'11213'
# using wordnet corpus to check for valid words

# write a function to eliminate repeated characters from tokens
from nltk.corpus import wordnet

def remove_repeated_characters(tokens):
    # characters that occur twice among other characters
    pattern = re.compile(r'(?!(w)(w)(w))')

    # using backreferences to construct the replacement for the pattern
    # e.g., 11 will be replaced with the substring matched by group 1 in the pattern
    # keep one instance of the repeating character in group 2
    repl = r'\1\2\3'

    def replace(tkn):
        # check for semantically correct word
        if wordnet.synsets(tkn):
            return tkn

```

Let us move forward, so another aspect that we might encounter in processing text and transforming text is correcting words. So, there could be spelling mistakes, there could be you know repeated characters sometimes. So, especially informal writing you know if somebody is writing in a social media platform or in the comment section of e-commerce website they might write you know who emphasize you know certain words they might in a repeat certain characters.

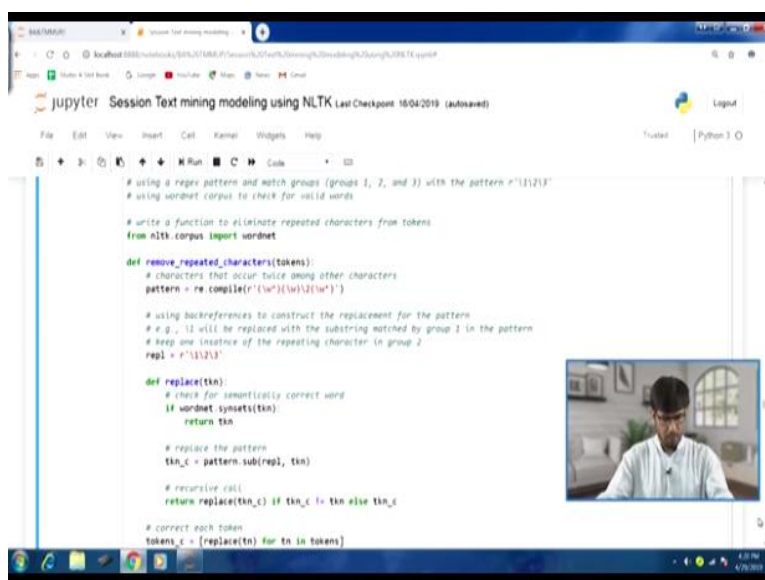
So, for our you know text a analytics purpose if we want to remove those you know we want to you know correct those words, so how we can perform that .So, let us talk about this, so we will pick up the repeated characters scenario, so here we will be using a rejects pattern and you know different match groups 1, 2, 3. So, you know pattern like R you know in within single quotes slash1, slash2, slash3 this is essentially we can use different groups in a pattern.

And we construct another pattern you know that this would be really useful for us in repeated character scenario. So, let us have a look at the function here, now here it will also be important for us to check the check for valid words. So, you know for many of these words for example repeated character for example finally and there we have Y repeating many times or L repeating many times.

So, where do we stop while we are trying to remove some of those repeated characters, so for

that we will take we will have to take the help of word net corpus to actually check whether you know in that removal process whether we have at least valid word you know case.

(Refer Slide Time: 12:00)



```
# using a regex pattern and match groups (groups 1, 2, and 3) with the pattern r'([a-z])\1\2\3'
# using wordnet corpus to check for valid words

# write a function to eliminate repeated characters from tokens
from nltk.corpus import wordnet

def remove_repeated_characters(tokens):
    # characters that occur twice among other characters
    pattern = re.compile(r'([a-z])\1\2\3')

    # using backreferences to construct the replacement for the pattern
    # e.g., \1 will be replaced with the substring matched by group 1 in the pattern
    # keep one instance of the repeating character in group 2
    repl = r'\1\2\3'

    def replace(tkn):
        # check for semantically correct word
        if wordnet.synsets(tkn):
            return tkn

        # replace the pattern
        tkn_c = pattern.sub(repl, tkn)

        # recursive call
        return replace(tkn_c) if tkn_c != tkn else tkn_c

    # correct each token
    tokens_c = [replace(tn) for tn in tokens]
```

So, that will have to perform here, so let us have a look at the definition and definition of this function remove underscore, peter underscore characters. So, we are passing tokens here, so now will be in the pattern that we are going to use here is this characters that occurred twice among other characters. So, if you look at the pattern we have used parentheses to identify different groups here.

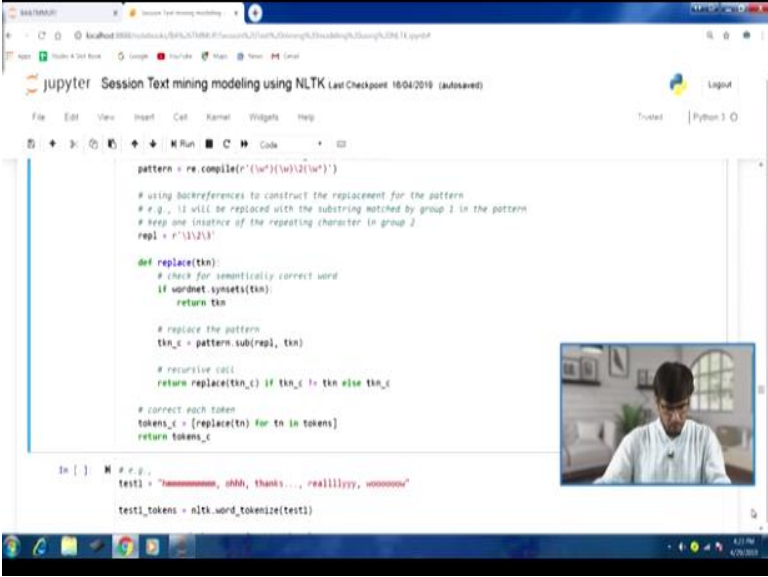
And you can see the bit the second group you know this is slash slash 2 indicating that you know occurring twice. Now these are you know are back differences, so will be using back references to construct the replacement for the pattern. So, one thing is to match a particular you know these kind of repeat words with repeating characters. The another thing is to replace them all eliminate the repeating characters.

So, here when we say slash 1, so slash 1 will be replaced with the substring matched by group 1 in the pattern. So, slash 1 is actually a back reference to the group 1, slash 2 is a back referenced to the group 2. So, when in the in the first line pattern Li dot compiled there you can see slash 2, it is actually referring to it as a back reference to a group 2 which is nothing but slash w, so the same what is repeating here.

So, that is why in the comment section we say characters that occurred twice, so in this in this fashion we have defined and that would be used for matching and then for the replacement. So, if you look at the you know pattern that we are using for replacement but keep one instance of the repeating character in group 2. So, you can see REPL this is our replacement pattern or within single quotes slash 1 that is group 1, then slash 2 that means just you know one instance of that you know that repeating character and then slash 3.

So, slash 1 and slash 3 they are back references the first and the third group they are being kept intact and slash 2 is just once, so just one occurrence we will like to keep here. So, this is how the pattern will use to match and the placement REPL will use to actually replace the matching occurrences.

(Refer Slide Time: 14:19)



```
pattern = re.compile(r'(\w*)(\w)(\w*)')

# using backreferences to construct the replacement for the pattern
# e.g., \1 will be replaced with the substring matched by group 1 in the pattern
# keep one instance of the repeating character in group 2
repl = r'\1\2\3'

def replace(tkn):
    # check for semantically correct word
    if wordnet.synsets(tkn):
        return tkn

    # replace the pattern
    tkn_c = pattern.sub(repl, tkn)

    # recursive call
    return replace(tkn_c) if tkn_c != tkn else tkn_c

# correct each token
tokens_c = [replace(tn) for tn in tokens]
return tokens_c

In [ ]: # e.g.,
test1 = "mmmmmmmm, ohh, thanks..., reallllyyy, nooooooo"
test1_tokens = nltk.word_tokenize(test1)
```

So, now let us have a look at the code you know more detail, so if you come down to this part correct each tokens. So, here you can see we are using a list comprehension we are calling replace function for TN in tokens. So, for each token in the list of tokens that we might be passing will be calling this replace you know function here. So, let us have a look at replace now.

So, once we receive a token will check in the world net whether that is syntactically correct word or not, so it might be so. So, we would like to stop at that point of time, so if world net dot sin

sets and in the begin the balance is TK and token then will return token. If it is true then we will get an token that means we have arrived at a syntactically correct word if not then we will like to replace the pattern.

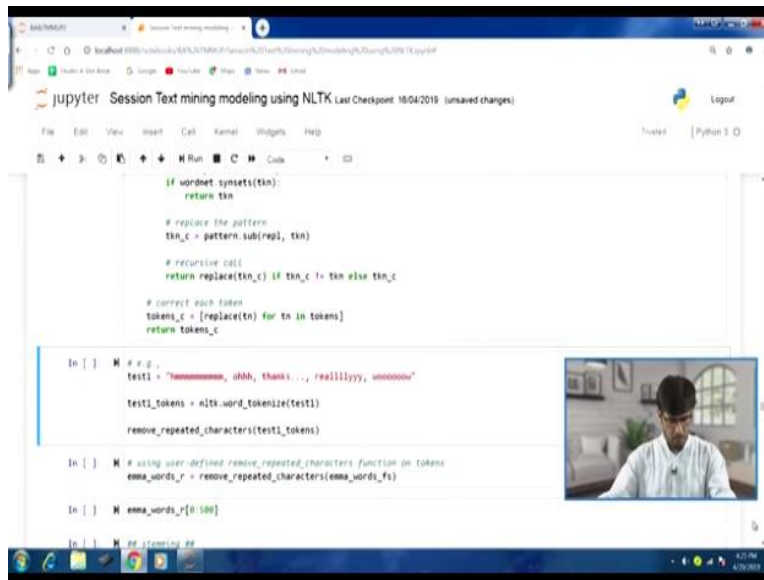
So, TK and underscore C and we are calling you know using this rejects object pattern dot sub and REPL will be use to replace. So, all the you know matching occurrences of the pattern they are going to be replaced with REPL in that particular token TK. Now then we are making a recursive call here, so you can see if token underscore C is not equal to token, in that case you know we are again calling replace TK and underscore C.

And if happens to be if it you know happens to be you know equal to then will written TK and underscore C itself. So, in this fashion recursively because 1 character might be repeated many times, so once we remove you know let us say it is repeating 4 times or we removed once. Then again we will you know we will check whether it is you know equal or not if it is you know not equal then again we will call replace.

In that fashion we will you know go on in this fashion till we you know arrive at a situation where in a recursive call the first if that we have Warner dot sin sets there, we are able to reach the syntactically correct word. So, somewhere this written will help us these decent syntactically correct word and that is the place where this if TK and underscore C not equal to token would actually become you know false and the else part would be written, so it will end this you know record C loop there.

So, this is how this code will work to actually remove the repeated characters, so let me run this.

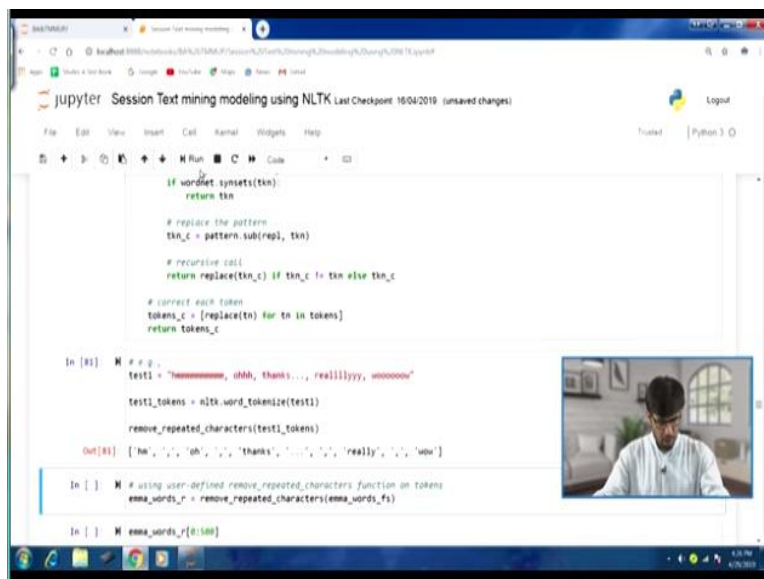
(Refer Slide Time: 17:01)



```
def remove_repeated_characters(token):  
    if wordnet.synsets(token):  
        return token  
  
    # replace the pattern  
    token_c = pattern.sub(repl, token)  
  
    # recursive call  
    return remove_repeated_characters(token_c) if token_c != token else token_c  
  
# correct each token  
tokens_c = [remove_repeated_characters(token) for token in tokens]  
return tokens_c  
  
In [ ]: # e.g.  
test1 = "mmmmmmmm, ohh, thanks ..., reallllyyy, unooooow"  
test1_tokens = nltk.word_tokenize(test1)  
remove_repeated_characters(test1_tokens)  
  
In [ ]: # using user-defined remove_repeated_characters function on tokens  
emma_words_r = remove_repeated_characters(emma_words_fs)  
  
In [ ]: # emma_words_r[0:100]
```

And let us have a let us take an example here, so here you can see test 1, this is the you know sentence this we are taking you can see all the words they are having repeated character. And then we are calling our you know we are tokenizing this you know sentence here, you are calling word underscore tokenize, then we are passing you know this token to our user our you know function that we have just defined remove underscore repeated characters.

(Refer Slide Time: 17:36)



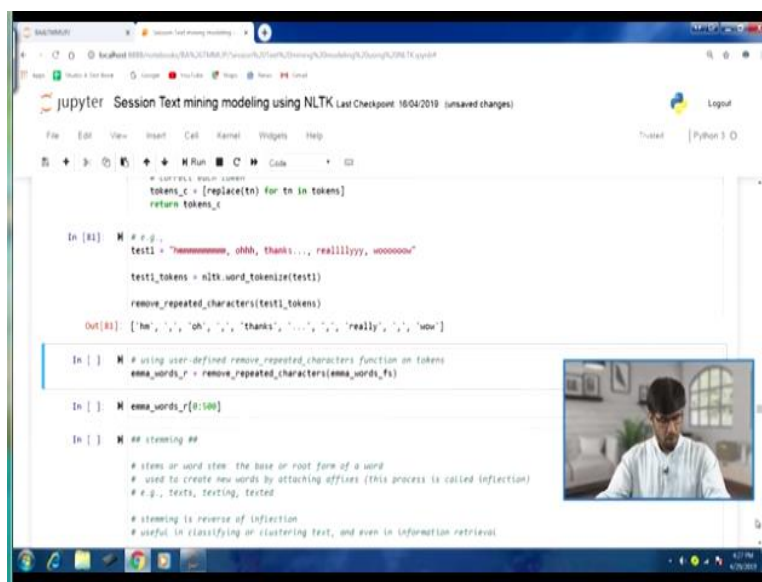
```
In [81]: # e.g.  
test1 = "mmmmmmmm, ohh, thanks ..., reallllyyy, unooooow"  
test1_tokens = nltk.word_tokenize(test1)  
remove_repeated_characters(test1_tokens)  
  
Out[81]: ['hm', ',', 'oh', ',', 'thanks', ',', '...', ',', 'really', ',', 'uow']  
  
In [ ]: # using user-defined remove_repeated_characters function on tokens  
emma_words_r = remove_repeated_characters(emma_words_fs)  
  
In [ ]: # emma_words_r[0:100]
```

So, if I run this and you can have a look at the output 81 you can see you know HM there is because we would not reach a syntactically correct word there. So, only will this you know function will be called recursive call will keep on happening till we are left with just one M. In the second case O triple H will be left with OH because you know that is also you know that you

know again same example it might not be syntactically correct word.

Then third one thanks, you can see there dot is repeating thrice, so that is gone and thanks is syntactically correct words, so it will stop there. Similarly really you can see you know 2 Ys and you know two L's are gone because only then will be able to reach the syntactically correct word similarly wall also, there also you know those in extra Os are gone. So, in this fashion our you know you can see that the function that we have justified is working quite well in terms of removing repeated characters.

(Refer Slide Time: 18:43)



```
def remove_repeated_characters(tokens):
    tokens_c = [replace(tn) for tn in tokens]
    return tokens_c

In [81]: # P.S.
test1 = "hmmmmmm, ohhh, thanks..., realllllyy, woowoooo"
test1_tokens = nltk.word_tokenize(test1)
remove_repeated_characters(test1_tokens)

Out[81]: ['hm', '.', 'h', 'oh', 'h', 'thanks', '...', 'r', 'e', 'a', 'l', 'l', 'y', 'y', 'w', 'o', 'o', 'w']

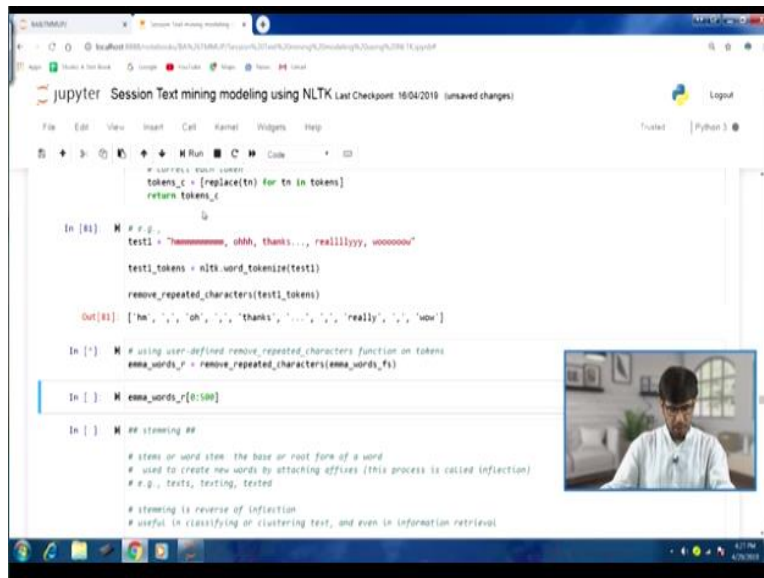
In [ ]: # using user-defined remove_repeated_characters function on tokens
emma_words_r = remove_repeated_characters(emma_words_fs)

In [ ]: # emma_words_r[0:100]

In [ ]: # stemming ##
# stems or word stem: the base or root form of a word
# used to create new words by attaching affixes (this process is called inflection)
# e.g., texts, testing, tested
# stemming is reverse of inflection
# useful in classifying or clustering text, and even in information retrieval
```

Now let us talk about let us apply this remove underscore repeated characters on Emma underscore words underscore FS as well.

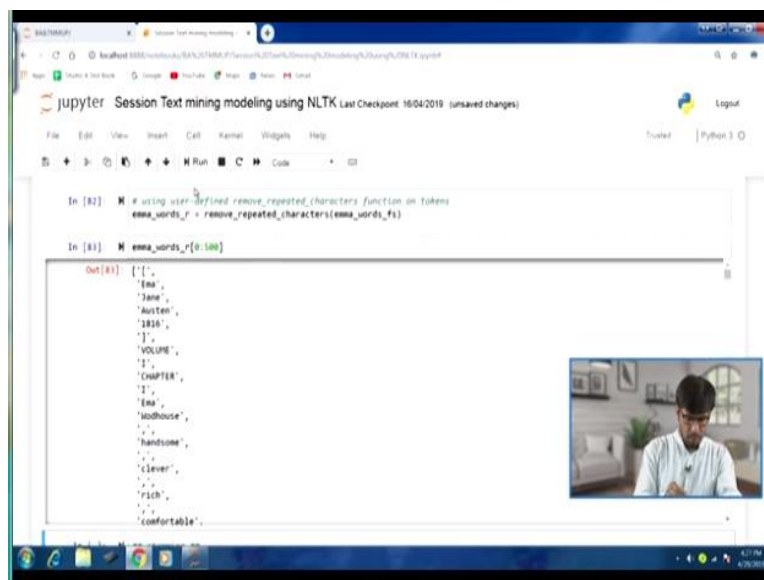
(Refer Slide Time: 18:54)



```
def remove_repeated_characters(tokens):  
    # corrects token content  
    tokens_c = [replace(tn) for tn in tokens]  
    return tokens_c  
  
In [81]: # e.g.,  
test1 = "hmmmmmm, ohh, thanks..., reallllyy, woowoooo"  
test1_tokens = nltk.word_tokenize(test1)  
remove_repeated_characters(test1_tokens)  
Out[81]: ['hm', ',', 'oh', ',', 'thanks', ',', '...', ',', 'really', ',', 'wow']  
  
In [ ]: # using user-defined remove_repeated_characters function on tokens  
emma_words_r = remove_repeated_characters(emma_words_fs)  
  
In [ ]: # emma_words_r[0:100]  
  
In [ ]: # stemming ##  
# stems or word stem: the base or root form of a word  
# used to create new words by attaching affixes (this process is called inflection)  
# e.g., tests, testing, tested  
# stemming is reverse of inflection  
# useful in classifying or clustering text, and even in information retrieval
```

So, let me run this and if there are any such instances in this particular you know text corpus then those would also be corrected.

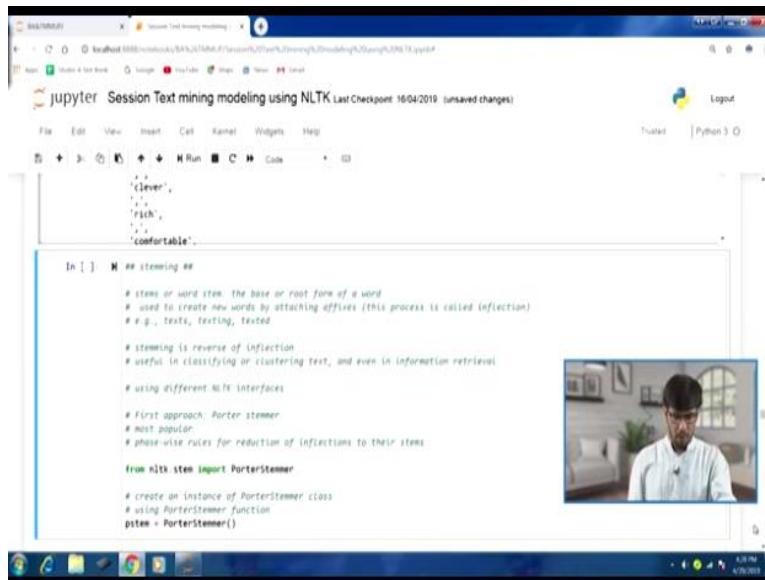
(Refer Slide Time: 19:07)



```
In [82]: # using user-defined remove_repeated_characters function on tokens  
emma_words_r = remove_repeated_characters(emma_words_fs)  
  
In [83]: # emma_words_r[0:100]  
Out[83]: ['.',  
          'Em',  
          'Jane',  
          'Austen',  
          '1818',  
          ']',  
          'VOLUME',  
          ']',  
          'CHAPTER',  
          ']',  
          'Em',  
          'Madhouse',  
          ',',  
          'handsome',  
          ',',  
          'clever',  
          ',',  
          'rich',  
          ',',  
          'comfortable',  
          '']
```

But because Emma was part of formal writing, so I do not expect this to happen in that case but anyway we can call this function.

(Refer Slide Time: 19:16)

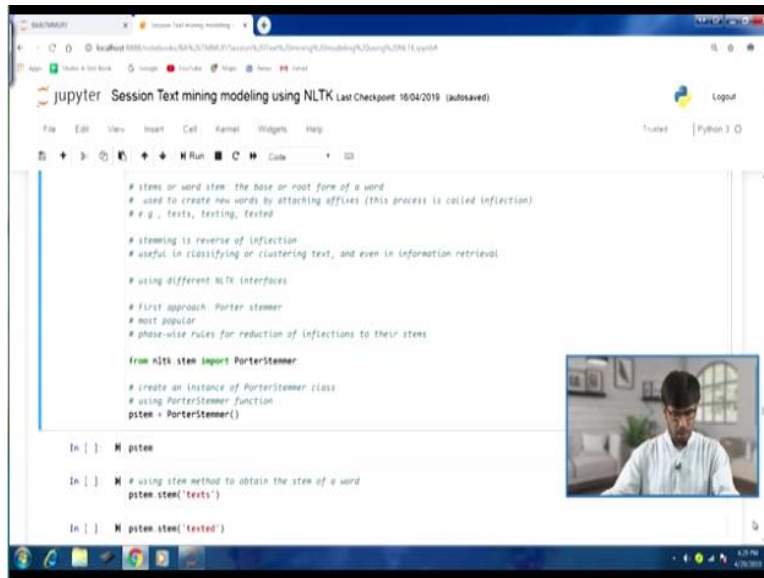


So, with this let us move onto the next aspect that is very important for text you know processing that is a stemming. So, what is stem or word you know or word stem, so essentially what we mean by this process stemming is we would like to reach the base or root form of a word. Because in you know there could be multiple variants of a root form of a word.

For example text we might in writing we might have texts, texting, texted, so we would like to just keep the root form of the word. So, essentially when we say stem or word stem we are essentially referring to that particular root form of the word. So, this might this kind of processing we are able to perform this kind of stemming in our text corpus, it will really help in our text analytics and text you know classification context, we will talk this in more detail.

So, stemming is all stemming can also be considered as a reverse of inflection, so as we said useful in classifying or clustering text and even in sometimes in information retrieval. There are so many available you know an NLTK interfaces for this, so we will talk about the first approach which is porter stemmer, porter algorithm this is the most popular approach. So, in this one we perform phase wise you know rules, we use phase wise rules for reduction of inflection to their stems.

(Refer Slide Time: 20:53)

A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (autosaved)". The notebook contains several code cells. The first cell is a comment block explaining the purpose of the code: "# stems or word stem: the base or root form of a word", "# used to create new words by attaching affixes (this process is called inflection)", "# e.g., tests, testing, tested", "# stemming is reverse of inflection", "# useful in classifying or clustering text, and even in information retrieval", "# using different NLTK interfaces", "# First approach: Porter stemmer", "# most popular", "# phrase-wise rules for reduction of inflections to their stems". The second cell imports the PorterStemmer class from the nltk.stem module: "from nltk.stem import PorterStemmer". The third cell creates an instance of the PorterStemmer class: "pstem = PorterStemmer()". The fourth cell uses the stem method to obtain the stem of the word 'texts': "pstem.stem('texts')". The fifth cell uses the stem method to obtain the stem of the word 'texted': "pstem.stem('texted')".

```
# stems or word stem: the base or root form of a word
# used to create new words by attaching affixes (this process is called inflection)
# e.g., tests, testing, tested

# stemming is reverse of inflection
# useful in classifying or clustering text, and even in information retrieval

# using different NLTK interfaces

# First approach: Porter stemmer
# most popular
# phrase-wise rules for reduction of inflections to their stems

from nltk.stem import PorterStemmer

# create an instance of PorterStemmer class
# using PorterStemmer function
pstem = PorterStemmer()

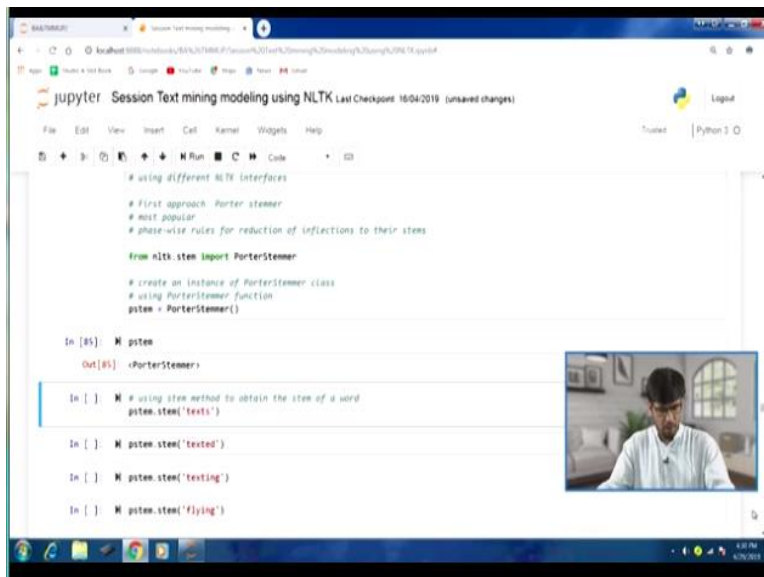
In [ ]: pstem

In [ ]: # using stem method to obtain the stem of a word
pstem.stem('texts')

In [ ]: pstem.stem('texted')
```

So, let us have a look at the code here, so for this we will have to import this from this and LD hit order NLTK dot stem module will have to import this class porter stemmer and then will define an instance of this porter stemmer class and let us run this.

(Refer Slide Time: 21:12)

A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (unsaved changes)". The notebook contains several code cells. The first cell is a comment block explaining the purpose of the code: "# using different NLTK interfaces", "# First approach: Porter stemmer", "# Most popular", "# phrase-wise rules for reduction of inflections to their stems". The second cell imports the PorterStemmer class from the nltk.stem module: "from nltk.stem import PorterStemmer". The third cell creates an instance of the PorterStemmer class: "pstem = PorterStemmer()". The fourth cell uses the stem method to obtain the stem of the word 'texts': "pstem.stem('texts')". The fifth cell uses the stem method to obtain the stem of the word 'texted': "pstem.stem('texted')". The sixth cell uses the stem method to obtain the stem of the word 'texting': "pstem.stem('texting')". The seventh cell uses the stem method to obtain the stem of the word 'flying': "pstem.stem('flying')".

```
# using different NLTK interfaces

# First approach: Porter stemmer
# Most popular
# phrase-wise rules for reduction of inflections to their stems

from nltk.stem import PorterStemmer

# create an instance of PorterStemmer class
# using PorterStemmer function
pstem = PorterStemmer()

In [85]: pstem
Out[85]: <PorterStemmer>

In [ ]: # using stem method to obtain the stem of a word
pstem.stem('texts')

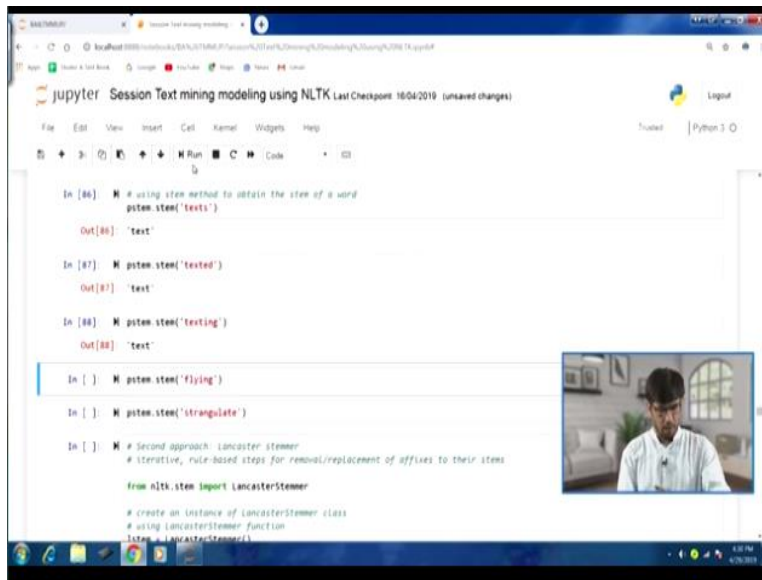
In [ ]: pstem.stem('texted')

In [ ]: pstem.stem('texting')

In [ ]: pstem.stem('flying')
```

If you look at the PSTM you can see this is porter stemmer instance of porter stemmer class. Now we can call this you know stem method using this instance and pass on the you know pass on the examples here for example text or textured or texting and it will give us the you know route form of the word.

(Refer Slide Time: 21:36)



A Jupyter Notebook interface titled "Session Text mining modeling using NLTK" with a last checkpoint of 16/04/2019. The notebook contains several code cells. The first three cells demonstrate the Porter stemmer's output for the words 'text', 'texted', and 'texting', all of which return 'text'. The fourth cell shows the output for 'flying' and 'strangulate', which also return 'text'. The fifth cell introduces the Lancaster stemmer as a second approach, showing the import of the class and the creation of an instance.

```
In [86]: # using stem method to obtain the stem of a word
         portem.stem('text')
Out[86]: 'text'

In [87]: # portem.stem('texted')
         portem.stem('texted')
Out[87]: 'text'

In [88]: # portem.stem('texting')
         portem.stem('texting')
Out[88]: 'text'

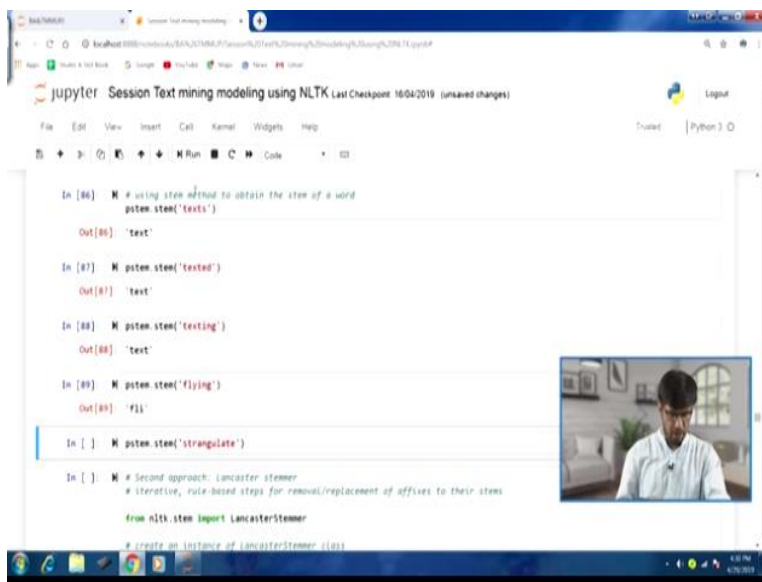
In [ ]: # portem.stem('flying')
         portem.stem('flying')
Out[ ]: 'text'

In [ ]: # portem.stem('strangulate')
         portem.stem('strangulate')
Out[ ]: 'text'

In [ ]: # Second approach: Lancaster stemmer
         # iterative, rule-based steps for removal/replacement of affixes to their stems
         from nltk.stem import LancasterStemmer
         # create an instance of LancasterStemmer class
         # using LancasterStemmer function
         lstem = LancasterStemmer()
```

So, if I run this you can see texted to get text texting again you will get text, now if I take an example like flying.

(Refer Slide Time: 21:45)



This screenshot shows the same Jupyter Notebook as before, but with the execution of the fifth code cell. The output for 'flying' is now 'fli' and the output for 'strangulate' is 'strangulate'.

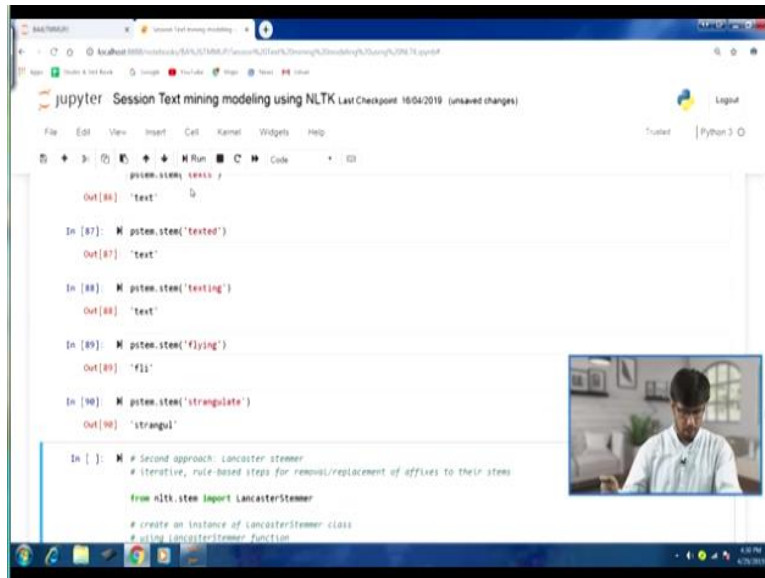
```
In [89]: # portem.stem('flying')
         portem.stem('flying')
Out[89]: 'fli'

In [ ]: # portem.stem('strangulate')
         portem.stem('strangulate')
Out[ ]: 'strangulate'

In [ ]: # Second approach: Lancaster stemmer
         # iterative, rule-based steps for removal/replacement of affixes to their stems
         from nltk.stem import LancasterStemmer
         # create an instance of LancasterStemmer class
```

So, if I run this we get FLI and you know for strangulate.

(Refer Slide Time: 21:53)



A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (unsaved changes)". The notebook is running on a local host. The code in the cells demonstrates the use of Porter's stemming algorithm. The first cell imports the necessary modules and creates a PorterStemmer object. The subsequent cells show the stemming of various words, with the output displayed in a separate window on the right. The words and their stems are: 'text' to 'text', 'texted' to 'text', 'texting' to 'text', 'flying' to 'fly', and 'strangulate' to 'strangul'.

```
from nltk.stem import PorterStemmer
stemmer = PorterStemmer()

In [86]: stemmer.stem('text')
Out[86]: 'text'

In [87]: stemmer.stem('texted')
Out[87]: 'text'

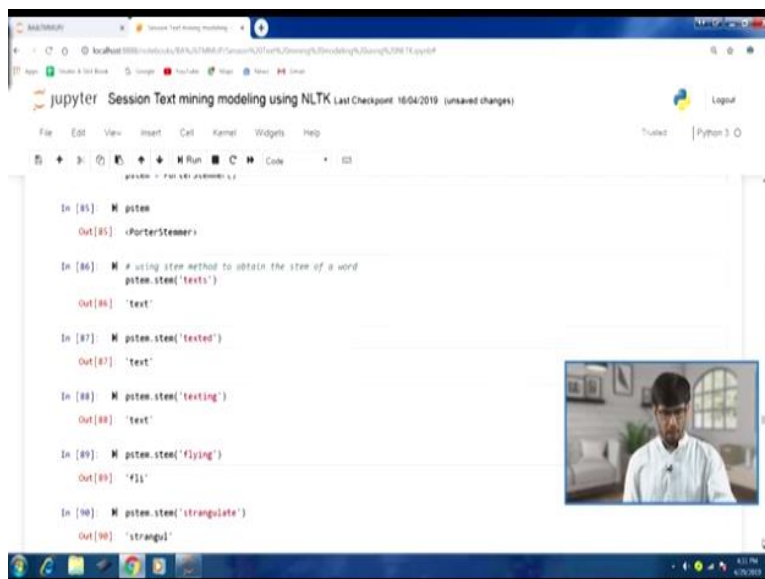
In [88]: stemmer.stem('texting')
Out[88]: 'text'

In [89]: stemmer.stem('flying')
Out[89]: 'fly'

In [90]: stemmer.stem('strangulate')
Out[90]: 'strangul'
```

If I run this you can see strangle, so in this fashion you can see you know different forms of a word would be there in writing in a pub would be there as part of a text. So, how we can use stemming algorithm suppose apple in this case we use Porter's method. So, how we can use these methods to actually reach arrive at a root form of the word and what it will how it will help us in text and analytics.

(Refer Slide Time: 22:20)



A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (unsaved changes)". The notebook is running on a local host. The code in the cells demonstrates the use of Porter's stemming algorithm. The first cell imports the necessary modules and creates a PorterStemmer object. The subsequent cells show the stemming of various words, with the output displayed in a separate window on the right. The words and their stems are: 'text' to 'text', 'texted' to 'text', 'texting' to 'text', 'flying' to 'fly', and 'strangulate' to 'strangul'.

```
from nltk.stem import PorterStemmer
stemmer = PorterStemmer()

In [86]: stemmer.stem('text')
Out[86]: 'text'

In [87]: stemmer.stem('texted')
Out[87]: 'text'

In [88]: stemmer.stem('texting')
Out[88]: 'text'

In [89]: stemmer.stem('flying')
Out[89]: 'fly'

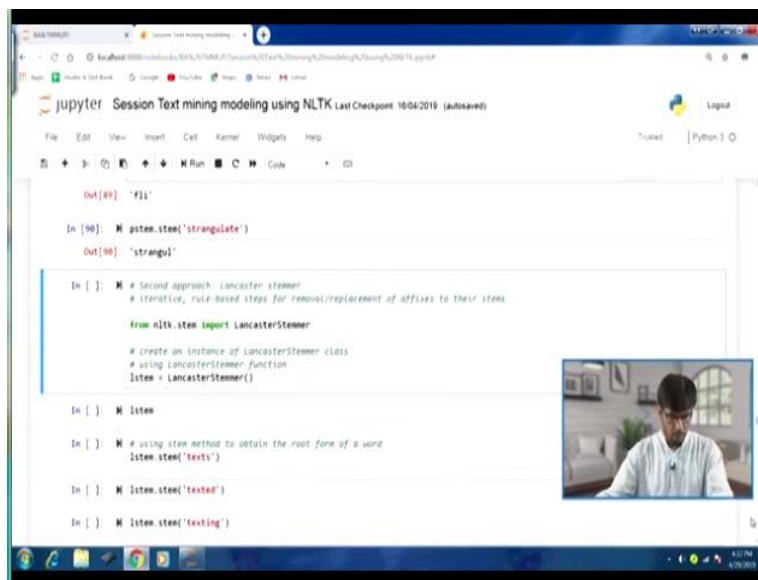
In [90]: stemmer.stem('strangulate')
Out[90]: 'strangul'
```

In that sense that if see if in your text if you have you know if you have words if you have variants like text, texture, texting and all of those instances would be converted into text. So, therefore frequency of text will increase and that will be really helpful in our you know factorization steps or later on you know when we apply a machine learning algorithms, so

because the frequency will increase.

So, and you would see that you know as we have talked about that in our table tabular layout you know will be counting you know for a particular term will be counting how many times it has appear in a particular document. So, that information is important now in this we perform stemming, so it will really help us in terms of having a you know fair distribution of you know root words or you know stems.

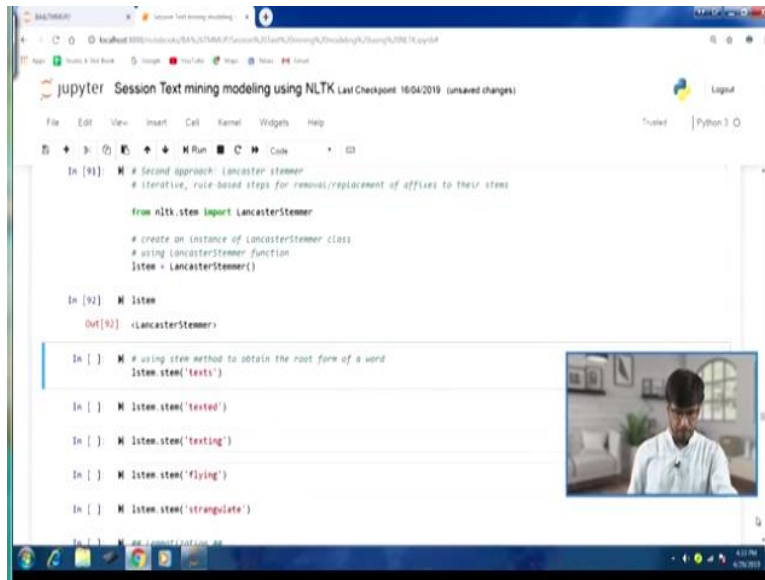
(Refer Slide Time: 23:18)

A screenshot of a Jupyter Notebook interface. The title bar says "jupyter Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (auto-save@0)". The notebook contains several code cells. The first cell shows a simple string replacement: `Out[89]: 'fil'` followed by `In [90]: pstem.stem('strangulate')` and `Out[90]: 'strangul'`. The second cell contains a comment: `# Second approach - Lancaster stemmer` followed by another comment: `# iterative, rule-based steps for removal/replacement of affixes to their stems`. The third cell imports the stemmer: `from nltk.stem import LancasterStemmer`. The fourth cell creates an instance: `# create an instance of LancasterStemmer class` followed by `# using LancasterStemmer function` and `lstem = LancasterStemmer()`. The fifth cell simply prints the instance: `In [ ]: lstem`. The sixth cell shows the `stem` method being used on a list: `In [ ]: # using stem method to obtain the root form of a word` followed by `lstem.stem(['text1'])`. The seventh cell shows the `stem` method being used on a single word: `In [ ]: lstem.stem('texted')`. The eighth cell shows the `stem` method being used on another word: `In [ ]: lstem.stem('texting')`. A small video inset in the bottom right corner shows a person speaking.

So, it will statistically it will help us in terms of distribution of words and therefore you know application of machine learning algorithms will also be you know useful. Now let us talk about the second approach this is the Lancaster's stemmer, so this is the iterative and follows rule-based steps for removal or replacement of affixes to their stems. So, you know when we talk about root form of root form of a word or stem.

So, there are going to be affixes there prefixes or suffixes and different variants of that words would be there in the writing. So, this particular you know algorithm you know is based on removal or replacement of these affixes. So, let us import this Lancaster stemmer class and we will be using an instance of this class.

(Refer Slide Time: 24:17)



```
# Second approach: Lancaster stemmer
# iterative, rule-based steps for removal/replacement of affixes to their stems

from nltk.stem import LancasterStemmer

# create an instance of LancasterStemmer class
# using LancasterStemmer function
lstem = LancasterStemmer()

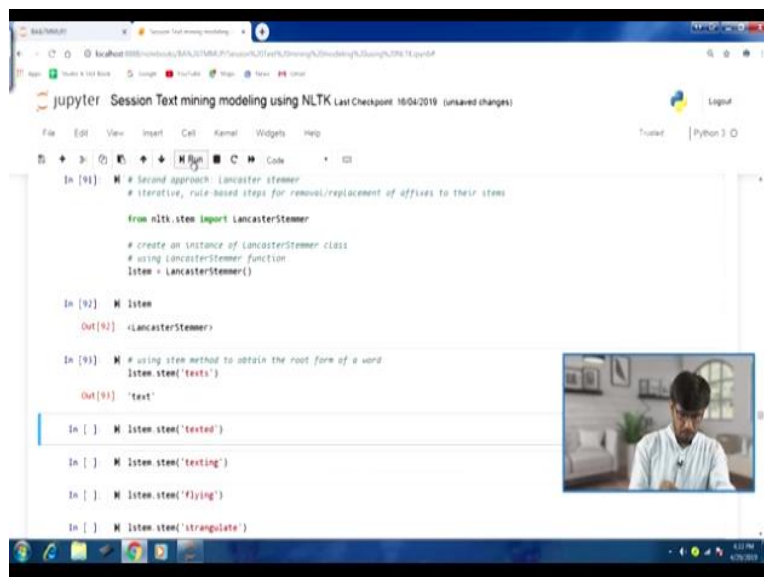
In [92]: lstem
Out[92]: <LancasterStemmer>

In [ ]: # using stem method to obtain the root form of a word
lstem.stem('texts')

In [ ]: lstem.stem('texted')
In [ ]: lstem.stem('texting')
In [ ]: lstem.stem('flying')
In [ ]: lstem.stem('strangulate')
```

So, let me run this and you can have a look at the class Lancaster stemmer this is an instance of that. Now will be again calling stem method of this particular class and will be passing on the same example same words we will be using here again and we will try to find out the root here.

(Refer Slide Time: 24:32)



```
# Second approach: Lancaster stemmer
# iterative, rule-based steps for removal/replacement of affixes to their stems

from nltk.stem import LancasterStemmer

# create an instance of LancasterStemmer class
# using LancasterStemmer function
lstem = LancasterStemmer()

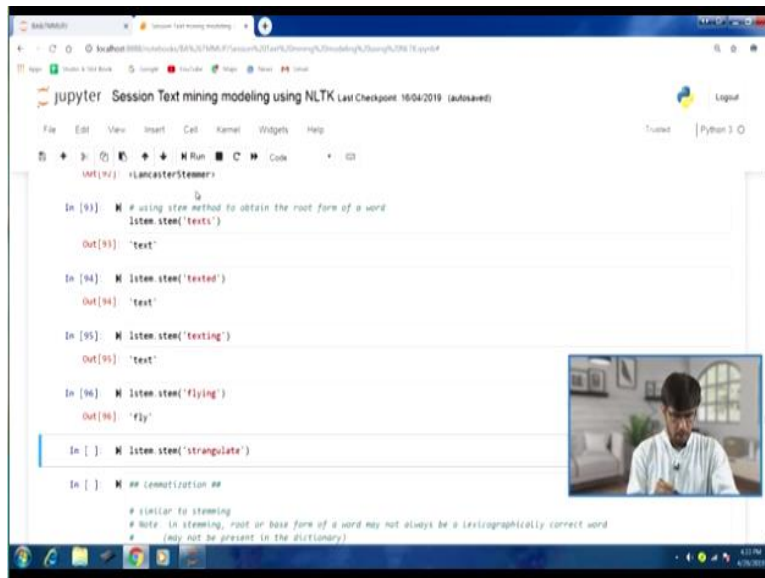
In [92]: lstem
Out[92]: <LancasterStemmer>

In [93]: # using stem method to obtain the root form of a word
lstem.stem('texts')
Out[93]: 'text'

In [ ]: lstem.stem('texted')
In [ ]: lstem.stem('texting')
In [ ]: lstem.stem('flying')
In [ ]: lstem.stem('strangulate')
```

So, you can see same output but texted also same output for texting also same output like you know Porter algorithm, how we run it for flying you can see we got fly here in the porters we got FLI.

(Refer Slide Time: 24:40)



```
from nltk.stem import LancasterStemmer

In [93]: # using stem method to obtain the root form of a word
         stemmer = LancasterStemmer()
         stemmer.stem('text')

Out[93]: 'text'

In [94]: # stemmer.stem('texted')
         stemmer.stem('texted')

Out[94]: 'text'

In [95]: # stemmer.stem('texting')
         stemmer.stem('texting')

Out[95]: 'text'

In [96]: # stemmer.stem('flying')
         stemmer.stem('flying')

Out[96]: 'fly'

In [97]: # stemmer.stem('strangulate')
         stemmer.stem('strangulate')

Out[97]: 'strangulate'

In [ ]: # Lemmatization #

# similar to stemming
# Note: in stemming, root or base form of a word may not always be a lexicographically correct word
# (may not be present in the dictionary)
# however, in lemmatization, root or base form of a word will always be present in the dictionary.

# root form or lemma is formed by removing the affixes from the word if and only if the lemma is present in the dictionary.

from nltk.stem import WordNetLemmatizer

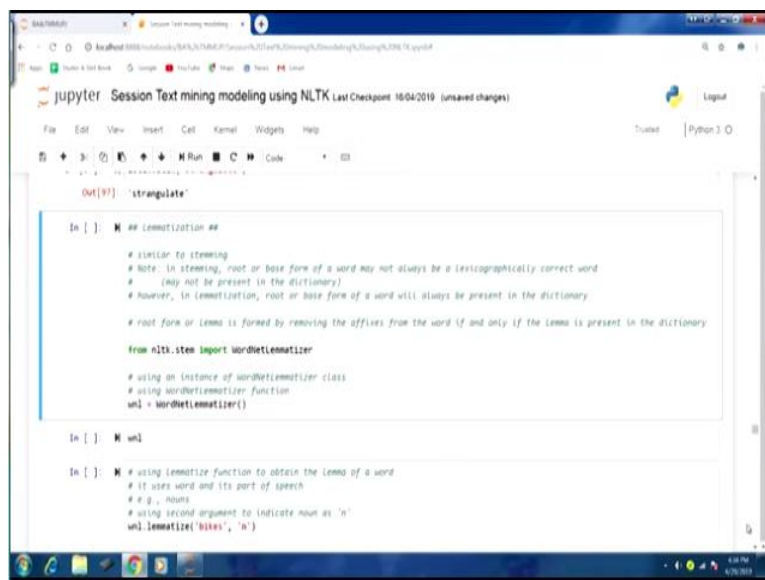
# using an instance of WordNetLemmatizer class
# using WordNetLemmatizer function
wnl = WordNetLemmatizer()

In [ ]: # wnl

In [ ]: # using lemmatize function to obtain the lemma of a word
# it uses word and its part of speech
# e.g., nouns
# using second argument to indicate noun as 'n'
wnl.lemmatize('bikes', 'n')
```

So, you can see different you know stemmer different stemming algorithm they might give you different output depending on what kind of you know steps they follow.

(Refer Slide Time: 25:00)



```
from nltk.stem import WordNetLemmatizer

In [97]: # wnl

In [98]: # using lemmatize function to obtain the lemma of a word
         wnl.lemmatize('strangulate')

Out[98]: 'strangulate'

In [ ]: # Lemmatization #

# similar to stemming
# Note: in stemming, root or base form of a word may not always be a lexicographically correct word
# (may not be present in the dictionary)
# however, in lemmatization, root or base form of a word will always be present in the dictionary.

# root form or lemma is formed by removing the affixes from the word if and only if the lemma is present in the dictionary.

from nltk.stem import WordNetLemmatizer

# using an instance of WordNetLemmatizer class
# using WordNetLemmatizer function
wnl = WordNetLemmatizer()

In [ ]: # wnl

In [ ]: # using lemmatize function to obtain the lemma of a word
# it uses word and its part of speech
# e.g., nouns
# using second argument to indicate noun as 'n'
wnl.lemmatize('bikes', 'n')
```

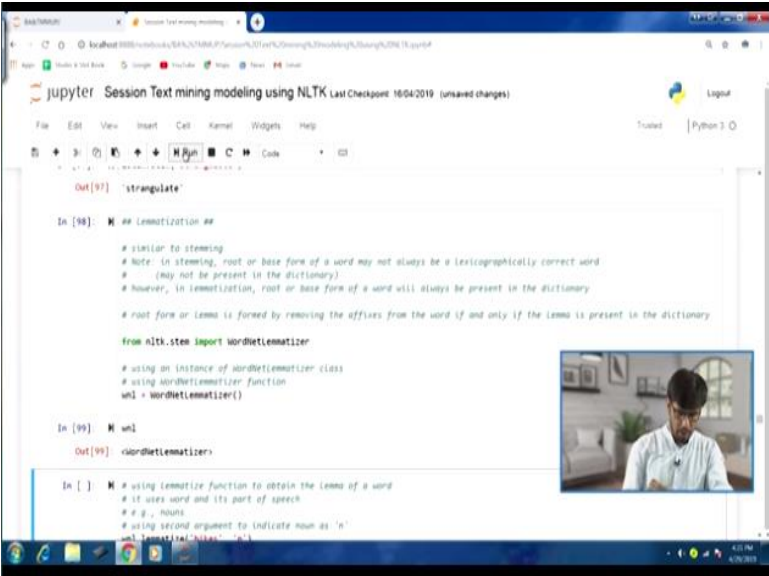
If I run this strangulate on this, so for this I think it is the again they there also difference in terms of Porter algorithm, we got a strangle, here we got a strangulated itself. So, you can see these 2 different algorithms they are working slightly differently when it comes to you know stemming. Now let us move to you know another aspect that is lambda lemmatization, so lambda lemmatization is quite similar to is to stemming.

And in stemming typically root or base form of a word you know that we typically obtain that

may not always be like lexicographically correct word. So, we are just trying to arrive at a stem or root form of a word which might not be present in the dictionary in another words. But when we perform lemmatization it is the root or base form of word you know that will also be present in dictionary that is taken, so slight difference is there.

So, root form of or lemma is formed by removing the affixes from the word if and only if lemma is present in the dictionary. So, there in stemming the stem might not be present in the dictionary however in lemmatization the lemma that we obtain the root form of the word that we obtain it has to be present in the dictionary. So, for this will be using again and NLTKdot stem module and word net lemmatizer, this is the class that we will be using, so we will define an instance of this class.

(Refer Slide Time: 26:33)



```
Out[97]: 'strangulate'

In [98]: ## Lemmatization ##

# similar to stemming
# Note: in stemming, root or base form of a word may not always be a lexicographically correct word
#       (may not be present in the dictionary)
# however, in lemmatization, root or base form of a word will always be present in the dictionary

# root form or lemma is formed by removing the affixes from the word if and only if the lemma is present in the dictionary

from nltk.stem import WordNetLemmatizer

# using an instance of WordNetLemmatizer class
# using wordnetlemmatizer function
wn_l = WordNetLemmatizer()

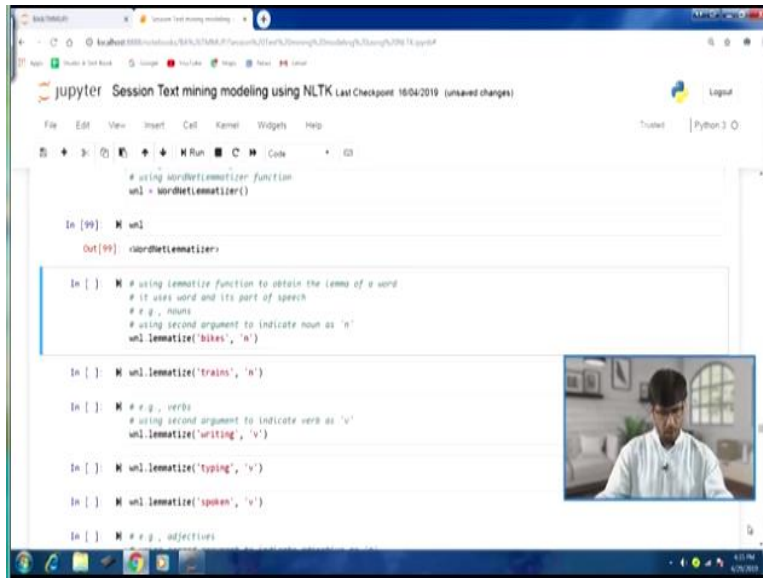
In [99]: wn_l

Out[99]: WordNetLemmatizer()

In [ ]: ## using lemmatize function to obtain the lemma of a word
# it uses word and its part of speech
# e.g., nouns
# using second argument to indicate noun as 'n'
wn_l.lemmatize('strangulate', 'n')
```

So, let me run this and will have WNL, this is an instance of word net lemmatizer class as you can see.

(Refer Slide Time: 26:40)



```
# using wordnetlemmatizer function
wnl = WordNetLemmatizer()

In [99]: wnl
Out[99]: WordNetLemmatizer()

In [ ]: # using lemmatize function to obtain the lemma of a word
# it uses word and its part of speech
# e.g., nouns
# using second argument to indicate noun as 'n'
wnl.lemmatize('bikes', 'n')

In [ ]: wnl.lemmatize('trains', 'n')

In [ ]: # e.g., verbs
# using second argument to indicate verb as 'v'
wnl.lemmatize('writing', 'v')

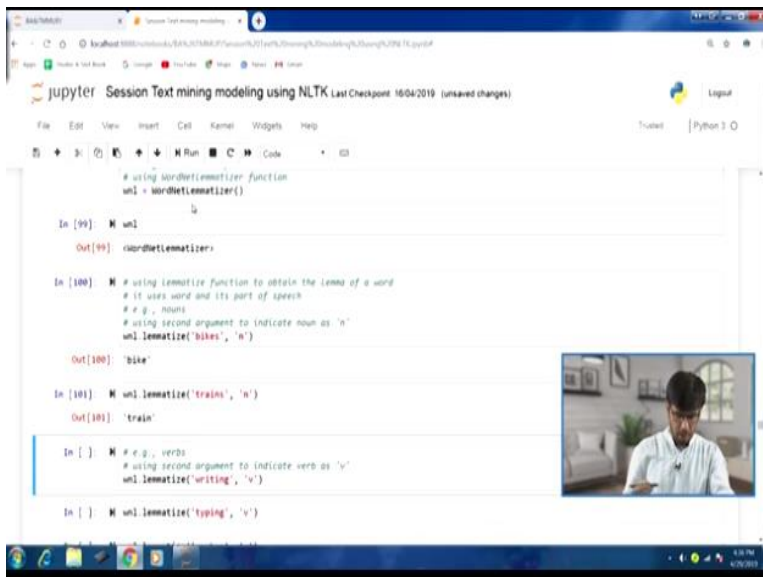
In [ ]: wnl.lemmatize('typing', 'v')

In [ ]: wnl.lemmatize('spoken', 'v')

In [ ]: # e.g., adjectives
```

Now what we will do we will take few examples, so we will use the limit lemmatize function to obtain lemma of a word here and you know it uses word and it is part of his speech. So, for example if we you know let us take the example of you know some nouns. So, second argument is typically use to indicate noun, so that is as n in single codes you can see WNL dot lemmatize within parentheses first argument is the word that we would like you lemmatize here bikes and comma and the part of a speech it is a noun.

(Refer Slide Time: 27:18)



```
# using wordnetlemmatizer function
wnl = WordNetLemmatizer()

In [99]: wnl
Out[99]: WordNetLemmatizer()

In [100]: # using lemmatize function to obtain the lemma of a word
# it uses word and its part of speech
# e.g., nouns
# using second argument to indicate noun as 'n'
wnl.lemmatize('bikes', 'n')

Out[100]: 'bike'

In [101]: wnl.lemmatize('trains', 'n')

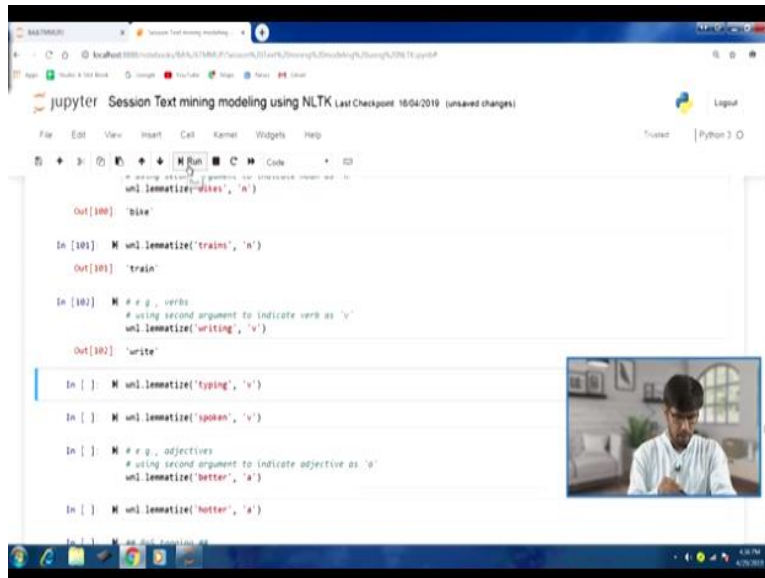
Out[101]: 'train'

In [ ]: # e.g., verbs
# using second argument to indicate verb as 'v'
wnl.lemmatize('writing', 'v')

In [ ]: wnl.lemmatize('typing', 'v')
```

So, let me run this you can see bikes has become pike similarly trains it has become train.

(Refer Slide Time: 27:27)



The screenshot shows a Jupyter Notebook titled "Session Text mining modeling using NLTK" with a last checkpoint of 16/04/2019. The notebook contains several code cells demonstrating the use of the `wnl.lemmatize()` function. The first cell shows `wnl.lemmatize('bikes', 'n')` returning `'bike'`. The second cell shows `wnl.lemmatize('trains', 'n')` returning `'train'`. The third cell shows `wnl.lemmatize('writing', 'v')` returning `'write'`. The fourth cell shows `wnl.lemmatize('typing', 'v')` returning `'type'`. The fifth cell shows `wnl.lemmatize('spoken', 'v')` returning `'speak'`. The sixth cell shows `wnl.lemmatize('better', 'a')` returning `'better'`. The seventh cell shows `wnl.lemmatize('hotter', 'a')` returning `'hotter'`. A small video inset in the bottom right corner shows a man speaking.

```
wnl.lemmatize('bikes', 'n')
Out[100]: 'bike'

In [101]: wnl.lemmatize('trains', 'n')
Out[101]: 'train'

In [102]: # e.g., verbs
# using second argument to indicate verb as 'v'
wnl.lemmatize('writing', 'v')
Out[102]: 'write'

In [ ]: wnl.lemmatize('typing', 'v')

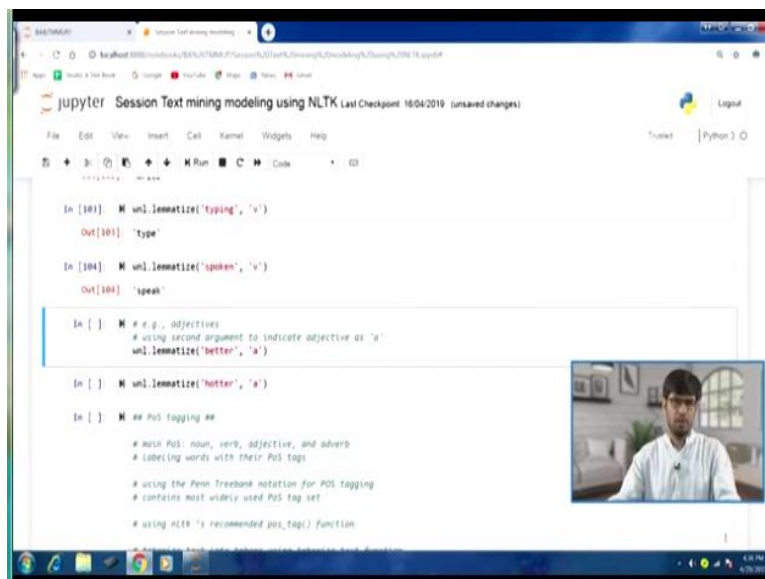
In [ ]: wnl.lemmatize('spoken', 'v')

In [ ]: # e.g., adjectives
# using second argument to indicate adjective as 'a'
wnl.lemmatize('better', 'a')

In [ ]: wnl.lemmatize('hotter', 'a')
```

So, you can see part of a speech is also is required to be passed on here because that will help in terms of identifying the you know root form of the word in the dictionary. Now let us take example of a few word, so again in the second argument instead of you know now will indicate v-twin you know to denote word here. So, writing will become write, typing will become type and a spoken will become a speak as you can see here.

(Refer Slide Time: 27:56)



The screenshot shows a Jupyter Notebook titled "Session Text mining modeling using NLTK" with a last checkpoint of 16/04/2019. The notebook contains several code cells demonstrating the use of the `wnl.lemmatize()` function. The first cell shows `wnl.lemmatize('typing', 'v')` returning `'type'`. The second cell shows `wnl.lemmatize('spoken', 'v')` returning `'speak'`. The third cell shows `wnl.lemmatize('better', 'a')` returning `'better'`. The fourth cell shows `wnl.lemmatize('hotter', 'a')` returning `'hotter'`. The fifth cell shows `wnl.lemmatize('better', 'a')` returning `'better'`. The sixth cell shows `wnl.lemmatize('hotter', 'a')` returning `'hotter'`. A small video inset in the bottom right corner shows a man speaking.

```
In [103]: wnl.lemmatize('typing', 'v')
Out[103]: 'type'

In [104]: wnl.lemmatize('spoken', 'v')
Out[104]: 'speak'

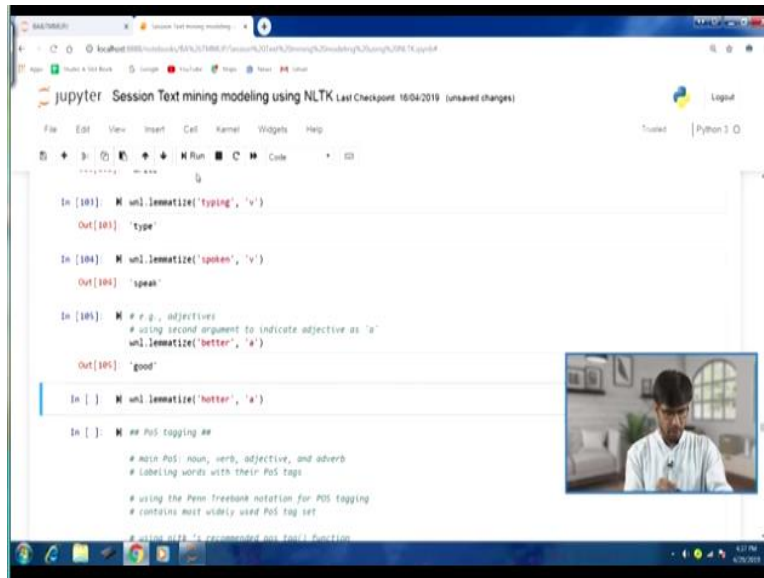
In [ ]: # e.g., adjectives
# using second argument to indicate adjective as 'a'
wnl.lemmatize('better', 'a')

In [ ]: wnl.lemmatize('hotter', 'a')

In [ ]: # POS tagging #
# main POS: noun, verb, adjective, and adverb
# labeling words with their POS tags
# using the Penn Treebank notation for POS tagging
# contains most widely used POS tag set
# using nltk's recommended pos_tag() function
```

So, you can see how lemmatization is slightly different from you know stem, now let us take another example this term adjective so better and A to indicate that this is an adjective.

(Refer Slide Time: 28:13)

A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (unsaved changes)". The notebook is running on a local host. The code in the cells shows the use of the `nltk.stem.WordNetLemmatizer` class. The first cell shows the lemmatization of the word "typing", which returns "type". The second cell shows the lemmatization of the word "spoken", which returns "speak". The third cell shows the lemmatization of the word "good", which returns "good". The fourth cell shows the lemmatization of the word "hotter", which returns "hot". The fifth cell shows the lemmatization of the word "hot", which returns "hot". The notebook interface includes a menu bar with options like File, Edit, View, Insert, Cell, Kernel, Widgets, and Help. There is also a toolbar with icons for running, saving, and other actions. A small video window in the bottom right corner shows a person speaking.

```
In [181]: M nltk.Lemmatizer('typing', 'v')
Out[181]: 'type'

In [184]: M nltk.Lemmatizer('spoken', 'v')
Out[184]: 'speak'

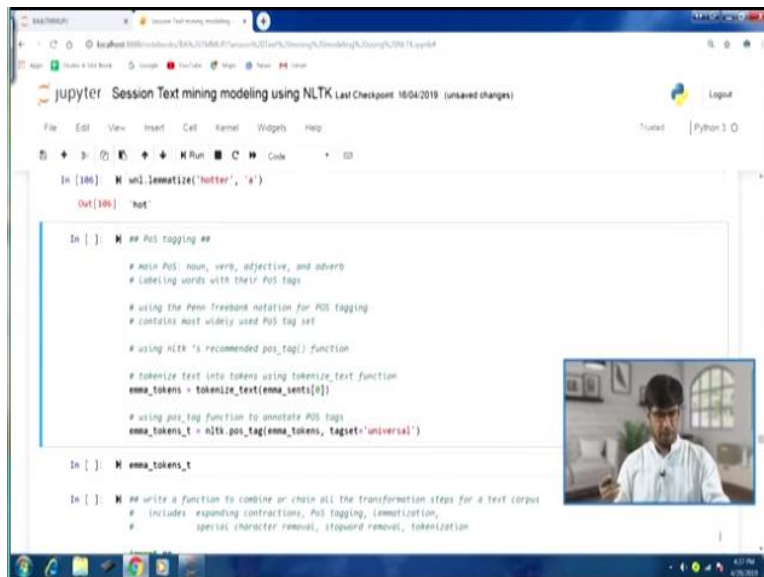
In [185]: M # e.g., adjectives
# using second argument to indicate adjective as 'a'
nltk.Lemmatizer('better', 'a')
Out[185]: 'good'

In [ ]: M nltk.Lemmatizer('hotter', 'a')

In [ ]: M # POS tagging ##
# main pos: noun, verb, adjective, and adverb
# labeling words with their pos tags
# using the Penn Treebank notation for POS tagging
# contains most widely used pos tag set
# using nltk's recommended pos_tag() function
```

So, if I run this you can see good and hotter we got hot.

(Refer Slide Time: 28:16)

A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK Last Checkpoint: 16/04/2019 (unsaved changes)". The notebook is running on a local host. The code in the cells shows the use of the `nltk.pos\_tag` function. The first cell shows the lemmatization of the word "hotter", which returns "hot". The second cell shows the lemmatization of the word "hot", which returns "hot". The third cell shows the lemmatization of the word "hot", which returns "hot". The fourth cell shows the lemmatization of the word "hot", which returns "hot". The fifth cell shows the lemmatization of the word "hot", which returns "hot". The notebook interface includes a menu bar with options like File, Edit, View, Insert, Cell, Kernel, Widgets, and Help. There is also a toolbar with icons for running, saving, and other actions. A small video window in the bottom right corner shows a person speaking.

```
In [186]: M nltk.Lemmatizer('hotter', 'a')
Out[186]: 'hot'

In [ ]: M # POS tagging ##
# main pos: noun, verb, adjective, and adverb
# labeling words with their pos tags
# using the Penn Treebank notation for POS tagging
# contains most widely used pos tag set
# using nltk's recommended pos_tag() function
# tokenize text into tokens using tokenize_text function
emma_tokens = tokenize_text(emma_sents[0])
# using pos_tag function to annotate POS tags
emma_tokens_t = nltk.pos_tag(emma_tokens, tagset='universal')

In [ ]: M emma_tokens_t

In [ ]: M # write a function to combine or chain all the transformation steps for a text corpus
# includes expanding contractions, pos tagging, lemmatization,
# special character removal, stopword removal, tokenization
```

So, in this fashion you can see that how lemmatizer is working here, now let us move to the another aspect here that is you know POS tagging. So, we just learned about you know that in lambda lemmatization that whenever we are lemmatizing a word or text or a particular token, we also indicate part of speech. So, how do we tag you know all the tokens with their corresponding part of a speech so for that we have will talk about this POS tagging aspect here.

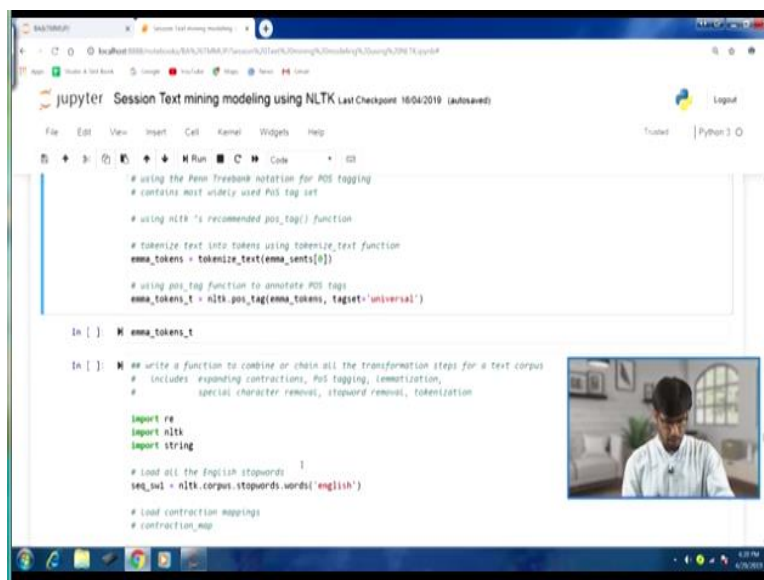
So, main part of a speech that are typically used in writing and here also noun, verb, adjective and adverb. So, it is important for us to label words with their POS tags and typically pantry

bank notation you know is used for you know he POS tagging and most it contains the most widely used pure stack set. So, will be using this analogy NLTKs POS underscore tag function to perform part of a speech tagging.

So, let us take an example, so we will use tokenize underscore text function here and we will pass on Emma underscore sense, it will you know perform tokenization here and then will call this NLTK dot. So, NLTK dot POS underscore tag, so first argument we need to pass on tokens there and tag set also something you know different tag sets are there or you can just not specify this, so universal is one example that can be used.

So, NLTK so in that code you can see that we first organize token, so we need to talk tokenize, for this I have written a function tokenize underscore text in the next section.

(Refer Slide Time: 30:00)

A screenshot of a Jupyter Notebook interface. The title bar says "jupyter Session Text mining modeling using NLTK Last Checkpoint 18/04/2019 (autosaved)". The code area contains several lines of Python code with comments. The first block of code uses Penn Treebank notation for POS tagging, mentions the most widely used POS tag set, and uses NLTK's recommended pos_tag() function. It then tokenizes text into tokens using tokenize_text function, creating a variable emma_tokens. The second block uses pos_tag function to annotate POS tags, creating a variable emma_tokens_t. The third block shows the import of re, nltk, and string modules, followed by loading all English stopwords and contraction mappings. A small video inset in the bottom right corner shows a person speaking.

```
# using the Penn Treebank notation for POS tagging
# contains most widely used POS tag set
# using nltk's recommended pos_tag() function
# tokenize text into tokens using tokenize_text function
emma_tokens = tokenize_text(emma_sents[0])
# using pos_tag function to annotate POS tags
emma_tokens_t = nltk.pos_tag(emma_tokens, tagset='universal')

In [ ]: emma_tokens_t

In [ ]: # we write a function to combine or chain all the transformation steps for a text corpus
# includes expanding contractions, pos tagging, lemmatization,
# special character removal, stopword removal, tokenization

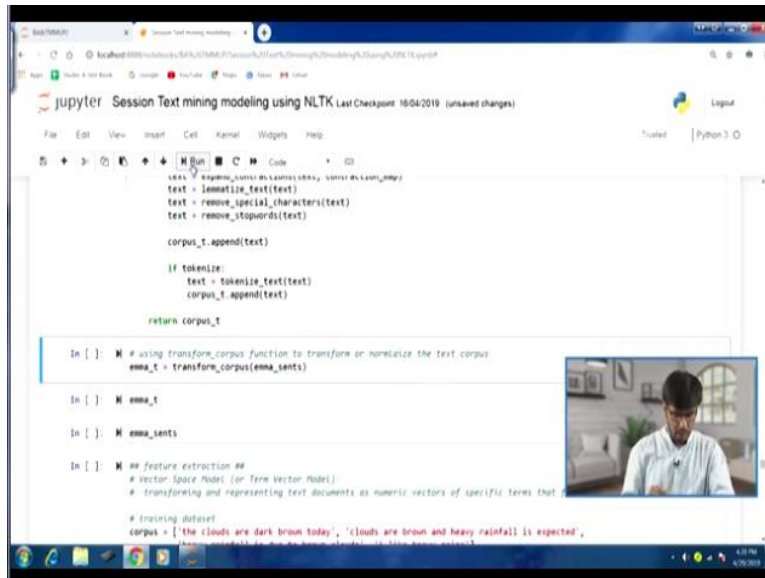
import re
import nltk
import string

# load all the English stopwords
stop_words = nltk.corpus.stopwords.words('english')

# load contraction mappings
# contraction_map
```

So, here I will just also run another part of the code right now, so that we are able to access that this particular function here.

(Refer Slide Time: 30:14)

A screenshot of a Jupyter Notebook titled "Session Text mining modeling using NLTK". The notebook shows a function definition for `transform\_corpus` which takes a list of sentences and a boolean `tokenize` parameter. The function processes each sentence by lowercasing, lemmatizing, removing special characters and stopwords, and tokenizing if requested. Below the function, there are several input-output pairs: `emma\_t = transform\_corpus(emma\_sents)`, printing `emma\_t`, `emma\_sents`, and a feature extraction step using a `VecSpaceModel` to transform the corpus into a matrix of term vectors. A training dataset is also defined as a list of sentences.

```
def transform_corpus(emma_sents, tokenize=False):
    corpus_t = []
    for text in emma_sents:
        text = text.lower()
        text = lemmatize_text(text)
        text = remove_special_characters(text)
        text = remove_stopwords(text)

        corpus_t.append(text)

    if tokenize:
        text = tokenize_text(text)
        corpus_t.append(text)

    return corpus_t

In [ ]: # using transform_corpus function to transform or normalize the text corpus
emma_t = transform_corpus(emma_sents)

In [ ]: # emma_t

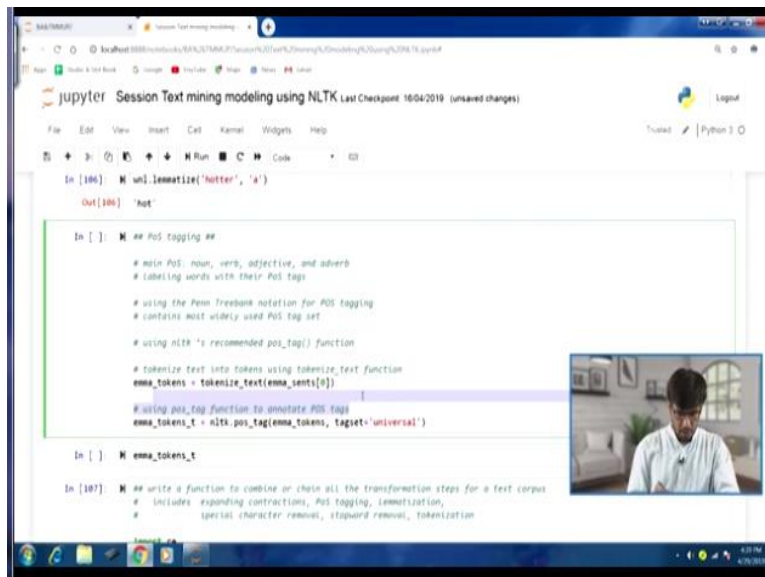
In [ ]: # emma_sents

In [ ]: ## feature extraction ##
# Vector Space Model (or Term Vector Model)
# transforming and representing text documents as numeric vectors of specific terms that

# training dataset
corpus = ['the clouds are dark brown today', 'clouds are brown and heavy rainfall is expected',
          'brown clouds are dark brown today', 'dark brown clouds are brown']
```

So, I will run this part as well and I will go back we will come back to this part and discuss in detail but because since I am using this function, so let me go back and use it here for POS tagging.

(Refer Slide Time: 30:28)

A screenshot of the same Jupyter Notebook, showing the implementation of POS tagging. It starts with a lemmatization example: `lemmatize('hotter', 'a')` resulting in `hot`. Then, it defines a `pos\_tag` function that uses NLTK's `pos\_tag` function to tag tokens. The function takes a list of tokens and returns a list of tuples where each tuple contains a token and its corresponding POS tag. Finally, it applies this function to the `emma\_tokens` list, creating `emma\_tokens\_t`.

```
In [106]: # unlemmatize('hotter', 'a')
Out[106]: 'hot'

In [ ]: ## pos tagging ##

# main pos: noun, verb, adjective, and adverb
# labeling words with their pos tags

# using the Penn Treebank notation for POS tagging
# contains most widely used POS tag set

# using nltk's recommended pos_tag() function

# tokenize text into tokens using tokenize_text function
emma_tokens = tokenize_text(emma_sents[0])

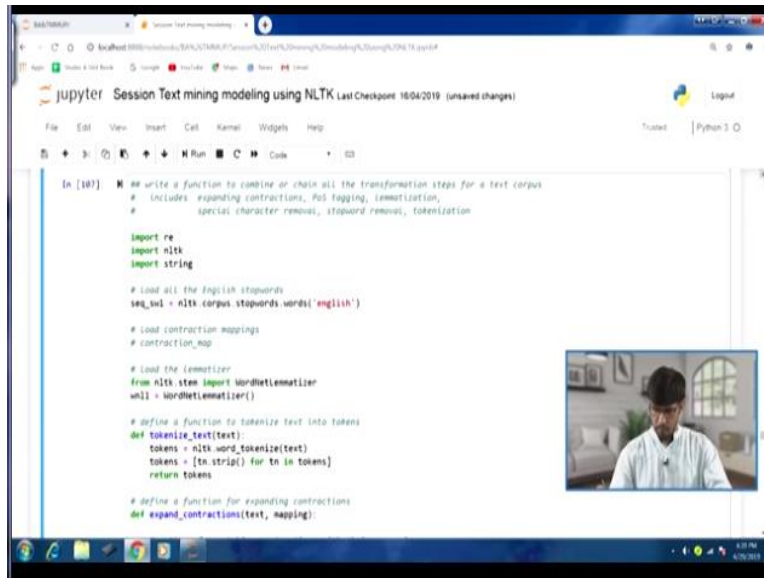
# using pos_tag function to annotate POS tags
emma_tokens_t = nltk.pos_tag(emma_tokens, tagset='universal')

In [ ]: # emma_tokens_t

In [107]: ## write a function to combine or chain all the transformation steps for a text corpus
# includes expanding contractions, pos tagging, lemmatization,
# special character removal, stopword removal, tokenization
```

So, if I run this and let us have a look at the tokens here.

(Refer Slide Time: 30:34)



```
In [187]: # We write a function to combine or chain all the transformation steps for a text corpus
# Includes: expanding contractions, POS tagging, lemmatization,
# special character removal, stopword removal, tokenization

import re
import nltk
import string

# Load all the English stopwords
stop_words = nltk.corpus.stopwords.words('english')

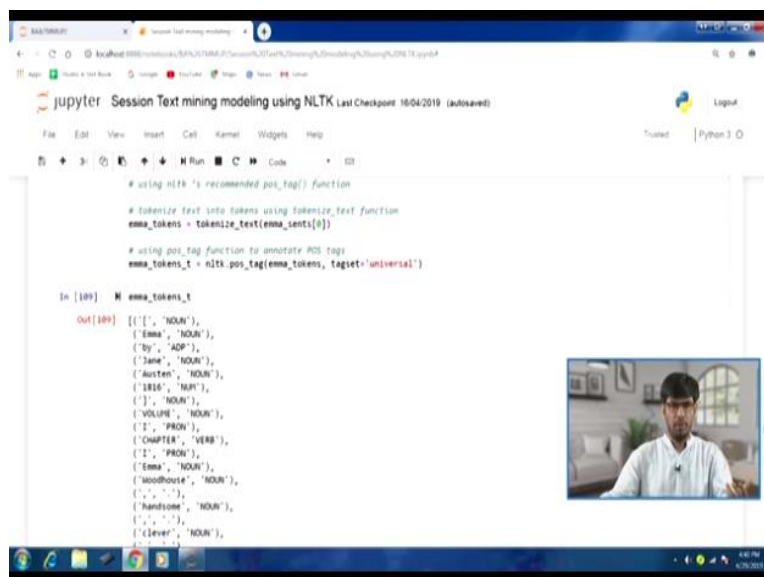
# Load contraction mappings
# contraction_map

# Load the lemmatizer
from nltk.stem import WordNetLemmatizer
wnl = WordNetLemmatizer()

# Define a function to tokenize text into tokens
def tokenize_text(text):
    tokens = nltk.word_tokenize(text)
    tokens = [tn.strip() for tn in tokens]
    return tokens

# Define a function for expanding contractions
def expand_contractions(text, mapping):
```

(Refer Slide Time: 30:39)



```
# using nltk's recommended pos_tag() function
# tokenize text into tokens using tokenize_text function
emma_tokens = tokenize_text(emma_sentences[0])

# using pos_tag function to annotate POS tags
emma_tokens_t = nltk.pos_tag(emma_tokens, tagset='universal')

In [189]: # emma_tokens_t
Out[189]: [(('I', 'NOUN'),
('Emma', 'NOUN'),
('by', 'ADP'),
('Jane', 'NOUN'),
('Austen', 'NOUN'),
('1816', 'NOUN'),
(';', 'NOUN'),
('VOLUME', 'NOUN'),
('I', 'PRON'),
('CHAPTER', 'VERB'),
('I', 'PRON'),
('Emma', 'NOUN'),
('Woodhouse', 'NOUN'),
(';', 'NOUN'),
('handsome', 'NOUN'),
(';', 'NOUN'),
('clever', 'NOUN'))]
```

So, you can see in the output here in the output 109 you can see Emma noun by ADP you know all those tokens they have been tagged with their corresponding part of a speech. So, you can clearly see here, so you can see how first we tokenized the text and those tokens were passed in the POS underscore tag function to perform the POS tagging and now you can see the output. So, we would like to stop here and in the next lecture we will move on to the you know some of the next aspect of you know transforming unstructured text, thank you.

Keywords: contraction, structured and unstructured text, tokens, reduction, inflation etc.