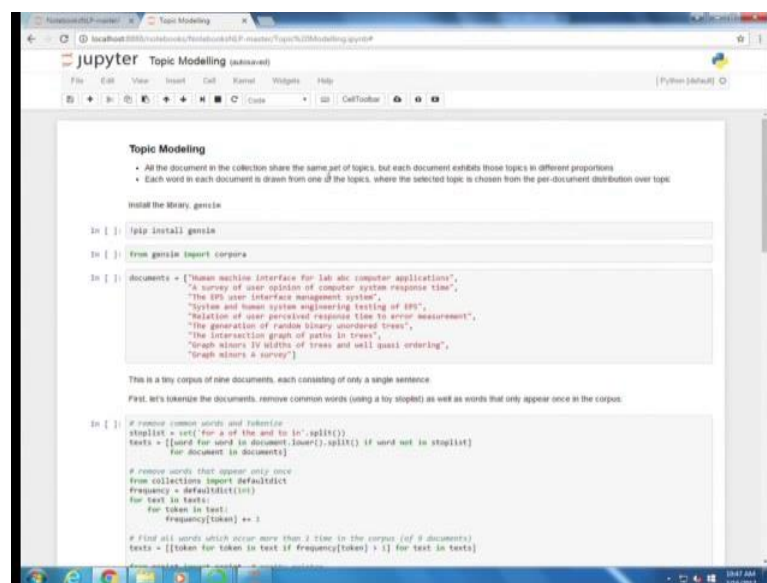


Natural Language Processing
Prof. Amrith Krishna
Department of Computer Science and Engineering
Indian Institute of Technology, Kharagpur

Lecture - 58
Tutorial

Hello everyone. So, welcome to today's tutorial session, today we will be discussing about topic modeling. So, for this tutorial, we assume that we already have Jupyter notebook installed and we will be starting with the topic modeling session.

(Refer Slide Time: 00:41)



```
Topic Modeling
• All the documents in the collection share the same set of topics, but each document exhibits those topics in different proportions
• Each word in each document is drawn from one of the topics, where the selected topic is chosen from the per-document distribution over topics

Install the library, gensim

In [ ]: !pip install gensim

In [ ]: from gensim import corpora

In [ ]: documents = ["Human machine interface for lab abc computer applications",
                    "A survey of user opinion of computer system response time",
                    "The EPS user interface management system",
                    "System and human system engineering testing of EPS",
                    "Relation of user perceived response time to error measurement",
                    "The generation of random binary ordered trees",
                    "The interaction graph of paths in trees",
                    "Graph minors IV Widths of trees and well quasi ordering",
                    "Graph minors V survey"]

This is a tiny corpus of nine documents, each consisting of only a single sentence.
First, let's tokenize the documents, remove common words (using a toy stoplist) as well as words that only appear once in the corpus.

In [ ]: # remove common words and tokenize
stoplist = set('for a of the and to in'.split())
texts = [[word for word in document.lower().split() if word not in stoplist]
          for document in documents]

# remove words that appear only once
from collections import defaultdict
frequency = defaultdict(int)
for text in texts:
    for token in text:
        frequency[token] += 1

# find all words which occur more than 1 time in the corpus (if 4 documents)
texts = [[token for token in text if frequency[token] > 1] for text in texts]
```

We have started Jupyter notebook and here we will be discussing about topic modeling. So, I assume you remember, what is topic modeling? So, it essentially is a way to model a corpus or a set of documents where all the documents in the collection share the same set of topics, but each document exhibits those topics in different proportions.

We essentially assume that we have a corpus and in that corpus that contains multiple documents, every document share a common set of topics. Now given those topics, they can be varying in different proportions in each of these documents. Now each word in the

Now, to infer all this, all we have is basically the observations or bay or the corpus itself. So, to run this codes that we are discussing today you might be requiring to install a library called Gensim. So, I assume that everyone is familiar with anaconda navigator.

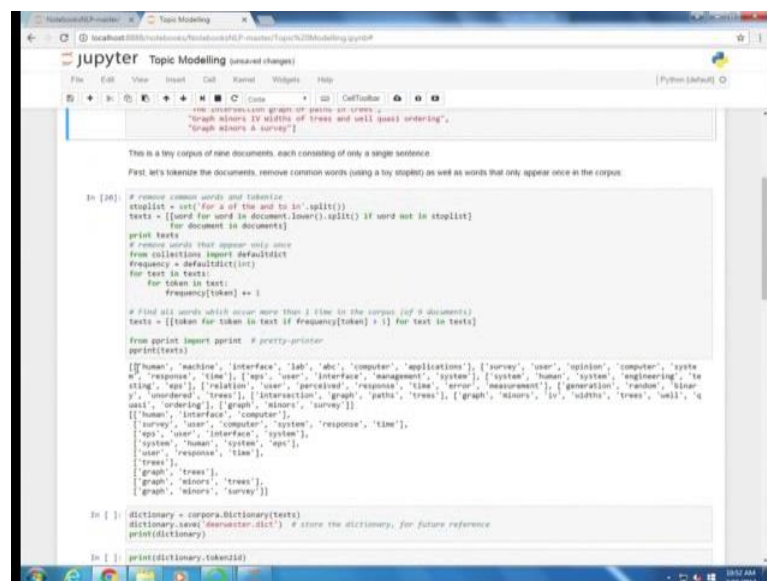
[illegible]

Once the library is installed, we will import the library and for this tutorial, we will be assuming a very small corpus which is basically a collection of 9 different sentences. So, here is a tiny corpus of 9 documents and each document essentially contains only one sentence. So, in order to view this corpus to the model we have to do some essential pre processing. So, we will be showing what all pre processing needs to be done. So, we

need to convert this to a suitable bag of words format and we will be showing how to do this.

Here are some pre processing steps that we are going to do. So, for this simplistic corpus we assume the only set of stop words are these few words, but in practicality you might have to consider other stop words as well, but there are pre defined sets that you can obtain we will show how to obtain those lists. So, we make a list of all those stop words we also make sure that all the words are converted to lowercase. So, here what we do is that we take each sentence in the documents and then we split it based on the space. So, we do tokenization again based on a very simple notion that if there is a space between 2 characters we assume that they belong to 2 different words.

(Refer Slide Time: 04:37)



```
from collections import defaultdict

# remove common words and tokenize
stoplist = set('for a of the and to in'.split())
texts = [[word for word in document.lower().split() if word not in stoplist]
          for document in documents]

# remove words that appear only once
from collections import defaultdict
frequency = defaultdict(int)
for text in texts:
    for token in text:
        frequency[token] += 1

# Find all words which occur more than 1 time in the corpus (of 8 documents)
texts = [[token for token in text if frequency[token] > 1] for text in texts]

from pprint import pprint # pretty-printer
pprint(texts)

[['human', 'machine', 'interface', 'lab', 'abc', 'computer', 'applications'], ['survey', 'user', 'opinion', 'computer', 'system', 'response', 'time'], ['api', 'user', 'interface', 'management', 'system'], ['system', 'human', 'system', 'engineering', 'testing', 'api'], ['relation', 'user', 'perceived', 'response', 'time', 'error', 'measurement'], ['generation', 'random', 'linear', 'unordered', 'trees'], ['interaction', 'graph', 'paths', 'trees'], ['graph', 'minors', 'ir', 'widths', 'trees', 'will', 'q', 'wall', 'ordering'], ['graph', 'minors', 'survey']]

[['human', 'interface', 'computer'],
 ['survey', 'user', 'computer', 'system', 'response', 'time'],
 ['api', 'user', 'interface', 'system'],
 ['system', 'human', 'system', 'api'],
 ['user', 'response', 'time'],
 ['trees'],
 ['graph', 'trees'],
 ['graph', 'minors', 'trees'],
 ['graph', 'minors', 'survey']]

In [ ]: dictionary = corpora.Dictionary(texts)
dictionary.save('news-corpus.dict') # store the dictionary, for future reference
print(dictionary)

In [ ]: print(dictionary.token2id)
```

Since our corpus is something that we are doing for representation purposes these simplistic method will help you in doing this, but again you can rely on any standard Tokenizers for this task. So, text will essentially contain all those individual words see if you see here the word the sentence human machine interface for lab, A B C computer applications here that document is now represented as the list of words where you can find the H in human being has been turned in this small h or the lowercase h, the stop word for is removed and we essentially got a list of list where each list shows a document

with it is important words.

After doing this what we find is that we find the count or the frequency of a particular word occurring in the entire corpus for example, the word human is appearing once in document one and it also appears one time in document 4. So, the total occurrence is 2 for the word human. So, this particular chunk of code basically tells you how to find the frequency of the individual words in your corpus once that is done, we filter those words which appear more than once. So, we keep those words only those words which appear more than once in our corpus and we remove them.

What you find here in the second printed statement are those words which are filtered after this particular criteria is applied. So, we can see that machine which was appearing only once in the entire corpus is removed and similarly other words like interface etcetera. So, once we have this, most often, what happens is that when we have large corp error, we may not be doing it every time like we may not be doing the pre processing every time because it itself might be time consuming.

(Refer Slide Time: 06:42)

```
File Edit View Insert Cell Kernel Widgets Help
Jupyter Topic Modeling (saved changes)
Python 3.6.0 [default]

In [21]: dictionary = corpora.Dictionary(texts)
         dictionary.save('dictionary.dict') # store the dictionary, for future reference
         print(dictionary)
         Dictionary12 unique tokens: ['u'minors', 'u'graph', 'u'system', 'u'trains', 'u'apps'],...

In [22]: print(dictionary.token2id)
         {'u'minors': 11, 'u'graph': 10, 'u'system': 0, 'u'trains': 0, 'u'apps': 0, 'u'computer': 1, 'u'survey': 0, 'u'user': 7, 'u'human': 2, 'u'i
         nt': 4, 'u'interface': 0, 'u'response': 0}

In [24]: new_doc = "human computer interaction computer"
         new_doc = dictionary.doc2id(new_doc.lower().split())
         print(new_doc) # the word "interaction" does not appear in the dictionary and is ignored
         [(1, 3), (3, 1)]

The function doc2id() simply counts the number of occurrences of each distinct word. Converts the word to the integer word id and returns the result as a
sparse vector (word, id). (word, id). Therefore results in the document "human computer interaction", the words "computer"
and "human", identified by an integer id given by the token dictionary, appear once; the other two documentary words appear (respectively) zero times. Check this at
the dictionary employed in the previous cell and see that they match.

In [25]: corpus = [dictionary.doc2id(text) for text in texts]
         corpora.Dictionary.save('dictionary.dict', corpus) # store to disk, for later use
         print(x)
         [(0, 1), (2, 1), (3, 1), (4, 1), (5, 1), (6, 1), (7, 1), (8, 1), (9, 1), (10, 1), (11, 1), (12, 1), (13, 1), (14, 1), (15, 1), (16, 1), (17, 1), (18, 1), (19, 1), (20, 1), (21, 1), (22, 1), (23, 1), (24, 1), (25, 1), (26, 1), (27, 1), (28, 1), (29, 1), (30, 1), (31, 1), (32, 1), (33, 1), (34, 1), (35, 1), (36, 1), (37, 1), (38, 1), (39, 1), (40, 1), (41, 1), (42, 1), (43, 1), (44, 1), (45, 1), (46, 1), (47, 1), (48, 1), (49, 1), (50, 1), (51, 1), (52, 1), (53, 1), (54, 1), (55, 1), (56, 1), (57, 1), (58, 1), (59, 1), (60, 1), (61, 1), (62, 1), (63, 1), (64, 1), (65, 1), (66, 1), (67, 1), (68, 1), (69, 1), (70, 1), (71, 1), (72, 1), (73, 1), (74, 1), (75, 1), (76, 1), (77, 1), (78, 1), (79, 1), (80, 1), (81, 1), (82, 1), (83, 1), (84, 1), (85, 1), (86, 1), (87, 1), (88, 1), (89, 1), (90, 1), (91, 1), (92, 1), (93, 1), (94, 1), (95, 1), (96, 1), (97, 1), (98, 1), (99, 1), (100, 1), (101, 1), (102, 1), (103, 1), (104, 1), (105, 1), (106, 1), (107, 1), (108, 1), (109, 1), (110, 1), (111, 1), (112, 1), (113, 1), (114, 1), (115, 1), (116, 1), (117, 1), (118, 1), (119, 1), (120, 1), (121, 1), (122, 1), (123, 1), (124, 1), (125, 1), (126, 1), (127, 1), (128, 1), (129, 1), (130, 1), (131, 1), (132, 1), (133, 1), (134, 1), (135, 1), (136, 1), (137, 1), (138, 1), (139, 1), (140, 1), (141, 1), (142, 1), (143, 1), (144, 1), (145, 1), (146, 1), (147, 1), (148, 1), (149, 1), (150, 1), (151, 1), (152, 1), (153, 1), (154, 1), (155, 1), (156, 1), (157, 1), (158, 1), (159, 1), (160, 1), (161, 1), (162, 1), (163, 1), (164, 1), (165, 1), (166, 1), (167, 1), (168, 1), (169, 1), (170, 1), (171, 1), (172, 1), (173, 1), (174, 1), (175, 1), (176, 1), (177, 1), (178, 1), (179, 1), (180, 1), (181, 1), (182, 1), (183, 1), (184, 1), (185, 1), (186, 1), (187, 1), (188, 1), (189, 1), (190, 1), (191, 1), (192, 1), (193, 1), (194, 1), (195, 1), (196, 1), (197, 1), (198, 1), (199, 1), (200, 1), (201, 1), (202, 1), (203, 1), (204, 1), (205, 1), (206, 1), (207, 1), (208, 1), (209, 1), (210, 1), (211, 1), (212, 1), (213, 1), (214, 1), (215, 1), (216, 1), (217, 1), (218, 1), (219, 1), (220, 1), (221, 1), (222, 1), (223, 1), (224, 1), (225, 1), (226, 1), (227, 1), (228, 1), (229, 1), (230, 1), (231, 1), (232, 1), (233, 1), (234, 1), (235, 1), (236, 1), (237, 1), (238, 1), (239, 1), (240, 1), (241, 1), (242, 1), (243, 1), (244, 1), (245, 1), (246, 1), (247, 1), (248, 1), (249, 1), (250, 1), (251, 1), (252, 1), (253, 1), (254, 1), (255, 1), (256, 1), (257, 1), (258, 1), (259, 1), (260, 1), (261, 1), (262, 1), (263, 1), (264, 1), (265, 1), (266, 1), (267, 1), (268, 1), (269, 1), (270, 1), (271, 1), (272, 1), (273, 1), (274, 1), (275, 1), (276, 1), (277, 1), (278, 1), (279, 1), (280, 1), (281, 1), (282, 1), (283, 1), (284, 1), (285, 1), (286, 1), (287, 1), (288, 1), (289, 1), (290, 1), (291, 1), (292, 1), (293, 1), (294, 1), (295, 1), (296, 1), (297, 1), (298, 1), (299, 1), (300, 1), (301, 1), (302, 1), (303, 1), (304, 1), (305, 1), (306, 1), (307, 1), (308, 1), (309, 1), (310, 1), (311, 1), (312, 1), (313, 1), (314, 1), (315, 1), (316, 1), (317, 1), (318, 1), (319, 1), (320, 1), (321, 1), (322, 1), (323, 1), (324, 1), (325, 1), (326, 1), (327, 1), (328, 1), (329, 1), (330, 1), (331, 1), (332, 1), (333, 1), (334, 1), (335, 1), (336, 1), (337, 1), (338, 1), (339, 1), (340, 1), (341, 1), (342, 1), (343, 1), (344, 1), (345, 1), (346, 1), (347, 1), (348, 1), (349, 1), (350, 1), (351, 1), (352, 1), (353, 1), (354, 1), (355, 1), (356, 1), (357, 1), (358, 1), (359, 1), (360, 1), (361, 1), (362, 1), (363, 1), (364, 1), (365, 1), (366, 1), (367, 1), (368, 1), (369, 1), (370, 1), (371, 1), (372, 1), (373, 1), (374, 1), (375, 1), (376, 1), (377, 1), (378, 1), (379, 1), (380, 1), (381, 1), (382, 1), (383, 1), (384, 1), (385, 1), (386, 1), (387, 1), (388, 1), (389, 1), (390, 1), (391, 1), (392, 1), (393, 1), (394, 1), (395, 1), (396, 1), (397, 1), (398, 1), (399, 1), (400, 1), (401, 1), (402, 1), (403, 1), (404, 1), (405, 1), (406, 1), (407, 1), (408, 1), (409, 1), (410, 1), (411, 1), (412, 1), (413, 1), (414, 1), (415, 1), (416, 1), (417, 1), (418, 1), (419, 1), (420, 1), (421, 1), (422, 1), (423, 1), (424, 1), (425, 1), (426, 1), (427, 1), (428, 1), (429, 1), (430, 1), (431, 1), (432, 1), (433, 1), (434, 1), (435, 1), (436, 1), (437, 1), (438, 1), (439, 1), (440, 1), (441, 1), (442, 1), (443, 1), (444, 1), (445, 1), (446, 1), (447, 1), (448, 1), (449, 1), (450, 1), (451, 1), (452, 1), (453, 1), (454, 1), (455, 1), (456, 1), (457, 1), (458, 1), (459, 1), (460, 1), (461, 1), (462, 1), (463, 1), (464, 1), (465, 1), (466, 1), (467, 1), (4
```

Once we have this pre processed state, we save it into any of the suitable formats. So, here Gensim provides a suitable format for saving this text called dictionary and we save

it in that format, see if you see here in Gensim, we use the dictionary method to convert the text to a suitable internal representation and we find that there are 12 unique tokens for the entire corpus, once we have that for ease of handling the data we convert each word to a unique individual representation. So, every unique word will be given an individual ID.

There arbitrarily provided there is no rational for which number goes to which word. So, once we apply this operation dictionary dot token 2 ID which is an in built attribute for this model as given by Gensim you can find that the words like minus is given an id eleven graph is given an id ten and so on. Now we assume that whatever words are going that we are going to use will come from only one of these words.

Now, let us take a new document which is human computer interaction. So, I have a new sentence or a new document human computer interaction. So, I have to apply all those pre processing that I have applied to the corpus, since there are now stop words I do not apply the stop words filtering explicitly, but it is also assumed to be that. Now we convert it to something called as b o w or bag of words representation. So, a document is now going to be represented in terms of the bag of words representation which is very similar to this one.

We can find that you get a list with 2 petals which show 1 comma 1 and 2 comma 1. So, what essentially is represented here? So, essentially what you can find here is that we have 3 terms human computer and interaction, but unfortunately interaction is a word that is not available in our dictionary. So, we ignore that word when we convert this document to a vectorial representation. So, for all practical purposes, you can assume your document to be human computer as this is what it is going to be fed to the topic model.

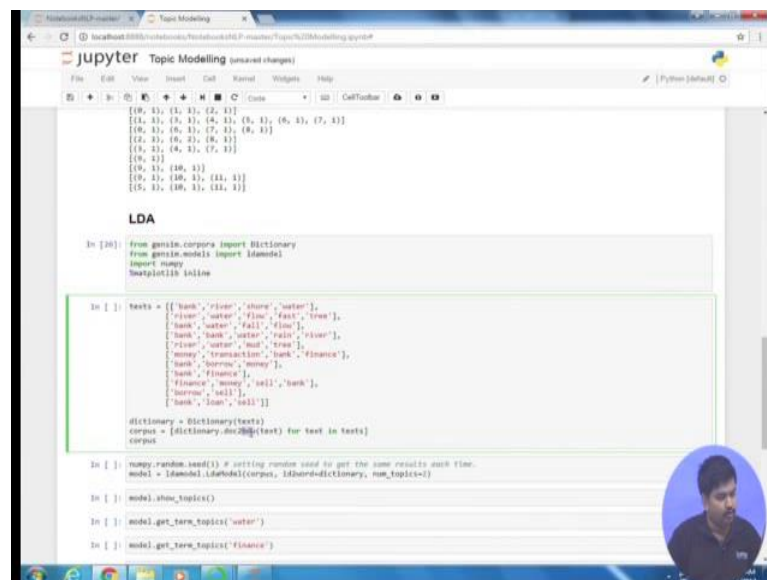
Now, what is this first entity that is 1 and 2 represent they basically represent each of the word human and computer see, if you see the identifier for human is 2 and the identifier for computer is 1 and what you see on the right is the number of times a particular word is appearing for example, if I add another term computer again. So, it is human computer

interaction computer the word with ID 1 which happens to be computer as increases frequency. So, this is at basically frequency count of each unique word that basically represents the documents.

As we have saved the dictionary, we can also save these representations also in a serializable format and we can store it so that we can load it, later whenever we require that so, these are all those lines sentences that we were dealing with here, they are first converted to a list format with it is relevant words. Now we convert it to the individual usage or the individual identifier along with it is frequency it. So, happen that in this corpus we do not have any of the word repeating twice other than this particular word that is the 6th word in the 3rd sentence. So, word 6 is system and in 3rd sentence, yes system; human system e p s. So, that is now reduced to 6 comma 2.

Till now we have not done anything specific topic modeling or what we call as LDA for this particular instance, what we have done is that we have used a means to represent in a format that is much more easier to handle for the model.

(Refer Slide Time: 11:27)

A screenshot of a Jupyter Notebook titled "Topic Modeling" with a "unsaved changes" indicator. The notebook interface includes a menu bar (File, Edit, View, Insert, Cell, Kernel, Windows, Help) and a toolbar. The code cell contains the following Python code:

```
[10, 1], [1, 1], [2, 1]]
[11, 1], [1, 1], [4, 1], [9, 1], [6, 1], [7, 1]]
[9, 1], [6, 1], [7, 1], [8, 1]]
[12, 1], [6, 1], [8, 1]]
[9, 1], [4, 1], [7, 1]]
[9, 1]]
[9, 1], [10, 1]]
[9, 1], [10, 1], [11, 1]]
[5, 1], [10, 1], [11, 1]]

LDA

In [20]: from gensim.corpora import Dictionary
         from gensim.models import LdaModel
         import numpy
         matplotlib inline

In [ ]: texts = [['bank', 'river', 'shore', 'water'],
                ['river', 'water', 'flow', 'fast', 'tree'],
                ['bank', 'water', 'fall', 'flow'],
                ['bank', 'bank', 'water', 'rain', 'river'],
                ['river', 'water', 'bank', 'tree'],
                ['money', 'transaction', 'bank', 'finance'],
                ['bank', 'surround', 'money'],
                ['bank', 'finance'],
                ['finance', 'money', 'sell', 'bank'],
                ['surround', 'sell'],
                ['bank', 'loan', 'sell']]

         dictionary = Dictionary(texts)
         corpus = [dictionary.doc2bow(text) for text in texts]

In [ ]: numpy.random.seed(1) # setting random seed to get the same results each time
         model = LdaModel(corpus, id2word=dictionary, num_topics=2)

In [ ]: model.show_topics()

In [ ]: model.get_term_topics('water')

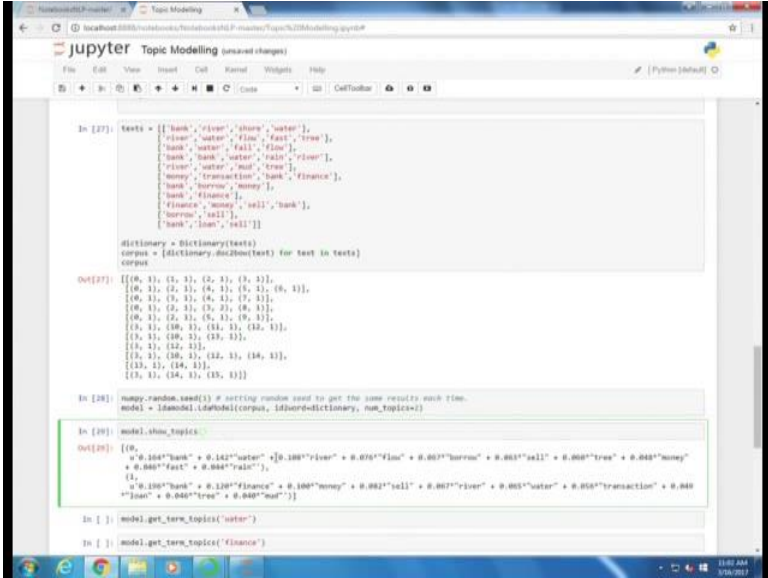
In [ ]: model.get_term_topics('finance')
```

Now, we invoked the LDA model or the latent Dirichlet allocation model for the topic models as you remember, there are different variations of topic model that has been used

as you have seen different methods in your class like the sequential mode sequential topic model and the relation topic model here we will be showing only the simple LDA model.

I import this library now let us assume another corpus where all the pre processing are initially done. So, we have something like a new corpus. So, again a tiny corpus with few sentences and the stop words are removed all the pre processing like that we convert it to the dictionary and we convert it to the bag of words format for each of this document. So, here we look at the representation for this corpus just to make things clear we have created a new dictionary based on the words in this corpus.

(Refer Slide Time: 12:35)



```
In [27]: texts = [['bank', 'river', 'shore', 'water'],
                ['river', 'water', 'flow', 'fast', 'tree'],
                ['bank', 'water', 'fall', 'flow'],
                ['bank', 'bank', 'water', 'rain', 'river'],
                ['river', 'water', 'mad', 'tree'],
                ['money', 'transaction', 'bank', 'finance'],
                ['bank', 'borrow', 'money'],
                ['finance', 'money', 'sell', 'bank'],
                ['borrow', 'sell'],
                ['bank', 'loan', 'sell']]

dictionary = Dictionary.from_instances(texts)
corpus = [dictionary.doc2bow(text) for text in texts]

Out[27]: [[(0, 1), (1, 1), (2, 1), (3, 1)],
          [(0, 1), (2, 1), (4, 1), (5, 1), (6, 1)],
          [(0, 1), (2, 1), (4, 1), (7, 1)],
          [(0, 1), (2, 1), (3, 1), (8, 1)],
          [(0, 1), (2, 1), (5, 1), (9, 1)],
          [(1, 1), (8, 1), (12, 1), (13, 1)],
          [(1, 1), (8, 1), (13, 1)],
          [(1, 1), (8, 1), (12, 1), (13, 1)],
          [(1, 1), (13, 1), (14, 1)],
          [(1, 1), (14, 1), (15, 1)]]

In [28]: numpy.random.seed(1) # setting random seed to get the same results each time.
model = LdaModel(corpus, id2word=dictionary, num_topics=2)

In [29]: model.show_topics()

Out[29]: [(0,
          "0.164*bank + 0.142*water + 0.188*river + 0.078*flow + 0.007*borrow + 0.001*sell + 0.000*tree + 0.000*money"
          + 0.000*fast + 0.000*rain"),
          (1,
          "0.100*bank + 0.130*finance + 0.100*money + 0.001*sell + 0.007*river + 0.001*water + 0.000*transaction + 0.000"
          + 0.000*tree + 0.000*mad")]

In [ ]: model.get_term_topics('water')

In [ ]: model.get_term_topics('finance')
```

Now, comes to the LDA model. So, in the LDA model what we have to essentially provide to the system is the number of topics. So, we assume that the user or the programmer has a he already he is aware of the number of topics that the certain collection of corpus is going to contain and this should be provided as a parameter or as an argument to the model. Now if we see what are the other possible options that you can use to customize this particular function, you can find what are the values the hyper parameter alpha can be given; the hyper parameter eta or the beta as you know it.

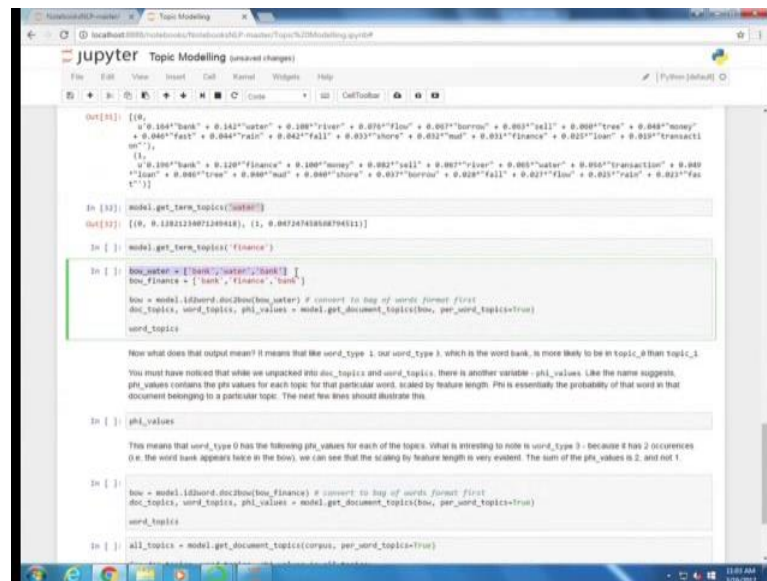
The decay and number of iterations and other convenient parameters that you can use, so

the alpha and eta are the important ones, others are from mostly from programming convenience. So, let us run this model. So, now, we have added this corpus into a doc 2 bow representation and we have provided that this purpose to the LDA model to get a model. This is stored in the variable named as model now if we see we can find since we have already told the system that there will be only 2 topics, it has provided 2 topics to us and it is showing the top terms that the topic represents. So, the topic is essentially a collection of words which is gathered from the different documents and you can see with what probability each word is belonging to that particular topic.

We can also find the same word might belong to multiple topics here the term bank is having a probability of 0.164 and here the term bank is having a probability of 0.196, we can find other words also repeat here like water has 0.065 in topic 1 or the topic 2 and water has a probability of 0.142 in the first topic which is index at 0, now if you see ideally the probability should add up to 1. So, all the words which are represented here when we add up the probability it should give a probability of 1.

But if you add up these values it is not getting adding up to one. So, here we have an option to show, what are the numbers of words or topic to be shown? So, here the model has shown only top 10 words for the given topic. So, in practical scenario we often use number of topics which are more than 10 also like that it can be 50 or 100 topics in that case this function just shows 10 topics; some 10 topics out of the 100 or 50 topics now since we have a small corpus and I am aware that the number of words is going to be less than 20, I just add the parameter number of words is because the argument number of words is equal to 20 and we can find all almost all the words that is present in the particular topic if we increase the topic it does not change which means these are the only number of words in that particular topic.

(Refer Slide Time: 16:24)

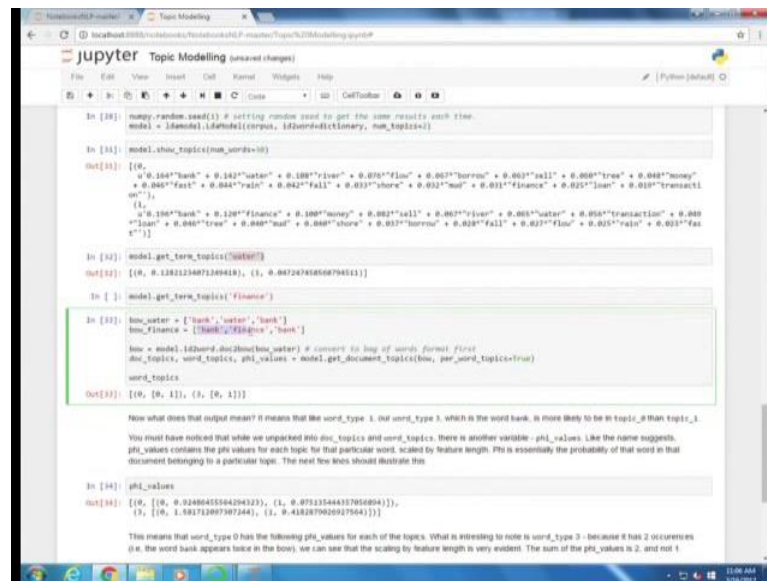


```
Out[11]: [(0, 0.164*'bank' + 0.142*'water' + 0.188*'river' + 0.878*'flow' + 0.887*'burrus' + 0.883*'cell' + 0.888*'tree' + 0.848*'money'  
+ 0.840*'fast' + 0.844*'rain' + 0.842*'fall' + 0.833*'shore' + 0.832*'mud' + 0.831*'finance' + 0.825*'loan' + 0.819*'transacti  
on'),  
(1, 0.106*'bank' + 0.128*'finance' + 0.189*'money' + 0.882*'sell' + 0.883*'river' + 0.881*'water' + 0.856*'transaction' + 0.849  
*'loan' + 0.888*'tree' + 0.888*'mud' + 0.888*'shore' + 0.837*'burrus' + 0.838*'fall' + 0.821*'flow' + 0.823*'rain' + 0.823*'fas  
t')]  
  
In [32]: model.get_term_topics('water')  
Out[32]: [(0, 0.1281238071280418), (1, 0.04724743858794311)]  
  
In [ ]: model.get_term_topics('finance')  
  
In [ ]: bow_water = ['bank', 'water', 'bank']  
bow_finance = ['bank', 'finance', 'bank']  
  
bow = model.id2word.doc2bow(bow_water) # convert to bag of words format first  
doc_topics, word_topics, phi_values = model.get_document_topics(bow, per_word_topics=True)  
word_topics  
  
Now what does that output mean? It means that the word_type 1, our word_type 1, which is the word bank, is more likely to be in topic_0 than topic_1.  
You must have noticed that while we unpacked into doc_topics and word_topics, there is another variable phi_values. Like the name suggests,  
phi_values contains the phi values for each topic for that particular word, scaled by feature length. Phi is essentially the probability of that word in that  
document belonging to a particular topic. The next few lines should illustrate this.  
  
In [ ]: phi_values  
  
This means that word_type 0 has the following phi_values for each of the topics. What is interesting to note is word_type 3 - because it has 2 occurrences  
(i.e. the word bank appears twice in the bow), we can see that the scaling by feature length is very evident. The sum of the phi_values is 2, and not 1.  
  
In [ ]: bow = model.id2word.doc2bow(bow_finance) # convert to bag of words format first  
doc_topics, word_topics, phi_values = model.get_document_topics(bow, per_word_topics=True)  
word_topics  
  
In [ ]: all_topics = model.get_document_topics(corpus, per_word_topics=True)
```

Now, if we want to know what is the probability of a particular word? So, we get the term topics from this function. So, we can find that water has a probability of 0.128 and a 0.047 in either of the topics, now let us see a new document that is bank; water bank which was not provided to the model while it was learning the probability distributions. So, we have a new test document or a new document where we have bank water bank and we first convert it to the b o w representation, once we have that here what we find is that the word 0 and word 3. So, both the terms which are like bank and water both the terms are more likely to belong to topic 0. So, here even if bank has a higher probability at topic one may be because of it is co occurrence with water which has as a high score in topic 0 both bank and water are given a higher probability of being in topic 0 for this particular document.

As you are aware in topic models given a word it can belong to multiple topic distributions and once you find the particular word in the document it might be coming from any one of those topics.

(Refer Slide Time: 18:22)



```
In [28]: numpy.random.seed(1) # setting random seed to get the same results each time
model = LatentDirichletAllocation(n_topics=10, random_state=1)

In [31]: model.show_topics(num_words=10)

Out[31]: [(0,
          '0.164*bank' + 0.142*water' + 0.189*river' + 0.079*flow' + 0.007*harbour' + 0.003*island' + 0.000*tree' + 0.048*money'
          + 0.045*fast' + 0.004*rain' + 0.042*fall' + 0.033*shore' + 0.032*sea' + 0.031*finance' + 0.025*loan' + 0.019*transaction'
          + 0.016*bank' + 0.109*finance' + 0.109*money' + 0.002*island' + 0.003*river' + 0.001*water' + 0.016*transaction' + 0.009
          *loan' + 0.000*tree' + 0.000*sea' + 0.000*shore' + 0.033*harbour' + 0.038*fall' + 0.033*flow' + 0.005*rain' + 0.033*fast'
          + 0.033*sea'
          ),
          (1,
          '0.106*bank' + 0.109*finance' + 0.109*money' + 0.002*island' + 0.003*river' + 0.001*water' + 0.016*transaction' + 0.009
          *loan' + 0.000*tree' + 0.000*sea' + 0.000*shore' + 0.033*harbour' + 0.038*fall' + 0.033*flow' + 0.005*rain' + 0.033*fast'
          + 0.033*sea'
          ),
          ...
          )

In [33]: model.get_topic_topics('000000')
Out[33]: [(0, 0.1381234071209418), (1, 0.007267450508794511)]

In [ ]: model.get_topic_topics('finance')

In [32]: bow_water = ['bank', 'water', 'bank']
bow_finance = ['000000', '000000', 'bank']

bow = model.id2word.doc2bow(bow_water) # convert to bag of words format first
doc_topics, word_topics, phi_values = model.get_document_topics(bow, per_word_topics=True)
word_topics

Out[33]: [(0, [0, 1]), (1, [0, 1])]

Now what does that output mean? It means that the word_type 1, our word_type 3, which is the word bank, is more likely to be in topic_0 than topic_1.
You must have noticed that while we unpacked into doc_topics and word_topics, there is another variable - phi_values. Like the name suggests,
phi_values contains the phi values for each topic for that particular word, scaled by feature length. Phi is essentially the probability of that word in that
document belonging to a particular topic. The next few lines should illustrate this.

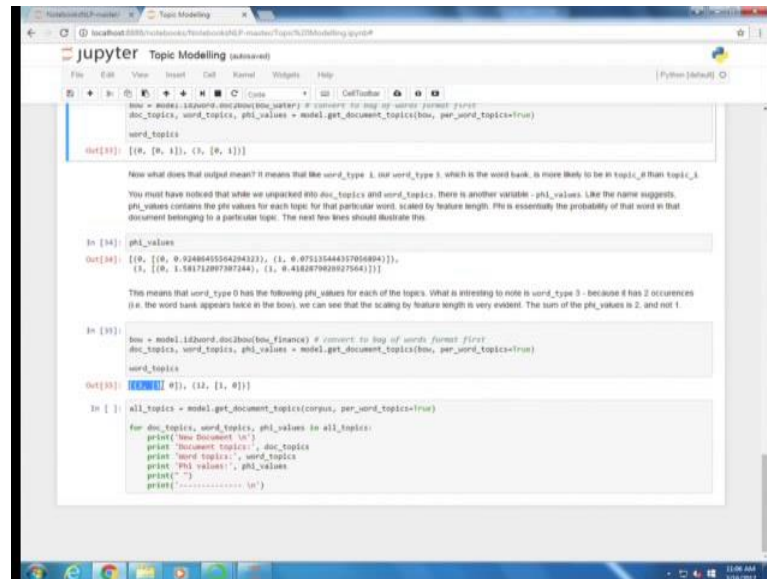
In [34]: phi_values
Out[34]: [(0, [(0, 0.0240045504204323), (1, 0.0723344437050804)]),
          (1, [(0, 1.584732007007244), (1, 0.438379000927044)])]

This means that word_type 0 has the following phi_values for each of the topics. What is interesting to note is word_type 3 - because it has 2 occurrences
i.e. the word bank appears twice in the bow, we can see that the scaling by feature length is very evident. The sum of the phi_values is 2, and not 1.
```

Now, if you see here the 5 values which is essentially the probability of a word in that document belonging to a particular doc topic. So, the 5 value essentially shows, what is the probability of each word belonging to each of those topics? So, if you see here the word type 3 when you add up the values here it has a probability it shows a score not a probability, but a score which adds up to 2.

If the sum of the 5 value is 2 which is and not 1 and this is indicative of the scaling by feature length. So, this always does not happen, but it is an indicative aspect here now what we can see here is that we have taken the second document here which is b o w finance. So, once we have bank finance bank as our document we can find that both the words bank and finance have a higher probability of belonging to topic 1. So, here if you see the word 3 which was bank had a higher probability of belonging to topic 0.

(Refer Slide Time: 19:41)



The screenshot shows a Jupyter Notebook titled 'Topic Modelling'. It contains several code cells and their outputs. The first code cell defines a function to get document topics for a given document ID. The output shows that for document ID 0, the topics are [0, 1] and [1, 0]. The second code cell explains that the output means the word 'bank' is more likely to be in topic 0 than topic 1. The third code cell prints the phi values for the topics, showing that for topic 0, the phi values are [0.9248645504294323, 0.0713144437056894] and for topic 1, they are [0.158712897187244, 0.438278689275641]. The fourth code cell prints the phi values for the topics, showing that for topic 0, the phi values are [0.9248645504294323, 0.0713144437056894] and for topic 1, they are [0.158712897187244, 0.438278689275641]. The fifth code cell prints the phi values for the topics, showing that for topic 0, the phi values are [0.9248645504294323, 0.0713144437056894] and for topic 1, they are [0.158712897187244, 0.438278689275641].

```
bow = model.id2word.doc2bow(bow_safety) # convert to bag of words format first
doc_topics, word_topics, phi_values = model.get_document_topics(bow, per_word_topics=True)
word_topics
Out[33]: [(0, [0, 1]), (1, [0, 1])]

In [34]: phi_values
Out[34]: [(0, [0.9248645504294323, 0.0713144437056894]),
          (1, [0.158712897187244, 0.438278689275641])]

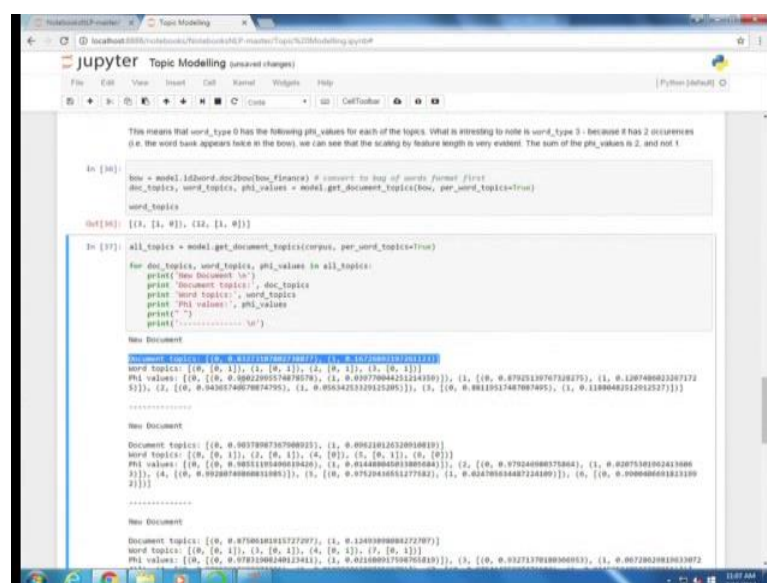
In [35]: bow = model.id2word.doc2bow(bow_finance) # convert to bag of words format first
doc_topics, word_topics, phi_values = model.get_document_topics(bow, per_word_topics=True)
word_topics
Out[35]: [(0, [1, 0]), (1, [1, 0])]

In [ ]: all_topics = model.get_document_topics(corpus, per_word_topics=True)

for doc_topics, word_topics, phi_values in all_topics:
    print('New Document is:')
    print('Document topics:', doc_topics)
    print('Word topics:', word_topics)
    print('Phi values:', phi_values)
    print('-----')
```

But in the second document, this is the model assumes that the word bank comes from topic 1 not topic 0, this is again assumed to be because it co occurs with the word finance. So, we should be talking about the financial institution bank rather than the river side or river bank. So, this becomes evident using this model.

(Refer Slide Time: 20:08)



The screenshot shows a Jupyter Notebook titled 'Topic Modelling'. It contains several code cells and their outputs. The first code cell defines a function to get document topics for a given document ID. The output shows that for document ID 0, the topics are [0, 1] and [1, 0]. The second code cell explains that the output means the word 'bank' is more likely to be in topic 0 than topic 1. The third code cell prints the phi values for the topics, showing that for topic 0, the phi values are [0.9248645504294323, 0.0713144437056894] and for topic 1, they are [0.158712897187244, 0.438278689275641]. The fourth code cell prints the phi values for the topics, showing that for topic 0, the phi values are [0.9248645504294323, 0.0713144437056894] and for topic 1, they are [0.158712897187244, 0.438278689275641]. The fifth code cell prints the phi values for the topics, showing that for topic 0, the phi values are [0.9248645504294323, 0.0713144437056894] and for topic 1, they are [0.158712897187244, 0.438278689275641].

```
bow = model.id2word.doc2bow(bow_finance) # convert to bag of words format first
doc_topics, word_topics, phi_values = model.get_document_topics(bow, per_word_topics=True)
word_topics
Out[36]: [(0, [1, 0]), (1, [1, 0])]

In [37]: all_topics = model.get_document_topics(corpus, per_word_topics=True)

for doc_topics, word_topics, phi_values in all_topics:
    print('New Document is:')
    print('Document topics:', doc_topics)
    print('Word topics:', word_topics)
    print('Phi values:', phi_values)
    print('-----')
```

New Document

Document topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1])]
Word Topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1])]
Phi values: [(0, [0.9248645504294323, 0.0713144437056894]), (1, [0.158712897187244, 0.438278689275641]), (2, [0.158712897187244, 0.438278689275641]), (3, [0.158712897187244, 0.438278689275641])]

New Document

Document topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1])]
Word Topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1])]
Phi values: [(0, [0.9248645504294323, 0.0713144437056894]), (1, [0.158712897187244, 0.438278689275641]), (2, [0.158712897187244, 0.438278689275641]), (3, [0.158712897187244, 0.438278689275641])]

New Document

Document topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1])]
Word Topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1])]
Phi values: [(0, [0.9248645504294323, 0.0713144437056894]), (1, [0.158712897187244, 0.438278689275641]), (2, [0.158712897187244, 0.438278689275641]), (3, [0.158712897187244, 0.438278689275641])]

Now, let us see how what are the different values that we get for each of these documents. So, we can find that the document topics. So, the first document which is essentially bank river shore water if the document representation has a higher tendency to or has a higher proportion of topic 0 rather than topic 1.

Now, individual words were all the 4 words 0, 1, 2 and 3 have a higher tendency to belong the topic 0 the 5 values which is like per the per word distribution for each word in that document also shows the similar branch. Now if you see next word per document which is the document river water flow fast tree what we can find is that again it is having a higher probability of belonging to topic 0 than topic one and we can find that some of the words have no, no occurrence or no probability of occurrence in topic one like the 4 and 6.

(Refer Slide Time: 21:28)

```

New Document
Document topics: [(0, 0.8758018915727297), (1, 0.1241980084272707)]
Word topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1]), (4, [0, 1]), (5, [0, 1]), (6, [0, 1]), (7, [0, 1])]
Phi values: [(0, [0, 0.9781708138617481], (1, 0.0218291750076189)]], (1, [0, 0.912711701380060953], (1, 0.08728823818613072
4)]], (4, [0, 0.9891069388274291], (1, 0.010893061007216078)]], (7, [0, 0.975414891007471182], (1, 0.024585040921280951)]])

New Document
Document topics: [(0, 0.8788892668878287), (1, 0.1211107333321853)]
Word topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1]), (4, [0, 1]), (5, [0, 1]), (6, [0, 1]), (7, [0, 1])]
Phi values: [(0, [0, 0.9750418611277985], (1, 0.02495813887701149)]], (2, [0, 0.9658210710642713], (1, 0.034177138837287
9)]], (3, [0, 1.014320840000706], (1, 0.0405703191005921)]], (4, [0, 0.9784420434029405], (1, 0.0215379575070535)]])

New Document
Document topics: [(0, 0.9544802055107282), (1, 0.0455197448027171)]
Word topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1]), (4, [0, 1]), (5, [0, 1]), (6, [0, 1]), (7, [0, 1])]
Phi values: [(0, [0, 0.9797038867853589], (1, 0.02029699124441355)]], (2, [0, 0.95838951422196477], (1, 0.04161047577803408
5)]], (3, [0, 0.938646428312288], (1, 0.0613543758878295)]], (4, [0, 0.9836861400441086], (1, 0.01631901313283993)]])

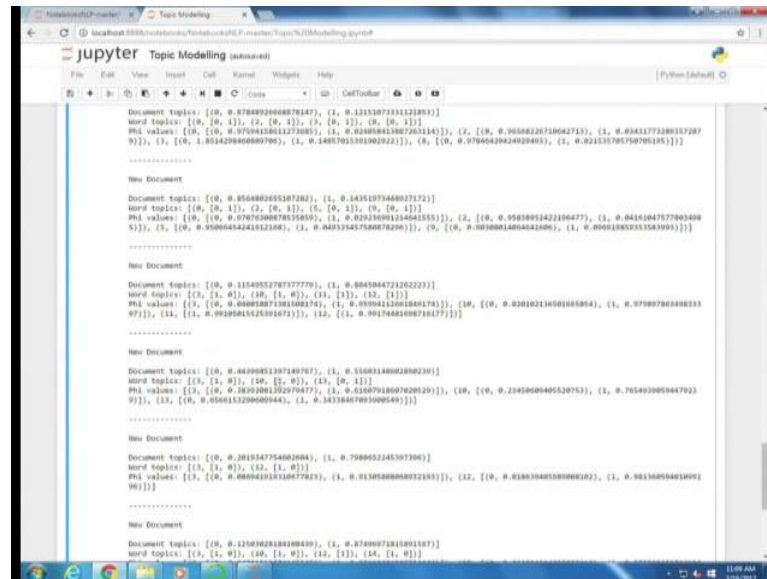
New Document
Document topics: [(0, 0.11548952787177779), (1, 0.8845104721282223)]
Word topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1]), (4, [0, 1]), (5, [0, 1]), (6, [0, 1]), (7, [0, 1])]
Phi values: [(1, [0, 0.940054091134108174], (1, 0.05994512618409174)]], (18, [0, 0.920182136581665854], (1, 0.07988786348833
97)]], (11, [0, 0.9918581525391671], (1, 0.00814188716177)]])

New Document
Document topics: [(0, 0.44398821397140767), (1, 0.556011800280239)]
Word topics: [(0, [0, 1]), (1, [0, 1]), (2, [0, 1]), (3, [0, 1]), (4, [0, 1]), (5, [0, 1]), (6, [0, 1]), (7, [0, 1])]
Phi values: [(1, [0, 0.98392081392079477], (1, 0.01607918007820529)]], (18, [0, 0.23430080485510753], (1, 0.7656930959447023
9)]], (11, [0, 0.938613270089944], (1, 0.0613846780399649)]])

```

Now, if you look into say then one of those documents where the it is probability of belonging to topic 1 is 0.88 while that of belonging to topic 0 is 0.11 and we can find almost all of the words have a higher probability to belong to a topic one it is also possible that suppose similarly for this document.

(Refer Slide Time: 22:01)



```
Document topics: [(0, 0.878892068878647), (1, 0.121107931121893)]
Word topics: [(0, [0, 1]), (2, [0, 1]), (3, [0, 1]), (0, [0, 1])]
Phi values: [(0, [(0, 0.970415861127388), (1, 0.0285841388726114)]), (2, [(0, 0.9658822671864271), (1, 0.034117728815728793)]), (3, [(0, 1.013429846880796), (1, 0.1405701539100922)]), (0, [(0, 0.9786413944020403), (1, 0.02135170575076185)])]

New Document

Document topics: [(0, 0.856880655107282), (1, 0.14311971488927172)]
Word topics: [(0, [0, 1]), (2, [0, 1]), (3, [0, 1]), (0, [0, 1])]
Phi values: [(0, [(0, 0.9787538877881558), (1, 0.0212461114441555)]), (2, [(0, 0.9581895142110477), (1, 0.041818475778034095)]), (3, [(0, 0.9580045424121388), (1, 0.04193545758878296)]), (0, [(0, 0.9810001480444108), (1, 0.018991905333583993)])]

New Document

Document topics: [(0, 0.1154055278737779), (1, 0.8845944721262223)]
Word topics: [(1, [1, 0]), (0, [1, 0]), (1, [1]), (1, [1])]
Phi values: [(1, [(0, 0.6880287338188174), (1, 0.3119712611881781)]), (0, [(0, 0.8201821368146584), (1, 0.17989786348833397)]), (1, [(1, 0.8918081512391871)]), (1, [(1, 0.90174681088718177)])]

New Document

Document topics: [(0, 0.8439081397148787), (1, 0.1560918802898239)]
Word topics: [(1, [1, 0]), (0, [1, 0]), (1, [1]), (1, [1])]
Phi values: [(1, [(0, 0.78392061392079477), (1, 0.21607918097020129)]), (0, [(0, 0.234506040518753), (1, 0.76549390584470239)]), (1, [(1, 0.9506132398898944), (1, 0.049384790909549)])]

New Document

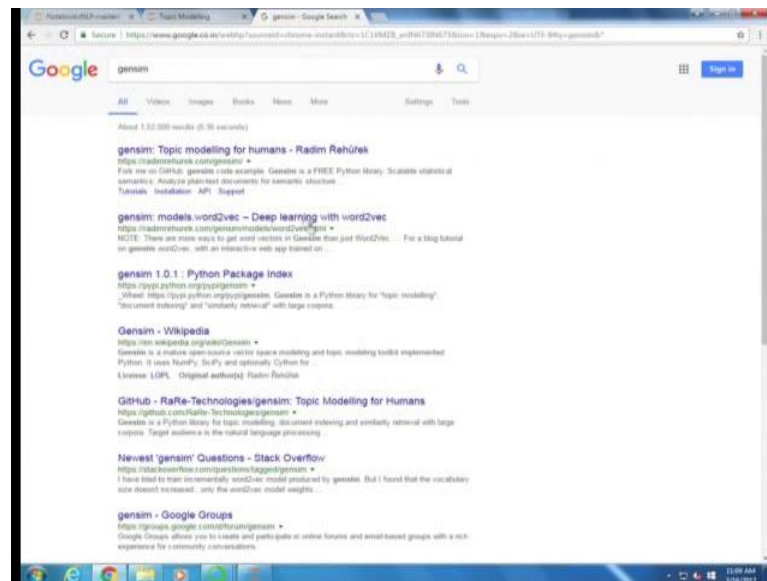
Document topics: [(0, 0.2819347754602684), (1, 0.718065245387306)]
Word topics: [(1, [1, 0]), (1, [1]), (1, [1]), (1, [1])]
Phi values: [(1, [(0, 0.88894193138677621), (1, 0.11105888888888888)]), (1, [(0, 0.81883848588888888)], (1, 0.18116151411111111)])]

New Document

Document topics: [(0, 0.12385808188888888), (1, 0.8761419181111111)]
Word topics: [(1, [1, 0]), (0, [1, 0]), (1, [1]), (1, [1])]
```

We have more or less similar weights for each of the topic like topic 0 it is 0.44 and topic 1, it is 0.55 where the confidence of this document belong to a particular topic is not that high to be decisive enough we can find that 2 of the words have a probability of belonging to topic 1 while that last third word has a higher probability of belonging to topic 0 and we can find the comparative values of each of those doc words with this topic proportion. So, this is how we analyze a corpus using topic modeling or the LDA.

(Refer Slide Time: 22:55)

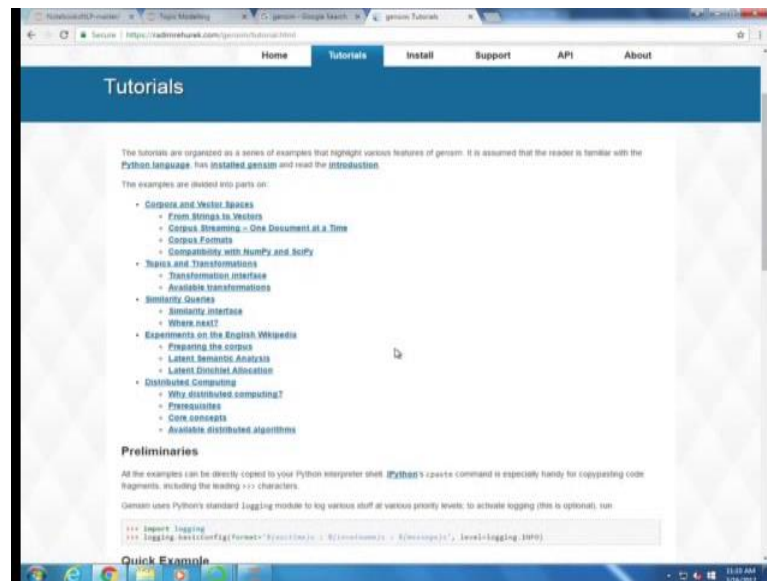


Gensim essentially provides a suit of packages that helps in different topic modeling task and they have a set of comprehensive tutorials. So, the tutorial that you have been seen here has been combined from multiple tutorials from Gensim.

(Refer Slide Time: 23:11)



(Refer Slide Time: 23:13)

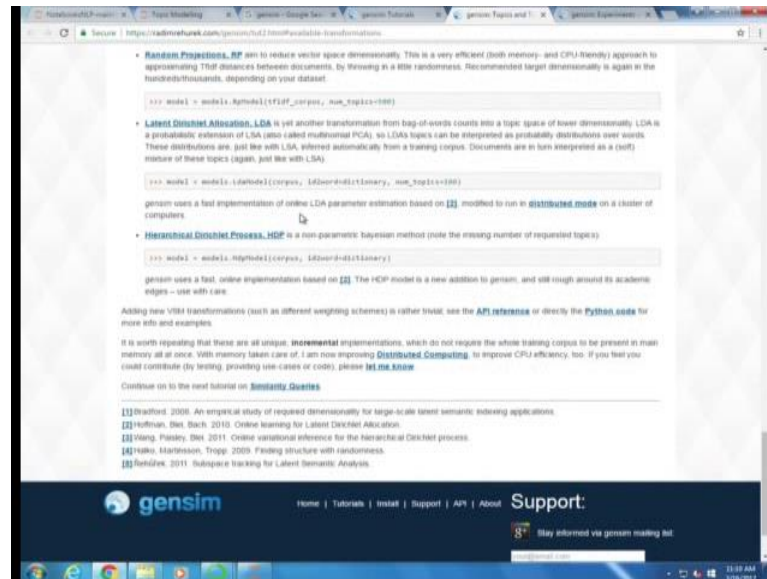


You can have a further look into different topics here, then in addition to LDA provide other topic models as well they have they show topic models e pitching load HDP which is hierarchical display process and they also have a model for sequential LDA. So, we can have a quick look into one of those.

(Refer Slide Time: 23:51)



(Refer Slide Time: 23:48)



Here we can find different transformations or different ways of representing a document what DF, IDF and LSA or LSI are some of the pops they use to be some of the popular ways of representing a document in a vectorial representation and we have LDA which is what we have seen we have HDP. So, HDP is a non parametric version where you do not need to provide the number of topics as well that is it is inferred by the model itself. So, this is pretty much for the topic modeling tutorial.

Thank you.