

Unsupervised Text Analysis

Gov 51: Section 9

Sima Biondi

Spring 2025

Overview

- 1 Housekeeping
- 2 Unsupervised learning overview
- 3 Implementation 1: polisci publishing
Results
- 4 Implementation 2: text analysis for description
Results
- 5 Summary
- 6 Office hours

- Reminder of deadlines

- Reminder of deadlines
 - April 11th → first draft of poster

- Reminder of deadlines
 - April 11th → first draft of poster
 - April 24th → final draft of poster

- Reminder of deadlines
 - April 11th → first draft of poster
 - April 24th → final draft of poster
 - April 25: Pset 3 due

- Reminder of deadlines
 - April 11th → first draft of poster
 - April 24th → final draft of poster
 - April 25: Pset 3 due
 - April 29rd → poster session

Terms and things

- Corpus - a collection of texts that we want to analyze
 - Complete works of Charlotte Bronte, newspaper articles from the AP
- Document - a unit within the corpus
 - Jane Eyre; an article from the AP

Unsupervised learning example 1: peer review bias

- Lots of criticism of the peer review model - send your article in to two double blind reviewers
 - Despite double blind, questions about anonymity - people post their papers online now
 - Reviewers can be quite biased as well!
- Two questions arise:
 1. Who is publishing in top political science journals? Are there ascriptive disparities?
 2. What is being published? Are some topics avoided?

- One advantage of unsupervised learning is pattern finding
 - Data could be high-dimensional, messy, huge observations (e.g. Jane Eyre) ...
- Patterns within the data are incredibly useful!
 - For example, the topics found in a flagship political science journal
- One application is **topic modeling**

Latent Dirichlet Allocation (LDA)

- When we topic model, we need to make some assumptions about how text is generated
 - Recall basic assumptions about topics → probability distribution of words
- Latent Dirichlet Allocation is just one model we can use to topic model
 - LDA relies on a similar assumption of the data generating process of text
- Steps:
 1. For each topic, draw a topic-word distribution → how likely is a topic given a word?
 2. For each document, draw a document-topic distribution → how likely is a document about certain topics?

“I had donuts this morning. I don’t have diabetes yet.”

LDA steps

1. For each topic, draw a topic-word distribution.
 - food: donuts (0.7), have (0.3), diabetes (0)
 - health: diabetes (0.7), have (0.3), donuts (0)
2. For each document in the corpus, draw a document-topic distribution.
 - Sentence 1: food (0.999), health (0.001)
 - Sentence 2: food (0.001), health (0.999)

- 'topicmodels' package is useful for implementing an LDA model
- In our example here, we will be taking a look at a sample of articles in the AP in 1992

```
1 library(topicmodels)
2 data("AssociatedPress")
```

Preview

Data is at the document-word level - for each document, each unique word is counted

```
1 head(tidy(AssociatedPress))
```

A tibble: 6 × 3

	document <int>	term <chr>	count <dbl>
	1	adding	1
	1	adult	2
	1	ago	1
	1	alcohol	1
	1	allegedly	1
	1	allen	1

6 rows

Function

- 'k': number of topics we want to specify
- 'control': setting a seed because probability distributions entail randomness

```
1 # set a seed so that the output of the model is  
  predictable  
2 ap_lda <- LDA(AssociatedPress, k = 2, control =  
  list(seed = 02138))  
3 ap_lda
```

A LDA_VEM topic model with 2 topics.

Results

- 'beta' refers to the probability that a given word is related to a topic
- Let's see which topic "harvard" is associated with

```
1 ap_topics <- tidy(ap_lda, matrix = "beta")  
2  
3 ap_topics |>  
4   filter(term == "harvard")
```

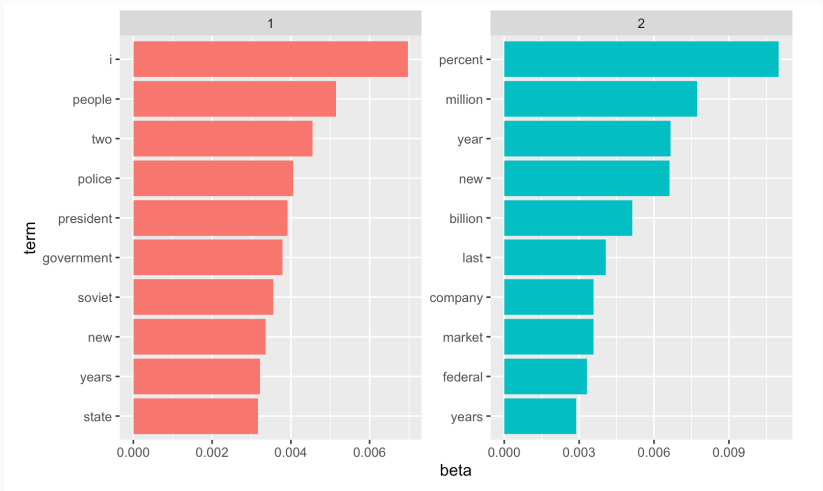
topic	term	beta
<int>	<chr>	<dbl>
1	harvard	4.018027e-05
2	harvard	1.202947e-04

Instead of looking for specific words, let's visualize the most likely terms per topic

```
1 top_terms <- ap_topics |>
2   group_by(topic) |>
3   slice_max(beta, n = 10) |>
4   ungroup() |>
5   arrange(topic, -beta)
6
7 top_terms <- top_terms |>
8   mutate(term = reorder_within(term, beta, topic))
```

```
1 ggplot(top_terms, aes(beta, term, fill =  
    factor(topic))) +  
2   geom_col(show.legend = FALSE) +  
3   facet_wrap(~ topic, scales = "free") +  
4   scale_y_reordered()
```

Visualization



Document level visualization

- How about topic likelihoods at the document level?
- 'gamma' gives us the likelihood of a topic given the words in a document

```
1 ap_documents <- tidy(ap_lda, matrix = "gamma") |>  
2   arrange(document, topic)  
3 head(ap_documents)
```

document <int>	topic <int>	gamma <dbl>
1	1	0.99932253
1	2	0.00067747
2	1	0.51357042
2	2	0.48642958
3	1	0.96200507
3	2	0.03799493

So, the motivating question

1. Who is publishing in top political science journals? Are there ascriptive disparities?
2. What is being published? Are some topics avoided?

So, the motivating question

1. Who is publishing in top political science journals? Are there ascriptive disparities?
2. **What is being published? Are some topics avoided?**

Saraceno (2020) did an analysis of publications in The Journal of Politics

Back to Saraceno (2020)

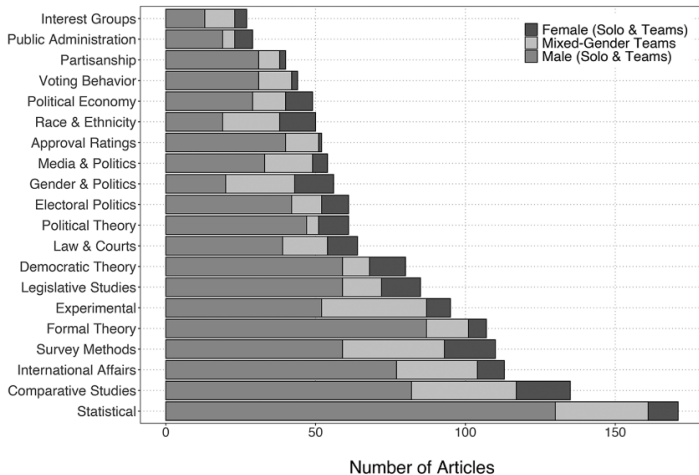


Table 2. Topic Model Output

Topic Label	High-Frequency Words
Interest Groups	Groups, interest, members, organizations, activity, influence
Political Theory	Man, life, human, nature, thought, Hobbes, philosophy
Gender and Politics	Women, social, participation, gender, female, men, politics
Political Economy	Economic, income, fiscal, growth, welfare, market, inequality
Race and Ethnicity	Black, white, racial, ethnic, school, minority, representation
Democratic Theory	Rights, democratic, moral, public, liberal, justice, freedom
The South	States, southern, county, Negro, governor, Virginia, Carolina
Electoral Politics	Elections, candidates, district, incumbent, office, primary
Survey Methods	Respondents, information, survey, attitudes, treatment
Postwar Politics	Soviet, war, world, national, social, united, communist
International Affairs	Conflict, foreign, military, international, war, aid, civil
Judicial Politics	Court, supreme, cases, judicial, law, courts, justice, legal
Formal Theory	Model, decision, equilibrium, preferences, choice, probability
Public Administration	Local, government, service, agencies, city, administer, board
Statistical	Models, data, variables, effect, results, analysis, significant
Legislative Politics	Congress, House, legislative, committee, Senate, majority, bill
Comparative Politics	Countries, international, institution, regime, corrupt, Latin
Partisanship	Issues, ideological, opinion, liberal, conservative, polarization
Media and Politics	Media, coverage, exposure, negative, television, advertising
Voting Behavior	Voter, campaign, candidate, turnout, electorate, ballot
Approval Ratings	President, economic, support, approval, performance
Experimental Methods	Experimental, subjects, control, condition, assigned, effect

- Reiterating caveats at the end of lecture
- Topic modeling is **not** a panacea!
 - Similar to other methods, it relies on assumptions, particularly about the DGP of text
 - Uncertainty is also a part of predictions - similar in respect to regression predictions which have their own standard errors

Packages and preamble

```
1 library(tm)
2 library(SnowballC)
```

- In what ways can we categorize and divide the Harvard Government faculty?
- Let's say we have a **corpus** with three variables in 'csv' form
 1. 'prof' - name of faculty member
 2. 'phd' - year of phd attainment
 3. 'bio' - biography on website

```
1 # Loading in the data
2 df <- read.csv("data/harvardgov.csv")
3
4 # Converting the .csv to document term matrix form
5 corpus <- Corpus(VectorSource(df$bio))
```

Pre-processing

```
1  # make everything lowercase
2  corpus <- tm_map(corpus, content_transformer(tolower))
3
4  # remove white space (e.g. spaces)
5  corpus <- tm_map(corpus, stripWhitespace)
6
7  # remove numbers
8  corpus <- tm_map(corpus, removeNumbers)
9
10 # remove stopwords
11 corpus <- tm_map(corpus, removeWords,
12                  stopwords("english"))
13
14 # stem words (e.g. remove "ing")
15 corpus <- tm_map(corpus, stemDocument)
```

Conversion to DTM

```
1 # Turning into a document term matrix
2 dtm <- DocumentTermMatrix(corpus)
3 dtm.mat <- as.matrix(dtm)
4
5 # Adding labels to each document
6 rownames(dtm.mat) <- df$prof
```

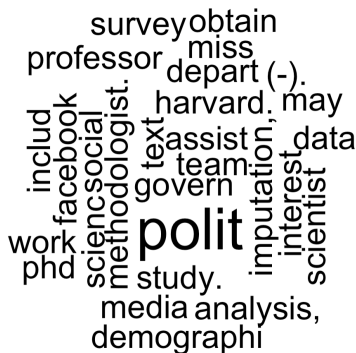
Normalize by document size

```
1 # Normalize by document size
2 tfidf <- weightTfIdf(dtm, normalize = TRUE)
3 tfidf.mat <- as.matrix(tfidf, normalize = TRUE)
4
5 # Adding labels to each document
6 rownames(tfidf.mat) <- df$prof
```

```
1 par(cex = 1.25)
2 library("wordcloud")
3 liu <- dtm.mat["naijia liu",]
4 liu.tfidf <- tfidf.mat["naijia liu",]
```

Visualizing example

```
1 wordcloud(colnames(dtm.mat), liu,  
2           min.freq = min(liu[liu > 0]))
```



Descriptive stats

```
1 sort(liu, decreasing = TRUE)[1:5]
2 sort(liu.tfidf, decreasing = T)[1:3]
```

polit	research	current	includ	professor
2	2	1	1	1
demographi	facebook	imputation,		
0.1745301	0.1745301	0.1745301		

K-means algorithm

- K-Means clustering is simply an exercise in partitioning n observations into k clusters
- In a text context, this means comparing the similarity of words between documents
- Here, we are looking to cluster professors based on similar word usage in their bios!
- This is an iterative process - the algorithm does initial groupings, then sees whether it can minimize error by another permutation

```
1 # Need to standardize so that each row sums to a unit
   length (e.g. 1)
2 tfidf.unit <- tfidf.mat / sqrt(rowSums(tfidf.mat^2))
3
4 set.seed(1234)
5 # centers indicates the number of clusters (e.g. k)
6 kconfour.out <- kmeans(tfidf.unit,
7                        centers = 5)
```

Descriptive stats

1

```
table(kconfour.out$cluster)
```

1	2	3	4	5
11	14	5	7	11



x
stephen ansolabehere
nara dillon
grzegorz ekiert
peter a. hall
jennifer hochschild
alastair iain johnston
taeku lee
elizabeth j. perry
michael rosen
james m. snyder, jr
richard tuck

x
danielle allen
eric beerbohm
daniel carpenter
timothy colton
katrina forrester
claudine gay
harvey c. mansfield
eric nelson
paul peterson
stephen peter rosen
michael sandel
theda skocpol
latanya sweeney
daniel ziblatt

```
knitr::kable(df$prof[kconfour.out$cluster == 1])
```

x

stephen ansolabehere

nara dillon

grzegorz ekiert

peter a. hall

jennifer hochschild

alastair iain johnston

taeku lee

elizabeth j. perry

michael rosen

james m. snyder, jr

richard tuck

K-means group 2

```
knitr::kable(df$prof[kconfour.out$cluster == 2])
```

x

danielle allen

eric beerbohm

daniel carpenter

timothy colton

katrina forrester

claudine gay

harvey c. mansfield

eric nelson

paul peterson

stephen peter rosen

michael sandel

theda skocpol

latanya sweeney

daniel ziblatt

```
knitr::kable(df$prof[kconfour.out$cluster == 3])
```

x

jeffry frieden
frances hagopian
alisha c. holland
torben iversen
steven levitsky

```
knitr::kable(df$prof[kconfour.out$cluster == 4])
```

x

melani cammett
stephen chaudoin
christina davis
joshua d. kertzer
christoph mikulaschek
pia raffler
yuhua wang

```
1 knitr::kable(df$prof[kconfour.out$cluster == 5])
```

x

matthew blackwell

peter buisseret

ryan enos

chase h. harrison

michael j. hiscox

kosuke imai

gary king

naijia liu

mashail malik

stephanie ternullo

dustin tingley

Overtime comparisons

- How similar are each bio to the professor with the earliest PhD, Harvey Mansfield?

```
1 # Isolate Harvey
2 harvey <- as.data.frame(tfidf.mat["harvey c.
   mansfield",])
3 # Isolate non-Harveys
4 nonhm.tfidf <-
   as.data.frame(tfidf.mat[rownames(tfidf.mat)
   != "harvey c. mansfield", ])
5
6 # Sort everything chronologically by PhD attainment
7 nonhm.tfidf$year.index <-
   df$phd[df$prof != "harvey c. mansfield"]
8
9 chron.tfidf <-
   nonhm.tfidf[order(nonhm.tfidf$year.index),]
10 years <- sort(unique(df$phd))
11 years <- years[-1]
```

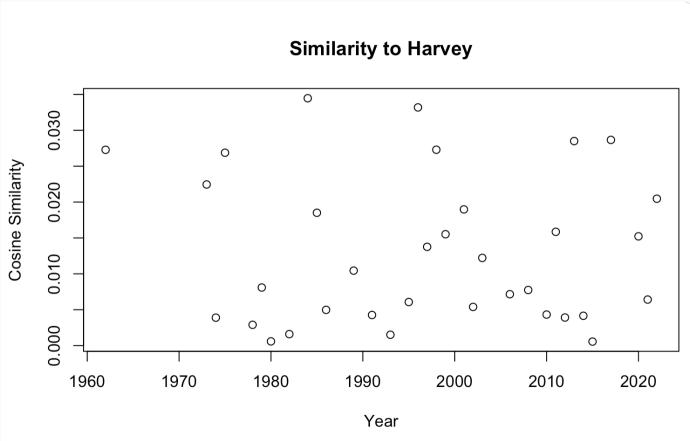
For-loop

```
1 cosine <- function(a, b) {  
2   ## t() transposes a matrix ensuring that vector `a`  
   is multiplied ## by each row of matrix `b`  
3   numer <- apply(a * t(b), 2, sum)  
4   denom <- sqrt(sum(a^2)) * sqrt(apply(b^2, 1, sum))  
5   return(numer / denom)  
6 }
```

```
1 avg.cosim <- rep(NA, length(years))  
2 for (i in 1:length(years)) {  
3   decade <- subset(chron.tfidf,  
4                     (year.index == years[i]))  
5   decade <- decade[, names(decade) != "year.index"]  
6   similarity <- cosine(harvey, decade)  
7   avg.cosim[i] <- mean(similarity)  
8 }
```

Plotting similarity to Harvey across time

```
1 plot(years, avg.cosim, main = "Similarity to Harvey",  
2       xlab = "Year", ylab = "Cosine Similarity")
```



Summary

- Pre-processing text in an easy-to-implement way
- Also, pre-pre-processing when our data isn't already a document term matrix
- Learned one way to group text based on similarity (e.g. k-means algorithm)
- Using for-loops and our own cosine similarity function, we can plot similarity over time

- Bring all your questions!
- Happy to help on code, identification, etc!