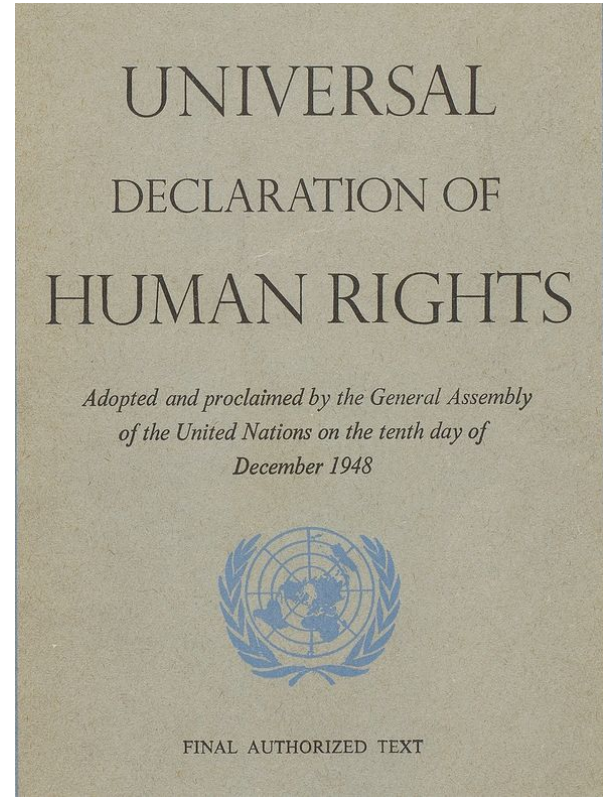


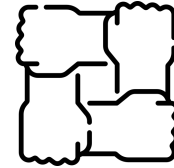
Data Innovation Lab | Team faktual

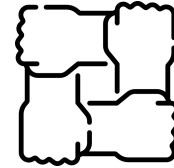
# Using NLP for Adaptive Fact Extraction and Text Summarization

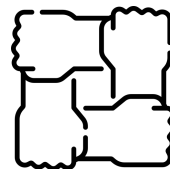
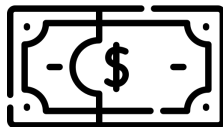
Afag Hasanli, Anelia Petrova, Benedikt Anselment, Marc Schneider and Sergii Poluektov | 31.07.2020

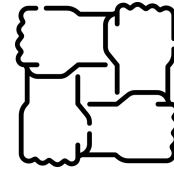
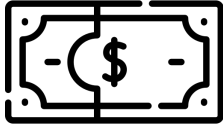


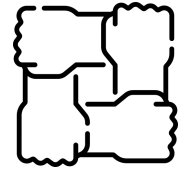
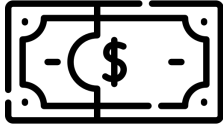




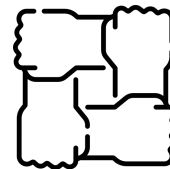
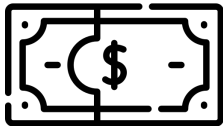








**Automated Summarization**





**Afag Hasanli**

Data Engineering and  
Analytics



**Anelia Petrova**

Informatics



**Benedikt Anselment**

Management &  
Technology



**Marc Schneider**

Student of teaching  
(Mathematics &  
Informatics)



**Sergii Poluektov**

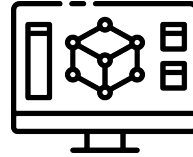
Robotics,  
Cognition,  
Intelligence



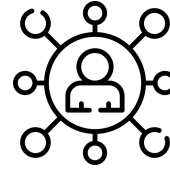
Data



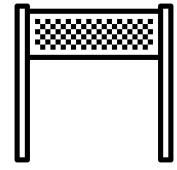
Methodology



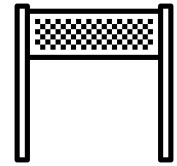
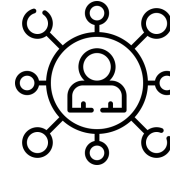
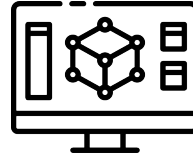
Models



Use cases



Conclusion



Data

Methodology

Models

Use cases

Conclusion

Acquisition

Pipeline

Exploration

Plenarprotokoll 19/126

Deutscher Bundestag

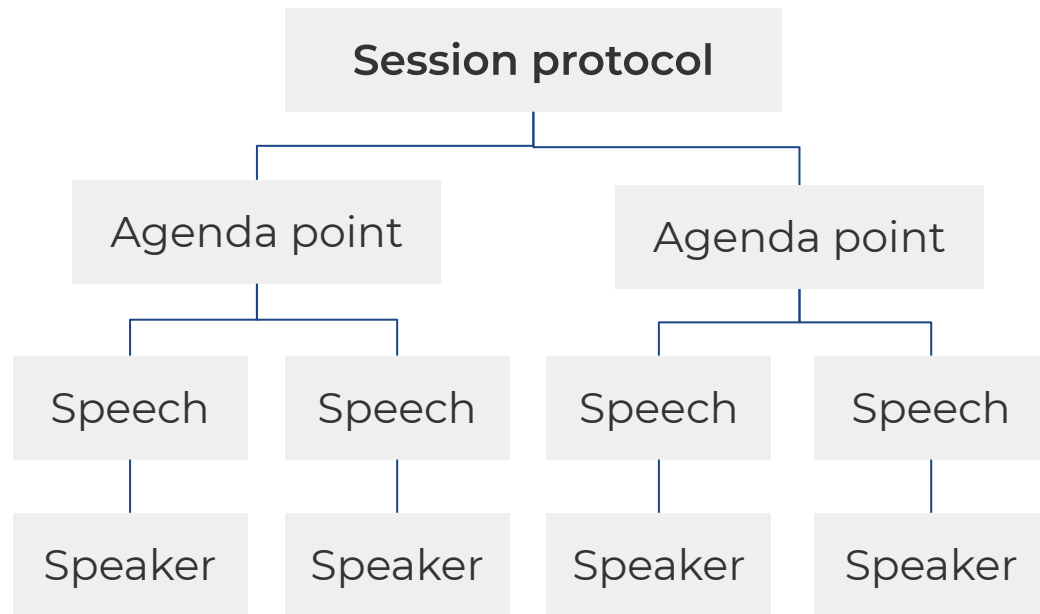
Stenografischer Bericht

126. Sitzung

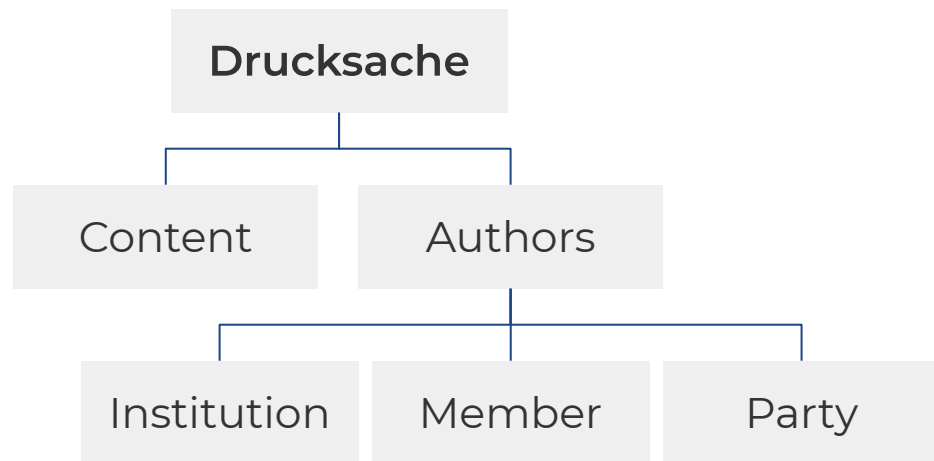
Berlin, Mittwoch, den 13. November 2019

Inhalt:

|                                             |         |                                            |         |
|---------------------------------------------|---------|--------------------------------------------|---------|
| Erweiterung der Tagesordnung .....          | 15689 A | Tagesordnungspunkt 6:                      |         |
| Ausschussbeurteilung .....                  | 15689 B | Erste Beratung des von der Bundesregierung |         |
|                                             |         | eingebrochenen Entwurfs eines Gesetzes zur |         |
|                                             |         | Modernisierung des Strafverfahrens         |         |
|                                             |         | Drucksachen 19/14972, 19/15082 .....       | 15690 A |
| Tagesordnungspunkt 1:                       |         | Tagesordnungspunkt 6:                      |         |
| Erste Beratung des von der Bundesregierung  |         | Befragung der Bundesregierung              |         |
| eingebrochenen Entwurfs eines Gesetzes zur  |         | Peter Altmaier, Bundesminister BAfG .....  | 15690 B |
| Einführung eines Bundes-Klimaschutzge-      |         | Steffen Kotté (AD) .....                   | 15691 B |
| setzes und zur Änderung weiterer Ver-       |         | Peter Altmaier, Bundesminister BAfG .....  | 15691 C |
| schriften                                   |         | Steffen Kotté (AD) .....                   | 15691 D |
| Drucksachen 19/14948, 19/15079 .....        | 15689 D | Peter Altmaier, Bundesminister BAfG .....  | 15692 A |
| Tagesordnungspunkt 2:                       |         | Klaus Kretz (DIE LINKE) .....              | 15692 B |
| Erste Beratung des von der Bundesregierung  |         | Peter Altmaier, Bundesminister BAfG .....  | 15692 B |
| eingebrochenen Entwurfs eines Gesetzes über |         | Oliver Kricher (BÜNDNIS 90                 |         |
| einen nationalen Zertifikatshandel für      |         | DIE GRÜNEN) .....                          | 15692 C |
| Brennstoffemissionen (Brennstoffemissionen- |         | Peter Altmaier, Bundesminister BAfG .....  | 15692 C |
| handelsgepäck – BEHG)                       |         | Johann Sartorius (SPD) .....               | 15693 A |
| Drucksachen 19/14949, 19/15081 .....        | 15689 D | Peter Altmaier, Bundesminister BAfG .....  | 15693 A |
| Tagesordnungspunkt 3:                       |         | Johann Sartorius (SPD) .....               | 15693 C |
| Erste Beratung des von der Bundesregierung  |         | Peter Altmaier, Bundesminister BAfG .....  | 15693 D |
| eingebrochenen Entwurfs eines Gesetzes zur  |         | Lorenz Götsche (DIE LINKE) .....           | 15693 D |
| Umsetzung des Klimaschutzprogramms          |         | Peter Altmaier, Bundesminister BAfG .....  | 15694 A |
| 2030 im Steuerrecht                         |         | Oliver Kricher (BÜNDNIS 90                 |         |
| Drucksachen 19/14957, 19/15080 .....        | 15690 A | DIE GRÜNEN) .....                          | 15694 C |
| Tagesordnungspunkt 4:                       |         | Peter Altmaier, Bundesminister BAfG .....  | 15694 D |
| Erste Beratung des von der Bundesregierung  |         | Dr. Julia Verlinden (BÜNDNIS 90            |         |
| eingebrochenen Entwurfs eines Gesetzes zur  |         | DIE GRÜNEN) .....                          | 15695 A |
| Änderung des Luftverkehrsmessnetzes         |         | Peter Altmaier, Bundesminister BAfG .....  | 15695 C |
| Drucksachen 19/14958, 19/15083 .....        | 15690 A | Manuel Hörlitz (FDP) .....                 | 15696 A |

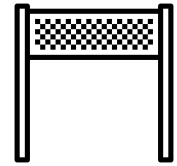
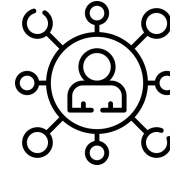
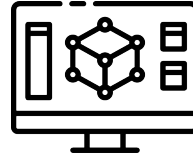


Approximately 100 pages



Approximately 15 pages

| Legislative Period | Years              | Protocol count | Drucksache count |
|--------------------|--------------------|----------------|------------------|
| 14                 | 1998 - 2002        | 253            | 10005            |
| 15                 | 2002 - 2005        | 187            | 6016             |
| 16                 | 2005 - 2009        | 233            | 14163            |
| 17                 | 2009 - 2013        | 253            | 14839            |
| 18                 | 2013 - 2017        | 245            | 13706            |
| 19                 | 2017 -             | 152+           | 21069+           |
| <b>14 - 19</b>     | <b>1998 - 2020</b> | <b>1323+</b>   | <b>79798+</b>    |



Data

Methodology

Models

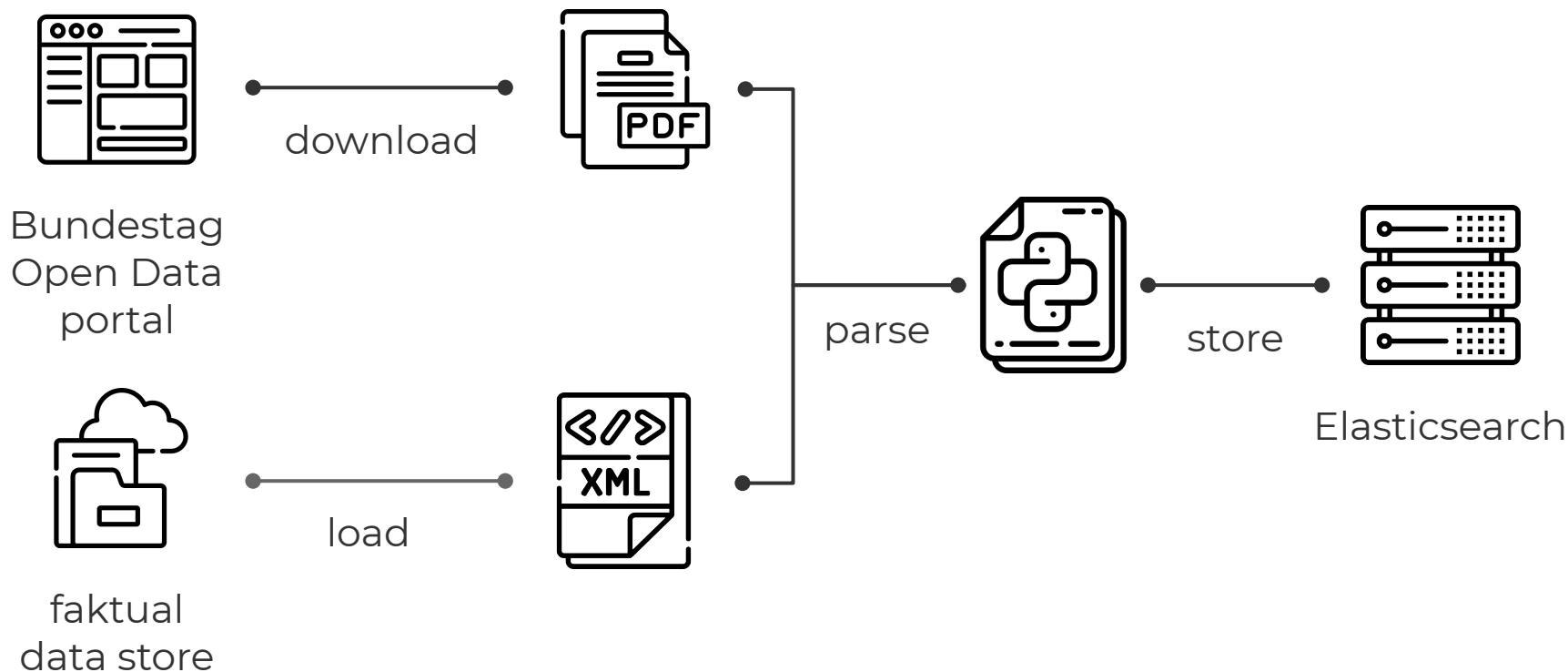
Use cases

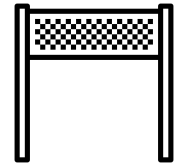
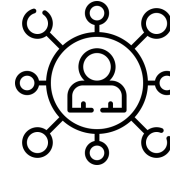
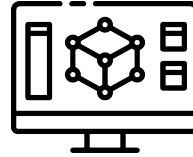
Outlook

Acquisition

Pipeline

Exploration





Data

Methodology

Models

Use cases

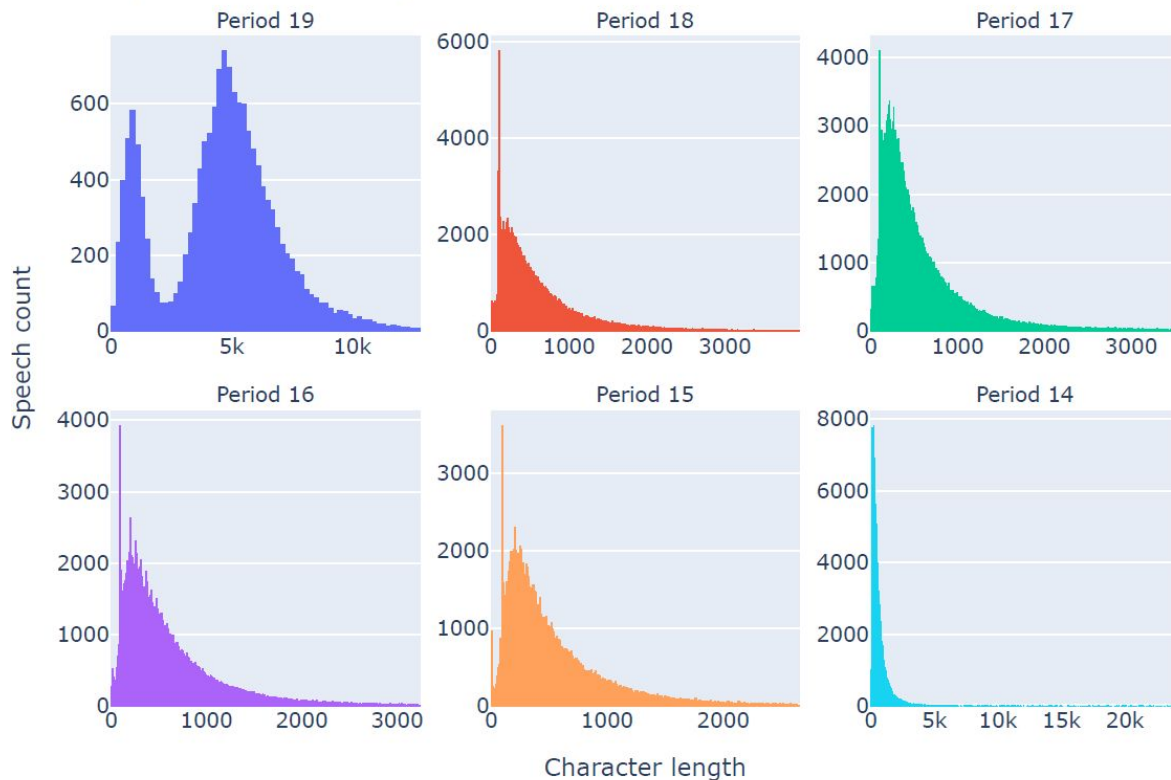
Conclusion

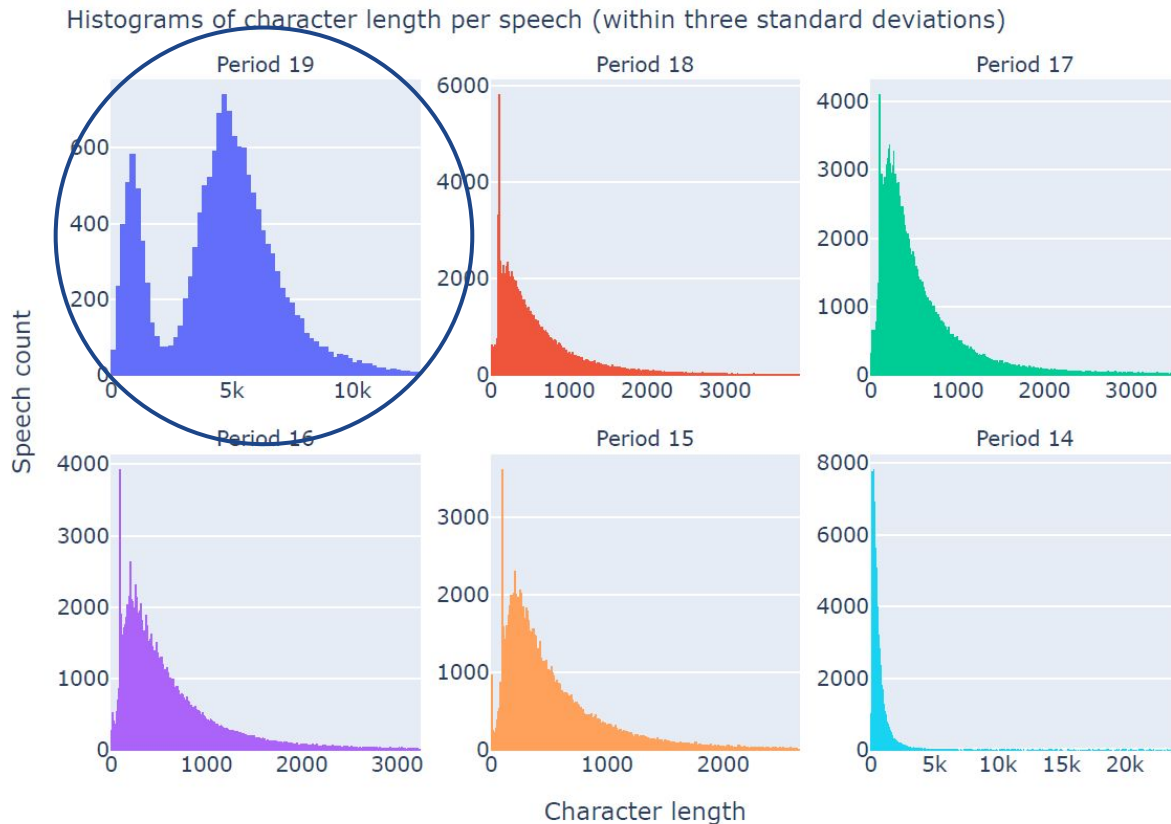
Acquisition

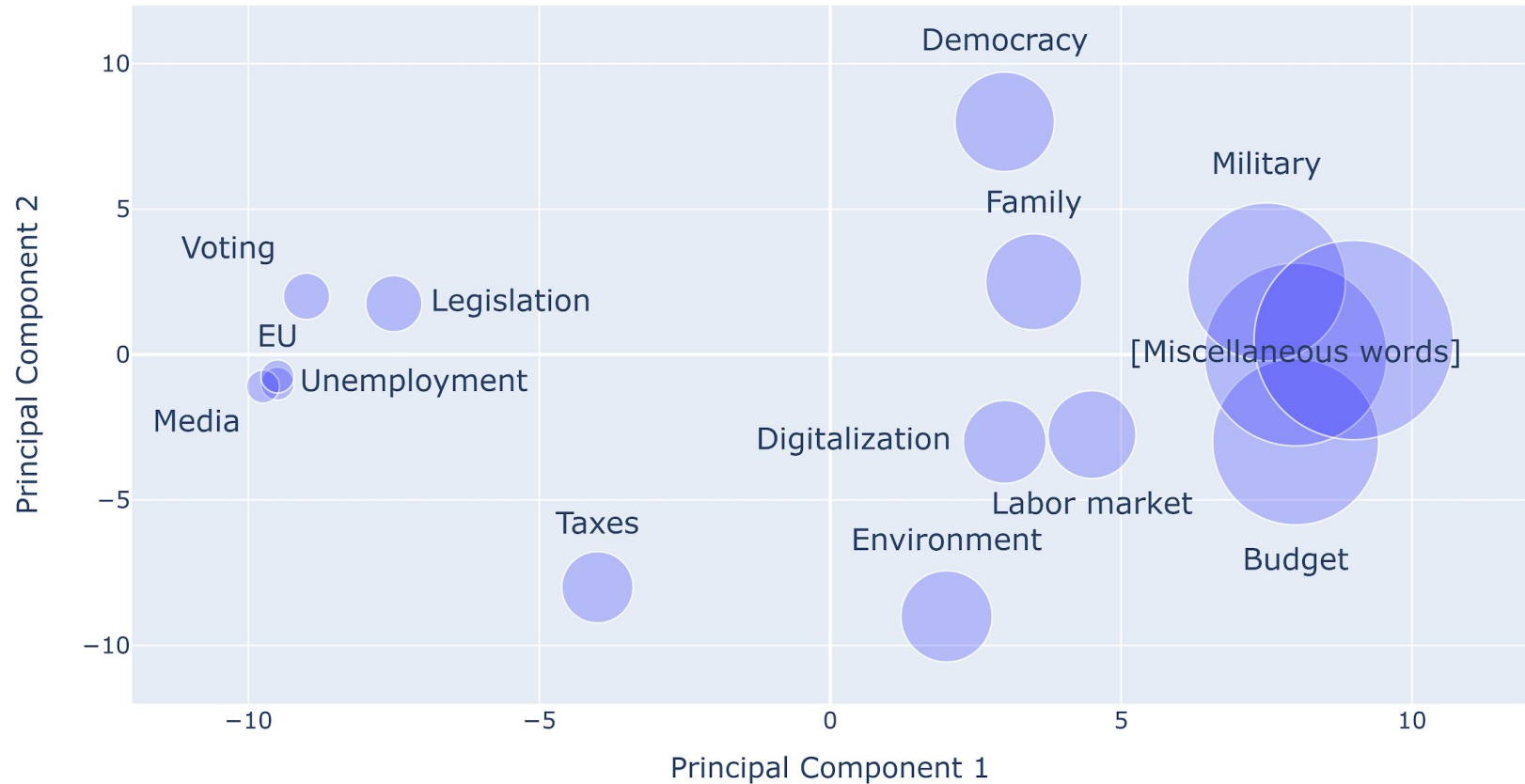
Pipeline

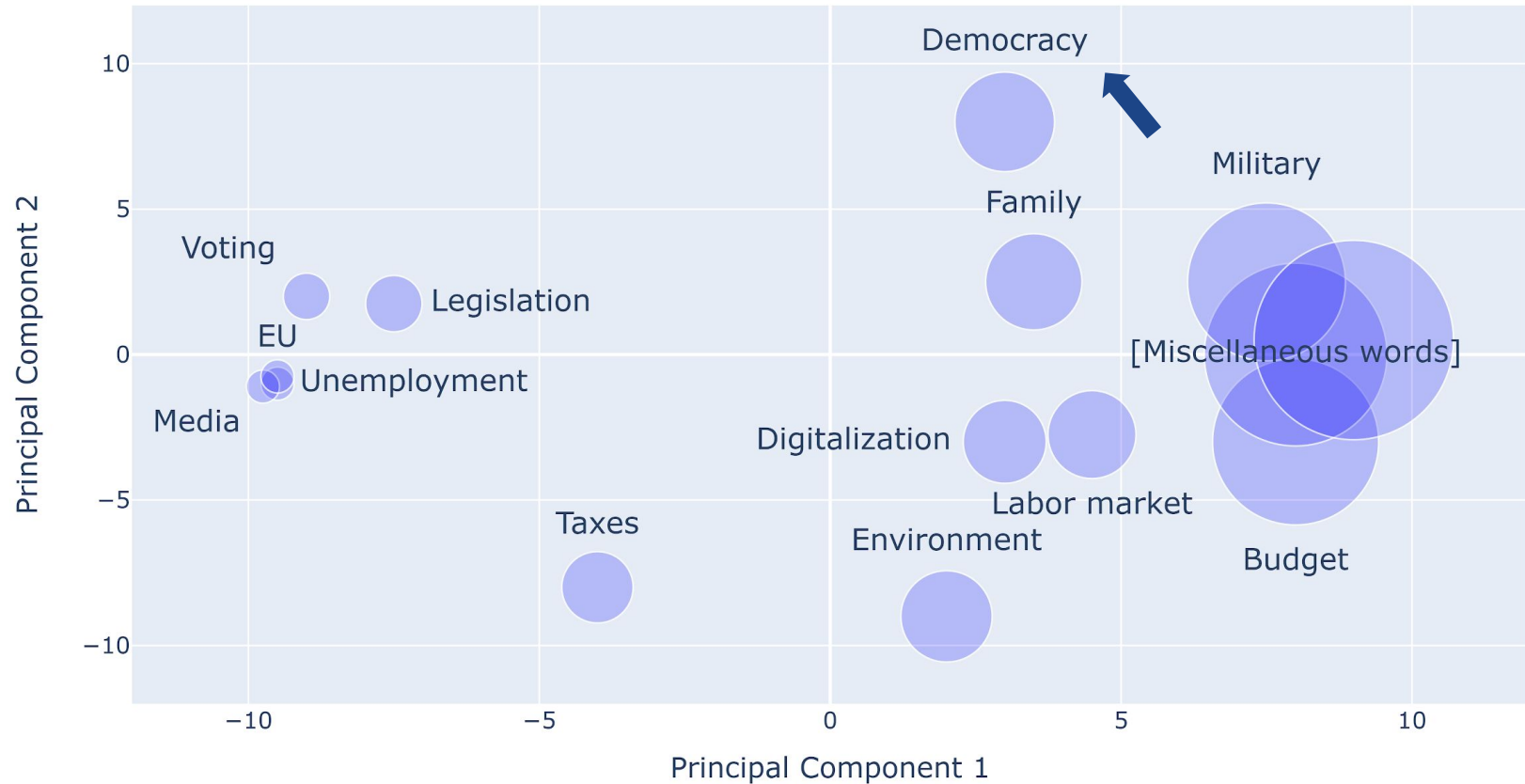
Exploration

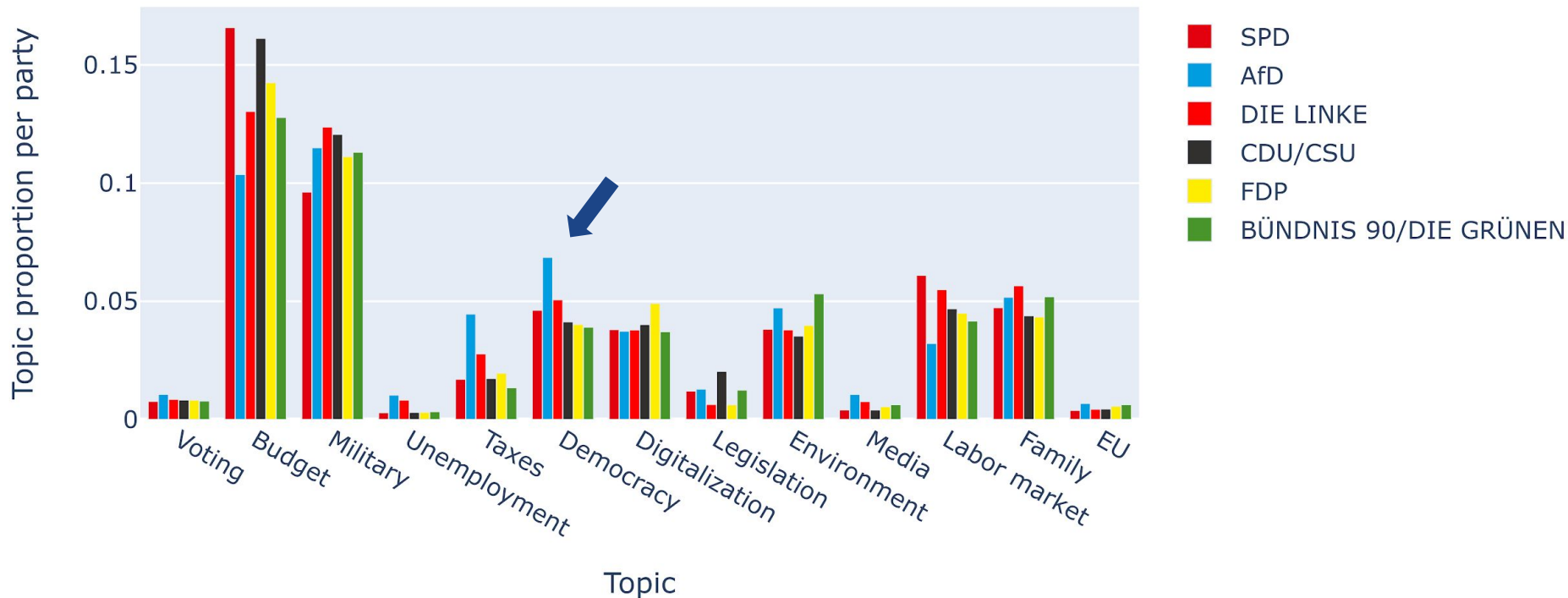
Histograms of character length per speech (within three standard deviations)

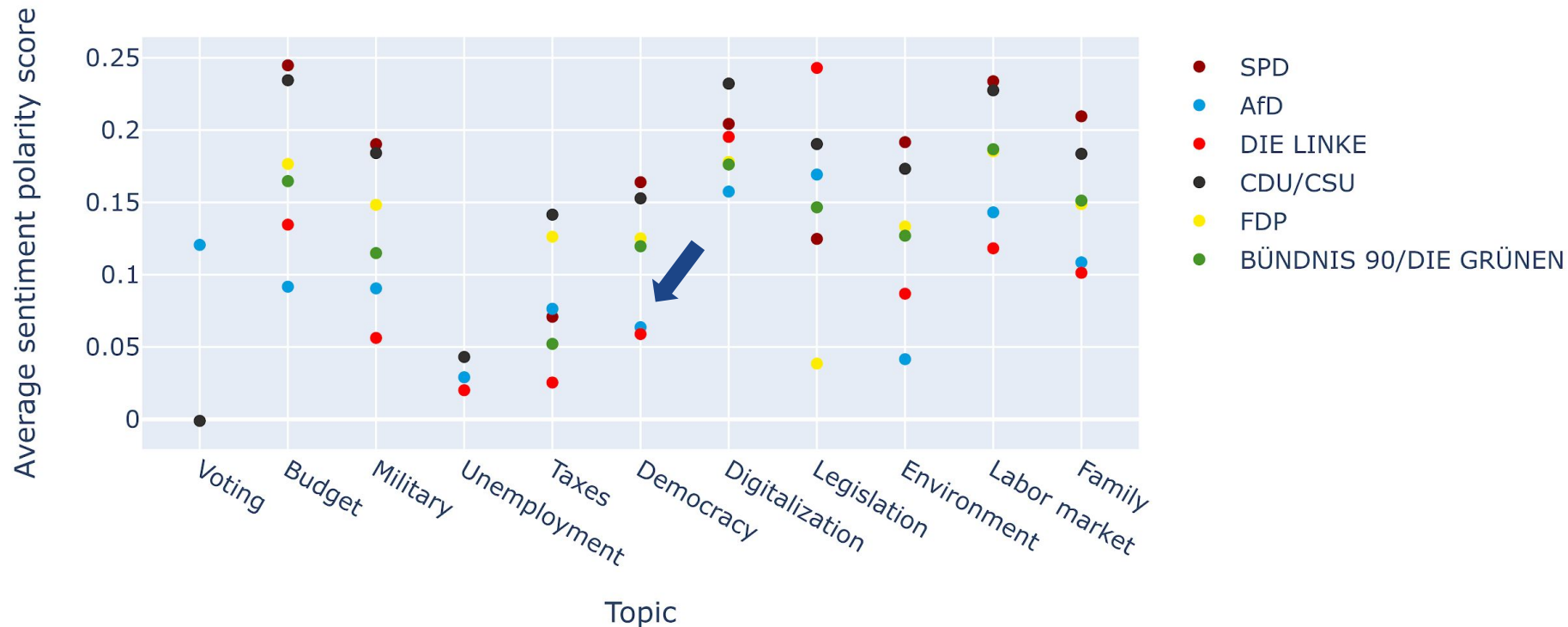










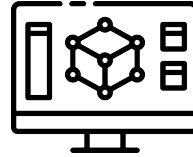




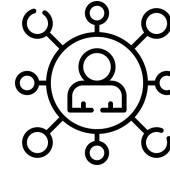
Data



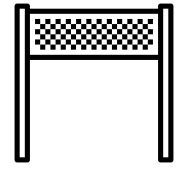
Methodology



Models



Use cases



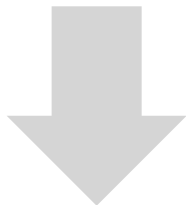
Conclusion

Model selection

Evaluation

## Extractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



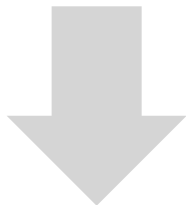
## Abstractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



## Extractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



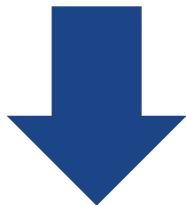
## Abstractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



## Extractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



*"The two main model-categories in automated summarization are extractive and abstractive summarization."*

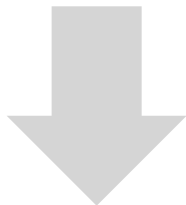
## Abstractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



## Extractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



*"The two main model-categories in automated summarization are extractive and abstractive summarization."*

## Abstractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



## Extractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



*"The two main model-categories in automated summarization are extractive and abstractive summarization."*

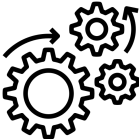
## Abstractive summarization

*"This is an example text. The two main model-categories in automated summarization are extractive and abstractive summarization. Thank you for reading."*



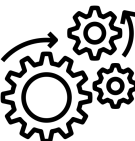
*"The field of automated text summarization can be divided in an extractive and an abstractive domain."*

## Abstractive summarization

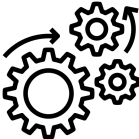
|                                    |  |  |  |
|------------------------------------|-----------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| Approach from scratch              | ●                                                                                 | ●                                                                                   | ●                                                                                   |
| Unsupervised approach (MeanSum)    | ●                                                                                 | ●                                                                                   | ●                                                                                   |
| German pretrained approach (BERT)  | ●                                                                                 | ●                                                                                   | ●                                                                                   |
| English pretrained approach (BART) | ●                                                                                 | ●                                                                                   | ●                                                                                   |

## Abstractive summarization



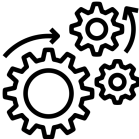
|                                    |  |  |  |
|------------------------------------|-----------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| Approach from scratch              |  |  |  |
| Unsupervised approach (MeanSum)    |  |  |  |
| German pretrained approach (BERT)  |  |  |  |
| English pretrained approach (BART) |  |  |  |

## Abstractive summarization

|                                    |  |  |  |
|------------------------------------|-----------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| Approach from scratch              |  |  |  |
| Unsupervised approach (MeanSum)    |  |  |  |
| German pretrained approach (BERT)  |  |  |  |
| English pretrained approach (BART) |  |  |  |

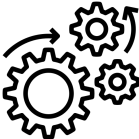


## Abstractive summarization

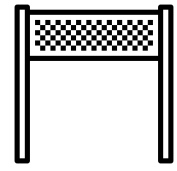
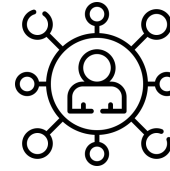
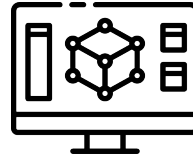
|                                    |  |  |  |
|------------------------------------|-----------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| Approach from scratch              |  |  |  |
| Unsupervised approach (MeanSum)    |  |  |  |
| German pretrained approach (BERT)  |  |  |  |
| English pretrained approach (BART) |  |  |  |



## Abstractive summarization

|                                    |  |  |  |
|------------------------------------|-----------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| Approach from scratch              |  |  |  |
| Unsupervised approach (MeanSum)    |  |  |  |
| German pretrained approach (BERT)  |  |  |  |
| English pretrained approach (BART) |  |  |  |





Data

Methodology

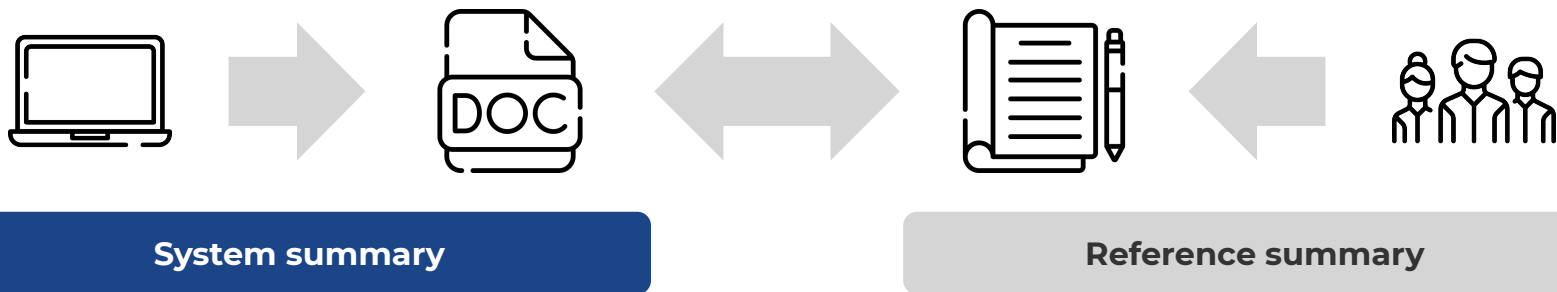
Models

Use cases

Conclusion

Model selection

Evaluation



ROUGE

Measures n-gram overlap



Computational effort



Abstraction capturing

BERTscore

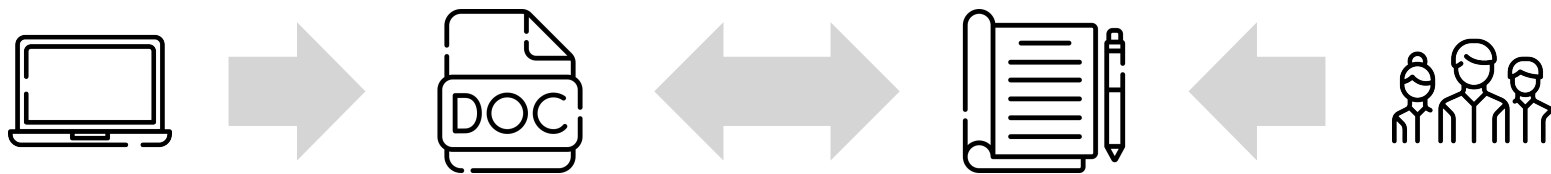
Compares sentence embeddings



Computational effort



Abstraction capturing



**System summary**

**Reference summary**

**ROUGE**

Measures n-gram overlap



**Computational effort**



**Abstraction capturing**

**BERTscore**

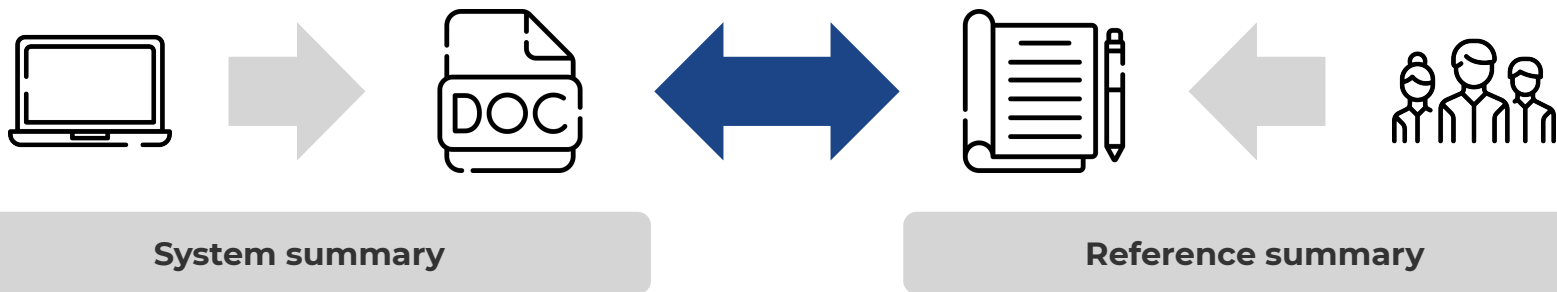
Compares sentence embeddings



**Computational effort**



**Abstraction capturing**



ROUGE

Measures n-gram overlap



Computational effort



Abstraction capturing

BERTscore

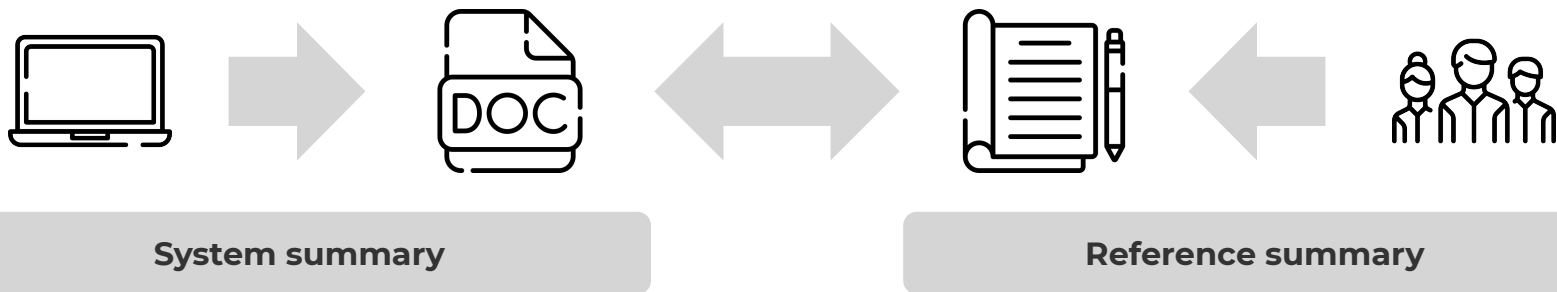
Compares sentence embeddings



Computational effort



Abstraction capturing



## ROUGE

Measures n-gram overlap



Computational effort



Abstraction capturing

## BERTscore

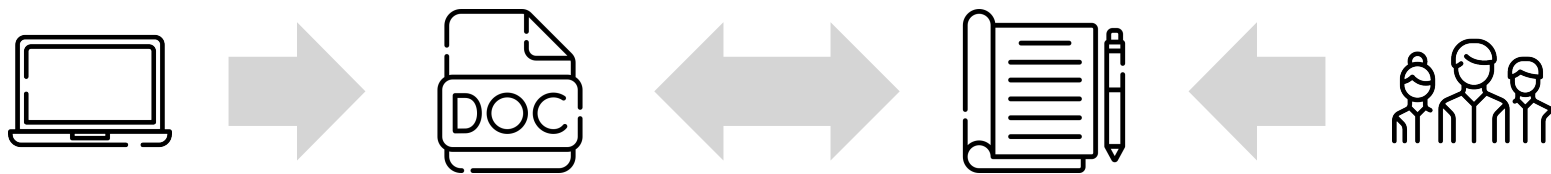
Compares sentence embeddings



Computational effort



Abstraction capturing



System summary

Reference summary

**ROUGE**

Measures n-gram overlap



Computational effort



Abstraction capturing

BERTscore

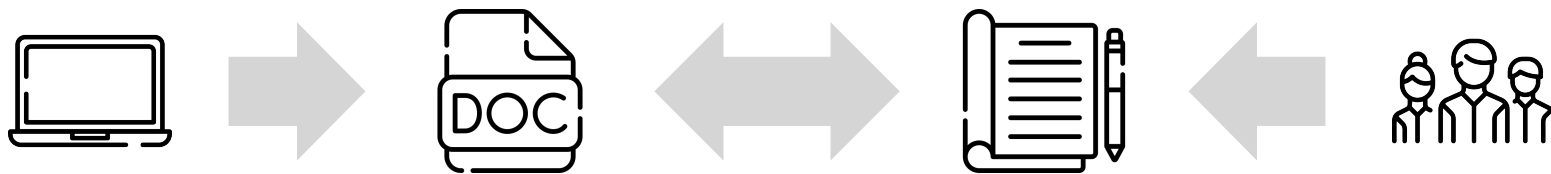
Compares sentence embeddings



Computational effort



Abstraction capturing



System summary

Reference summary

ROUGE

Measures n-gram overlap



Computational effort



Abstraction capturing

BERTscore

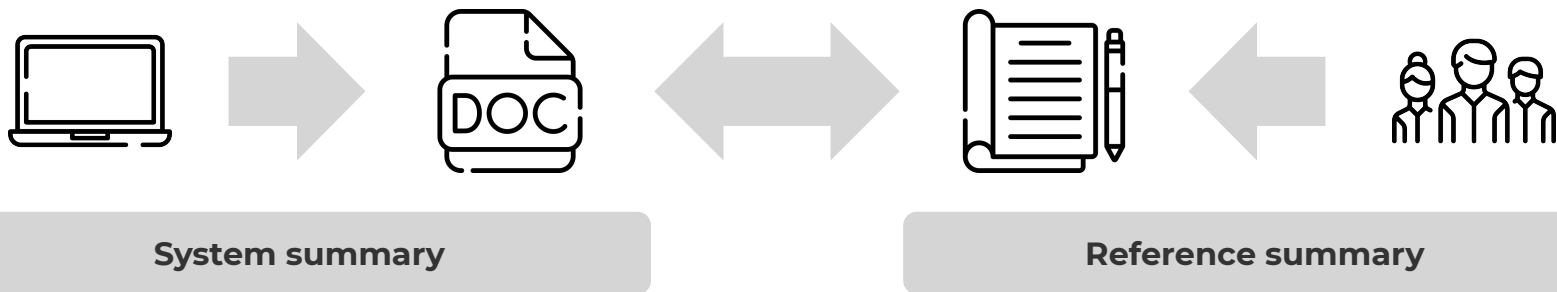
Compares sentence embeddings



Computational effort



Abstraction capturing



ROUGE

Measures n-gram overlap



Computational effort



Abstraction capturing

BERTscore

Compares sentence embeddings



Computational effort



Abstraction capturing



## Language-related questions

---

1 Grammatical correctness

---

2 Reading-fluency

---

## Content-related questions

---

3 Key-point capturing

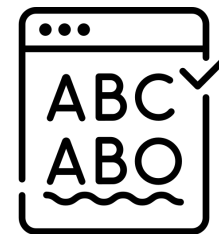
---

4 Content mix-up

---

5 Misleading information added

---

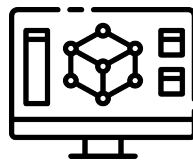




Data



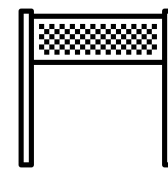
Methodology



Models



Use cases



Conclusion

Model from scratch

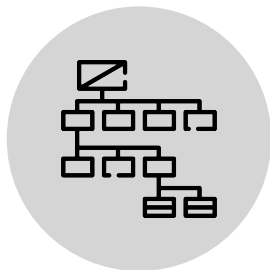
MeanSum

BERT & BART

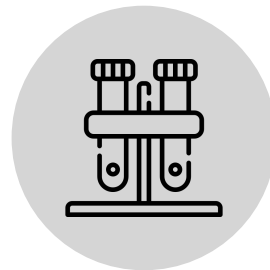
Overall comparison



Dataset



Architecture



Evaluation



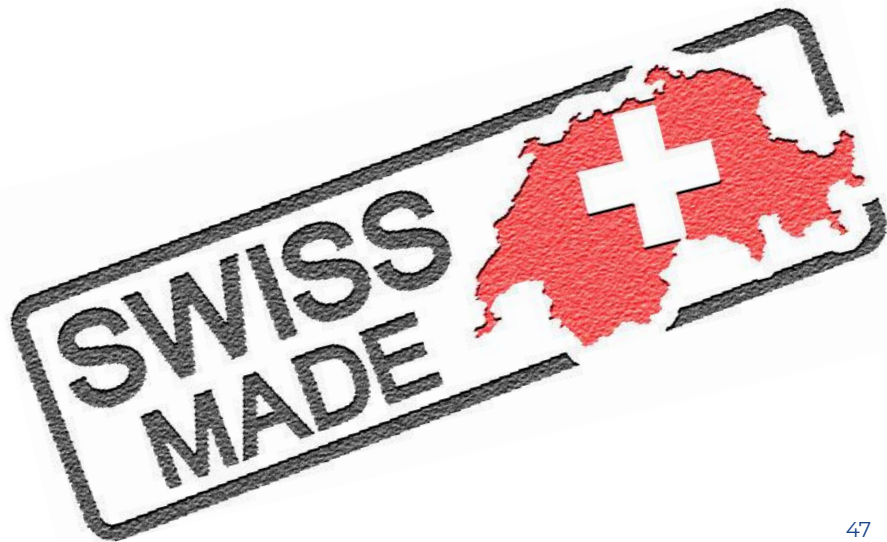
Benchmark

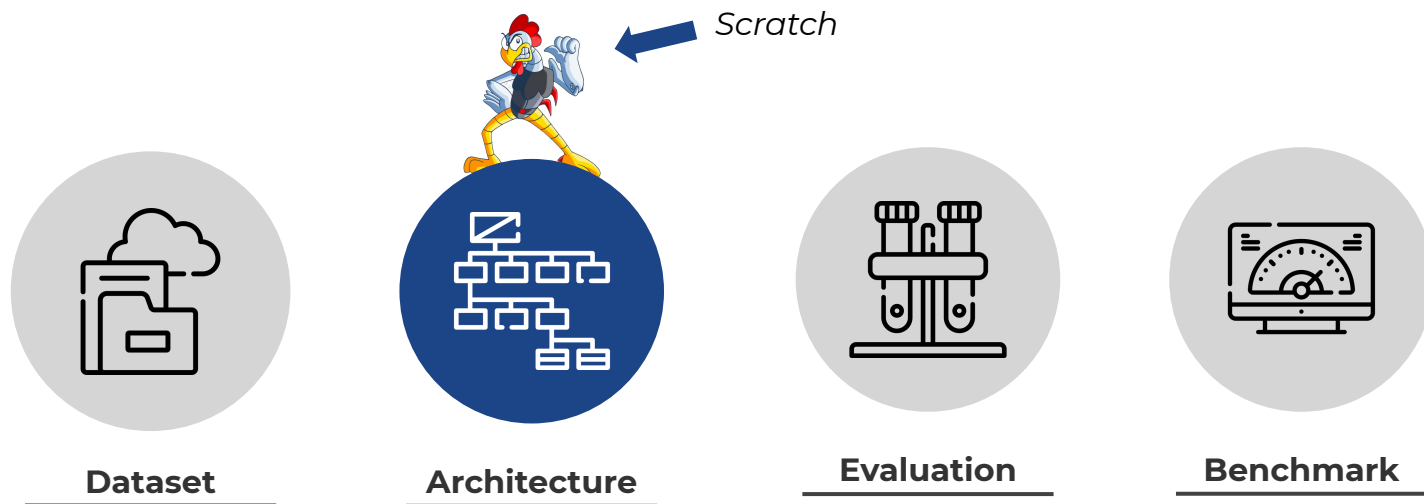
# The “SwissText”-dataset

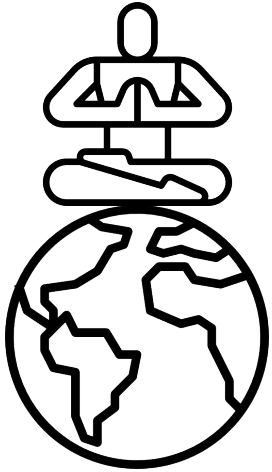
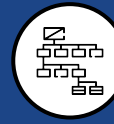


Created for “German text summarization challenge 2019” by researchers from UA Zurich

100,000 texts with corresponding reference summaries extracted from Wikipedia

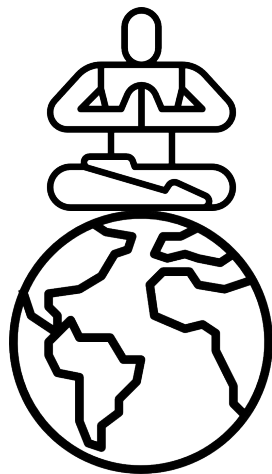






**Attention Layer**

**Encoder-Decoder**



**Attention Layer**

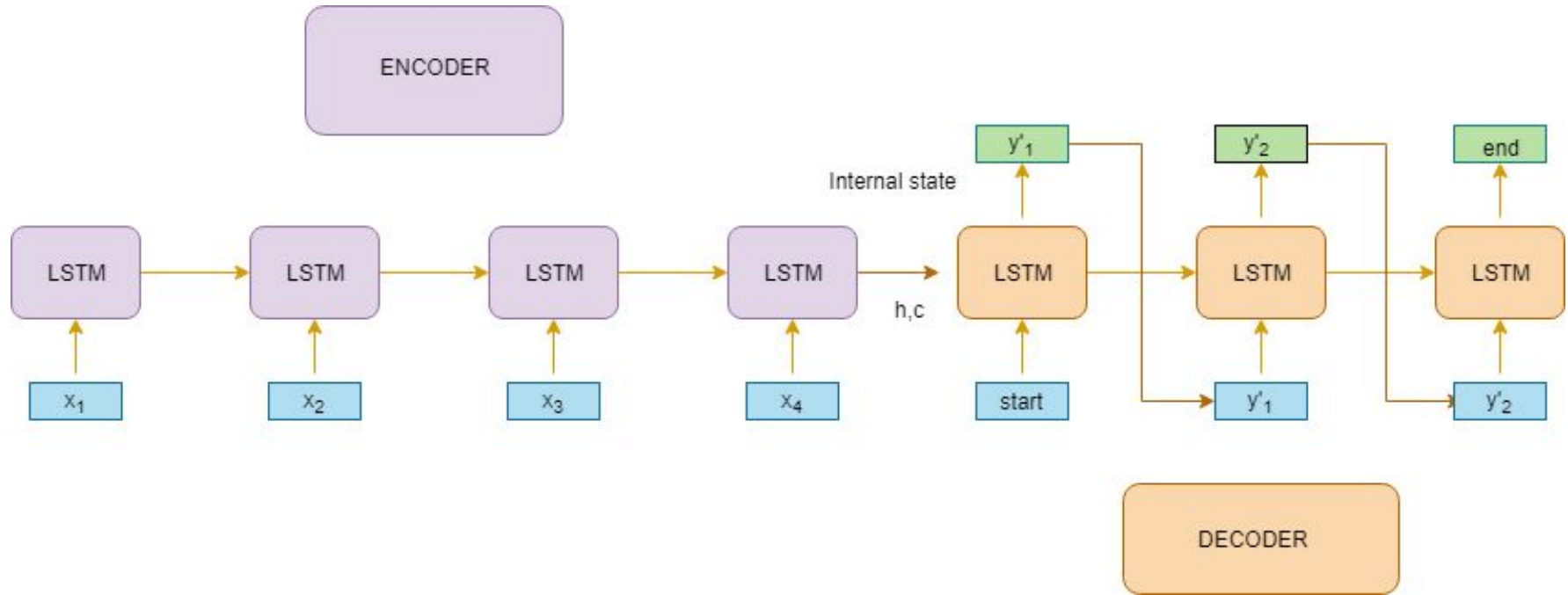


**Attention layer  
implemented in keras  
package**

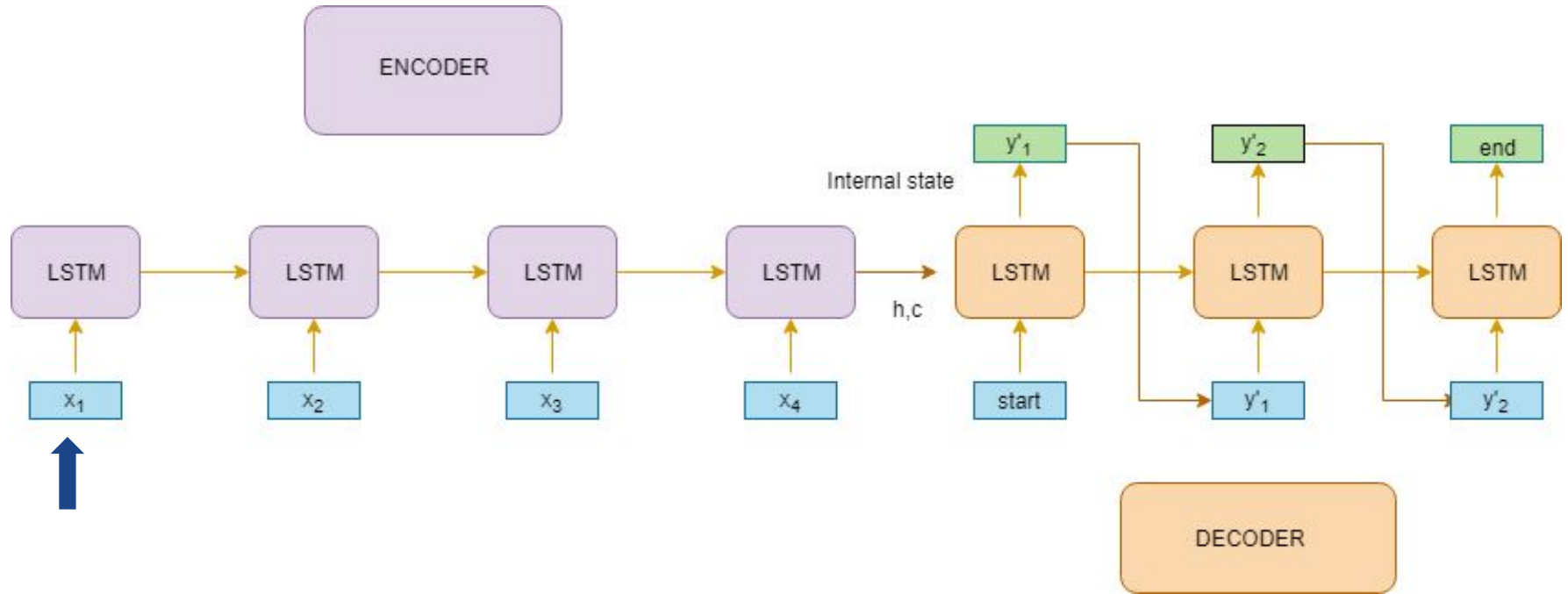
**Encoder-Decoder**



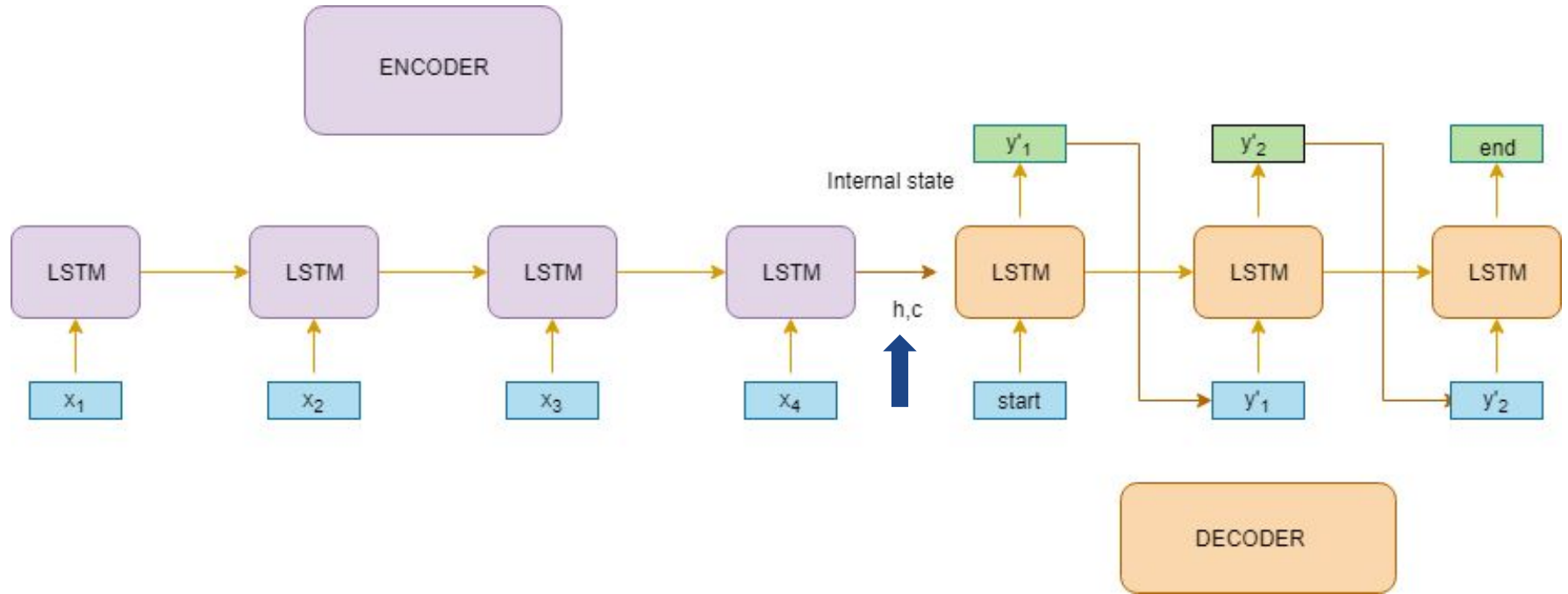
**Stacked LSTM layers**



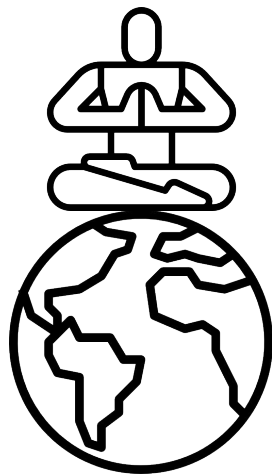
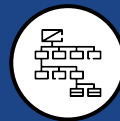
## LSTM-Architecture of our model built from scratch [1]



LSTM-Architecture of our model built from scratch [1]



LSTM-Architecture of our model built from scratch [1]



**Attention Layer**

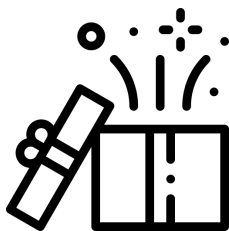


**Attention layer  
implemented in keras  
package**

**Encoder-Decoder**



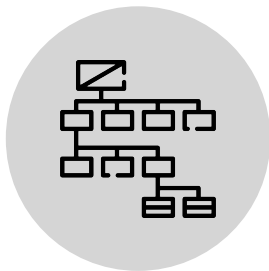
**Stacked LSTM layers**



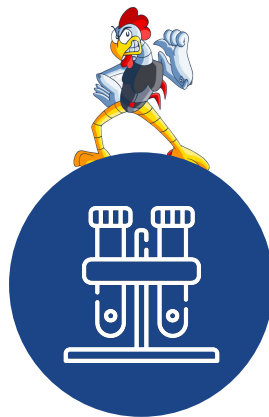
**380 020 868 trainable parameters.**



Dataset



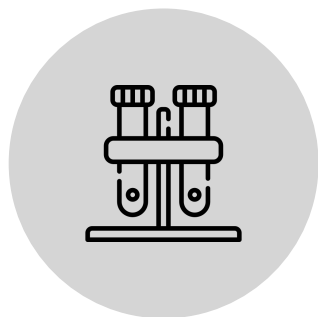
Architecture



Evaluation



Benchmark



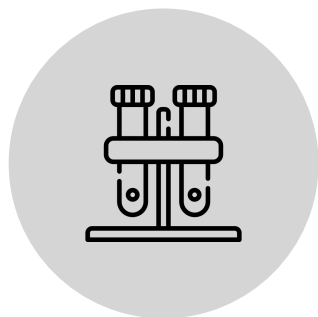
**few number  
of epochs**



**small training  
sample size**



**model does not scale for  
larger amount of text**



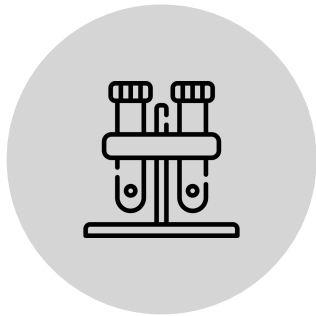
**few number  
of epochs**



**small training  
sample size**



**model does not scale for  
larger amount of text**



**few number  
of epochs**



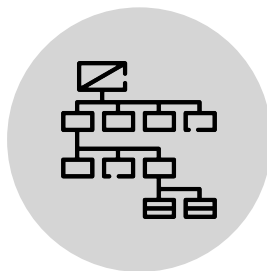
**small training  
sample size**



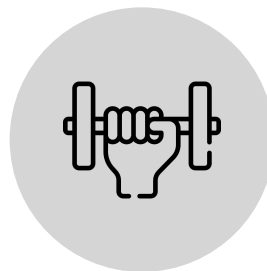
**model does not scale for  
larger amount of text**



Dataset



Architecture



Training



Benchmark



Speech given by Hermann Färber (CDU/CSU) on 11<sup>th</sup> of November 2019

Bundestag is discussing an agricultural bill regarding financial support for sheep and goat farmers



**Hermann Färber (CDU/CSU)**



Speech given by Hermann Färber (CDU/CSU) on 11<sup>th</sup> of November 2019

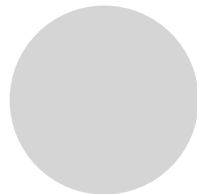
Bundestag is discussing an agricultural bill regarding financial support for sheep and goat farmers



**Hermann Färber (CDU/CSU)**



*“vereins vereins vereins vereins verbandsfreie verbands-freie verbandsfreie ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer ratmeyer rat-meyer ratmeyer ratmeyer ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer ratmeyerratmeyer ratmeyer ratmeyer ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer rat-meyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer”*

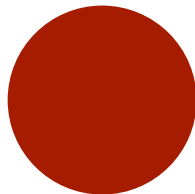


Evaluation





*“vereins vereins vereins vereins verbandsfreie verbands-freie verbandsfreie ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer ratmeyer rat-meyer ratmeyer ratmeyer ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer ratmeyerratmeyer ratmeyer ratmeyer ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer rat-meyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer  
ratmeyer ratmeyer ratmeyer”*



Evaluation

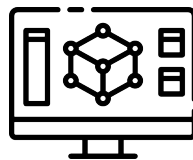




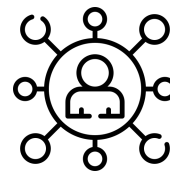
Data



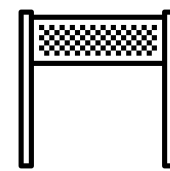
Methodology



Models



Use cases



Conclusion

Model from scratch

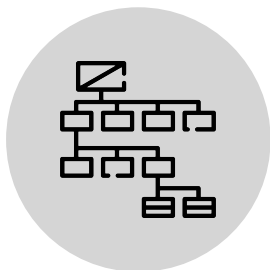
MeanSum

BERT & BART

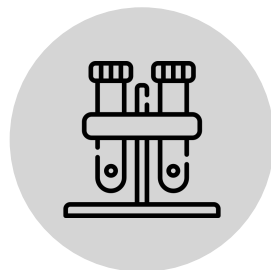
Overall comparison



Dataset



Architecture



Evaluation



Benchmark



## MeanSum : A Neural Model for Unsupervised Multi-Document Abstractive Summarization

Eric Chu<sup>\*†1</sup> Peter J. Liu<sup>\*‡2</sup>

### Abstract

Abstractive summarization has been studied using neural sequence transduction methods with datasets of large, paired document-summary examples. However, such datasets are rare and the models trained from them do not generalize to other domains. Recently, some progress has been made in learning sequence-to-sequence mappings with only unpaired examples. In our work, we consider the setting where there are only documents (product or business reviews) with no summaries provided, and propose an end-to-end, neural model architecture to perform unsupervised abstractive summarization. Our proposed model consists of an auto-encoder where the mean of the representations of the input reviews decodes to a reasonable summary-review while not relying on any review-specific features. We consider variants of the proposed architecture and perform an ablation study to show the importance of specific components. We show through automated metrics and human evaluation that the generated summaries are highly abstractive, fluent, relevant, and representative of the average sentiment of the input reviews. Finally, we collect a reference evaluation dataset and show that our model outperforms a strong extractive baseline.

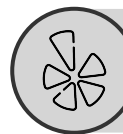
or recognition of short utterances, for which there is an abundance of parallel data. The application of such models to longer sequences (multi-sentence documents or long audio) works reasonably well in production systems because the sequences can be naturally decomposed into the shorter ones the models are trained on and thus sequence-transduction can be done piece-meal.

Similar neural models have also been applied to abstractive summarization, where large numbers of document-summary pairs are used to generate news headlines (Rush et al., 2015) or bullet-points (Nallapati et al., 2016; See et al., 2017). Work in this vein has been extended by Liu et al. (2018) to the multi-document<sup>1</sup> case to produce Wikipedia article text from references documents.

However, unlike translation or speech recognition, adapting such summarization models to different types of documents without re-training is much less reasonable; for example, in general documents do not decompose into parts that look like news articles, nor can we expect our idea of saliency or desired writing style to correspond with that of particular news publishers. Re-training or at least fine-tuning such models on many in-domain document-summary pairs should be expected to get desirable performance. Unfortunately, it is very expensive to create a large parallel summarization corpus and the most common case in our experience is that we have many documents to summarize, but have few or no examples of summaries.



**Multi-document summarization model**



**Summarizes multiple Yelp reviews of the same business**

[3] Eric Chu and Peter Liu (2019)



## MeanSum : A Neural Model for Unsupervised Multi-Document Abstractive Summarization

Eric Chu <sup>\*†1</sup> Peter J. Liu <sup>\*‡2</sup>

### Abstract

Abstractive summarization has been studied using neural sequence transduction methods with datasets of large, paired document-summary examples. However, such datasets are rare and the models trained from them do not generalize to other domains. Recently, some progress has been made in learning sequence-to-sequence mappings with only unpaired examples. In our work, we consider the setting where there are only documents (product or business reviews) with no summaries provided, and propose an end-to-end, neural model architecture to perform unsupervised abstractive summarization. Our proposed model consists of an auto-encoder where the mean of the representations of the input reviews decodes to a reasonable summary-review while not relying on any review-specific features. We consider variants of the proposed architecture and perform an ablation study to show the importance of specific components. We show through automated metrics and human evaluation that the generated summaries are highly abstractive, fluent, relevant, and representative of the average sentiment of the input reviews. Finally, we collect a reference evaluation dataset and show that our model outperforms a strong extractive baseline.

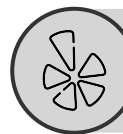
or recognition of short utterances, for which there is an abundance of parallel data. The application of such models to longer sequences (multi-sentence documents or long audio) works reasonably well in production systems because the sequences can be naturally decomposed into the shorter ones the models are trained on and thus sequence-transduction can be done piece-meal.

Similar neural models have also been applied to abstractive summarization, where large numbers of document-summary pairs are used to generate news headlines (Rush et al., 2015) or bullet-points (Nallapati et al., 2016; See et al., 2017). Work in this vein has been extended by Liu et al. (2018) to the multi-document<sup>1</sup> case to produce Wikipedia article text from references documents.

However, unlike translation or speech recognition, adapting such summarization models to different types of documents without re-training is much less reasonable; for example, in general documents do not decompose into parts that look like news articles, nor can we expect our idea of saliency or desired writing style to correspond with that of particular news publishers. Re-training or at least fine-tuning such models on many in-domain document-summary pairs should be expected to get desirable performance. Unfortunately, it is very expensive to create a large parallel summarization corpus and the most common case in our experience is that we have many documents to summarize, but have few or no examples of summaries.



**Multi-document summarization model**



**Summarizes multiple Yelp reviews of the same business**

[3] Eric Chu and Peter Liu (2019)



## MeanSum : A Neural Model for Unsupervised Multi-Document Abstractive Summarization

Eric Chu <sup>\*†1</sup> Peter J. Liu <sup>\*‡2</sup>

### Abstract

Abstractive summarization has been studied using neural sequence transduction methods with datasets of large, paired document-summary examples. However, such datasets are rare and the models trained from them do not generalize to other domains. Recently, some progress has been made in learning sequence-to-sequence mappings with only unpaired examples. In our work, we consider the setting where there are only documents (product or business reviews) with no summaries provided, and propose an end-to-end, neural model architecture to perform unsupervised abstractive summarization. Our proposed model consists of an auto-encoder where the mean of the representations of the input reviews decodes to a reasonable summary-review while not relying on any review-specific features. We consider variants of the proposed architecture and perform an ablation study to show the importance of specific components. We show through automated metrics and human evaluation that the generated summaries are highly abstractive, fluent, relevant, and representative of the average sentiment of the input reviews. Finally, we collect a reference evaluation dataset and show that our model outperforms a strong extractive baseline.

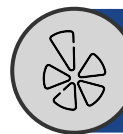
or recognition of short utterances, for which there is an abundance of parallel data. The application of such models to longer sequences (multi-sentence documents or long audio) works reasonably well in production systems because the sequences can be naturally decomposed into the shorter ones the models are trained on and thus sequence-transduction can be done piece-meal.

Similar neural models have also been applied to abstractive summarization, where large numbers of document-summary pairs are used to generate news headlines (Rush et al., 2015) or bullet-points (Nallapati et al., 2016; See et al., 2017). Work in this vein has been extended by Liu et al. (2018) to the multi-document<sup>1</sup> case to produce Wikipedia article text from references documents.

However, unlike translation or speech recognition, adapting such summarization models to different types of documents without re-training is much less reasonable; for example, in general documents do not decompose into parts that look like news articles, nor can we expect our idea of saliency or desired writing style to correspond with that of particular news publishers. Re-training or at least fine-tuning such models on many in-domain document-summary pairs should be expected to get desirable performance. Unfortunately, it is very expensive to create a large parallel summarization corpus and the most common case in our experience is that we have many documents to summarize, but have few or no examples of summaries.

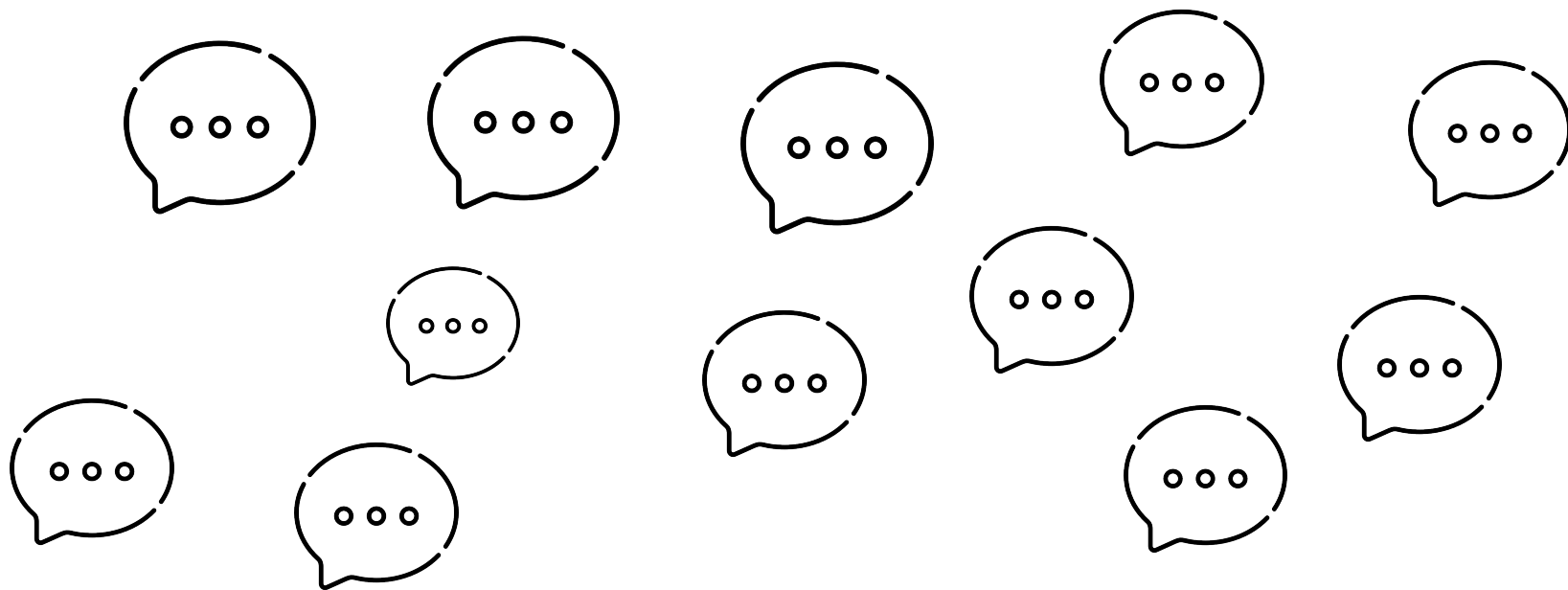


**Multi-document summarization model**



**Summarizes multiple Yelp reviews of the same business**

[3] Eric Chu and Peter Liu (2019)



Cluster PP19-speeches by agenda point

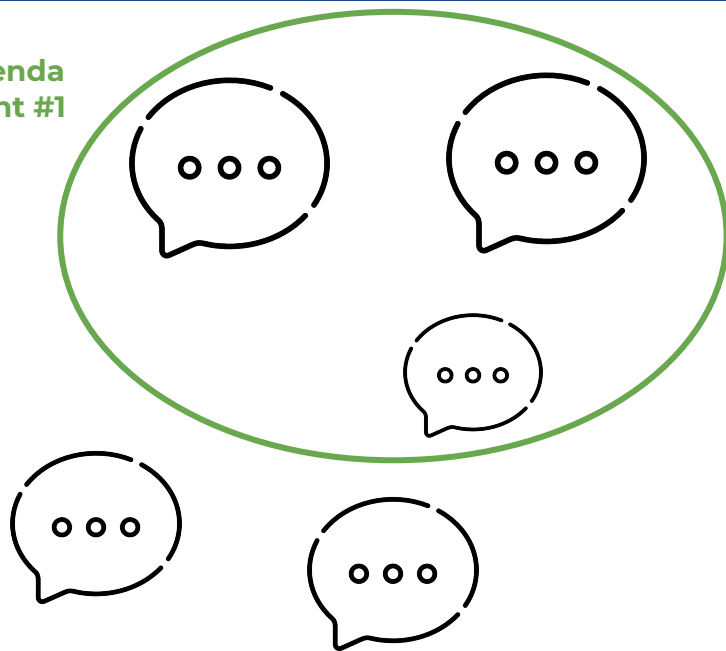


1289 agenda points

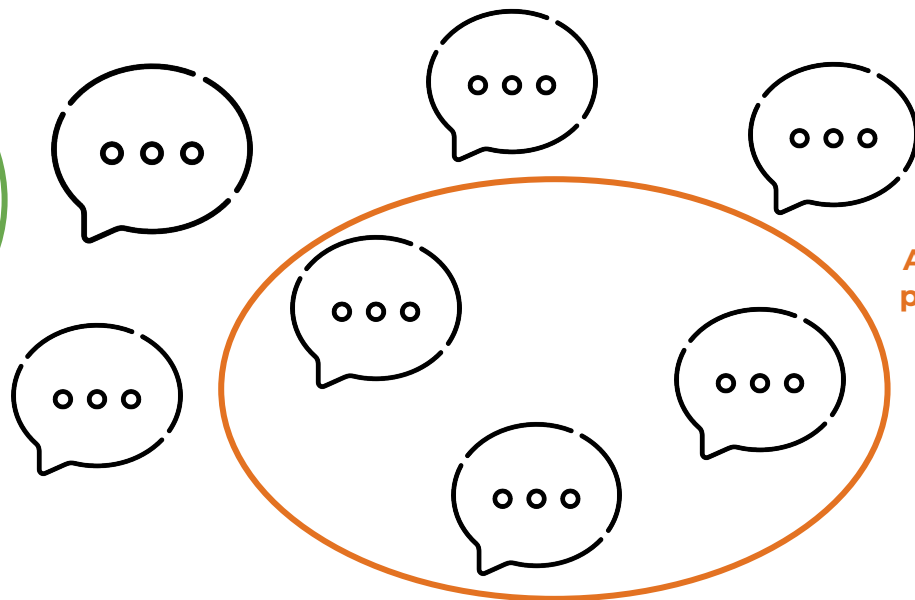
22 speeches on average



Agenda  
point #1



Agenda  
point #2



Cluster PP19-speeches by agenda point

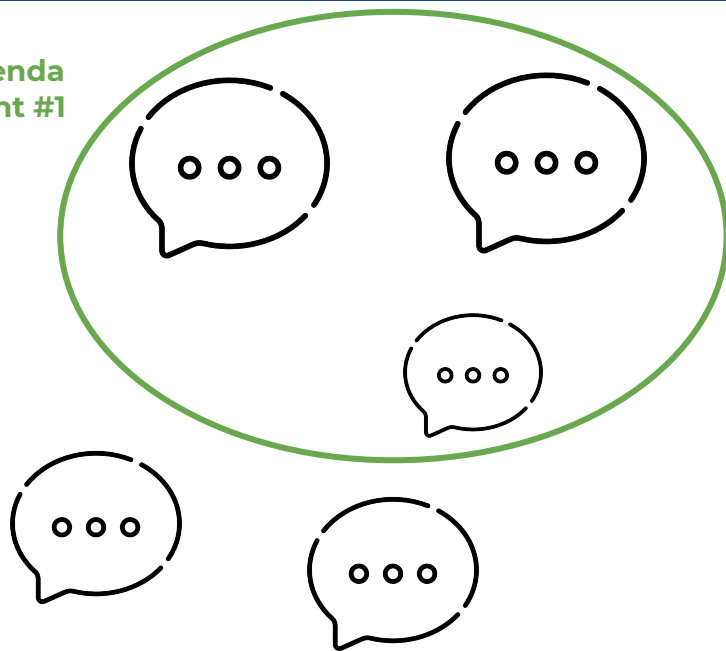


1289 agenda points

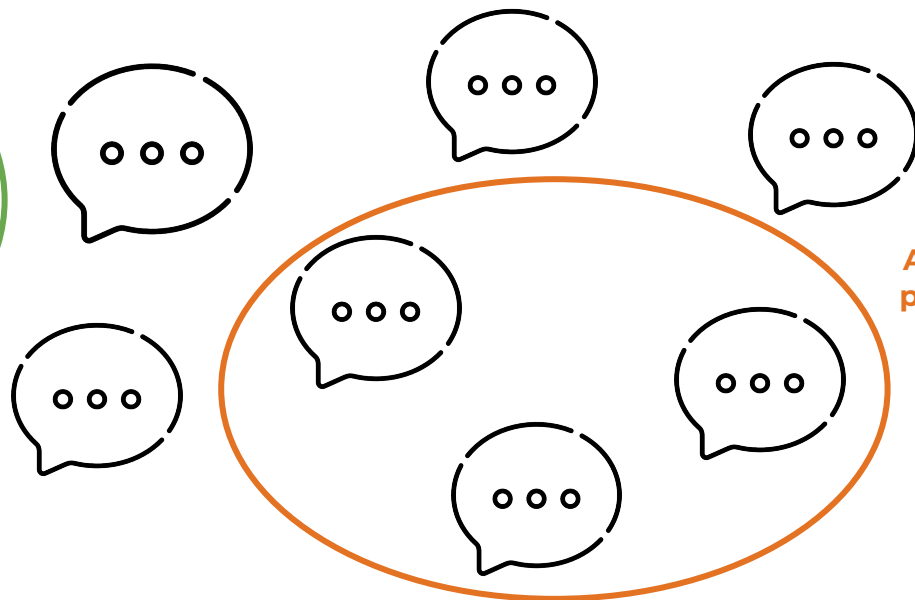
22 speeches on average



Agenda  
point #1



Agenda  
point #2

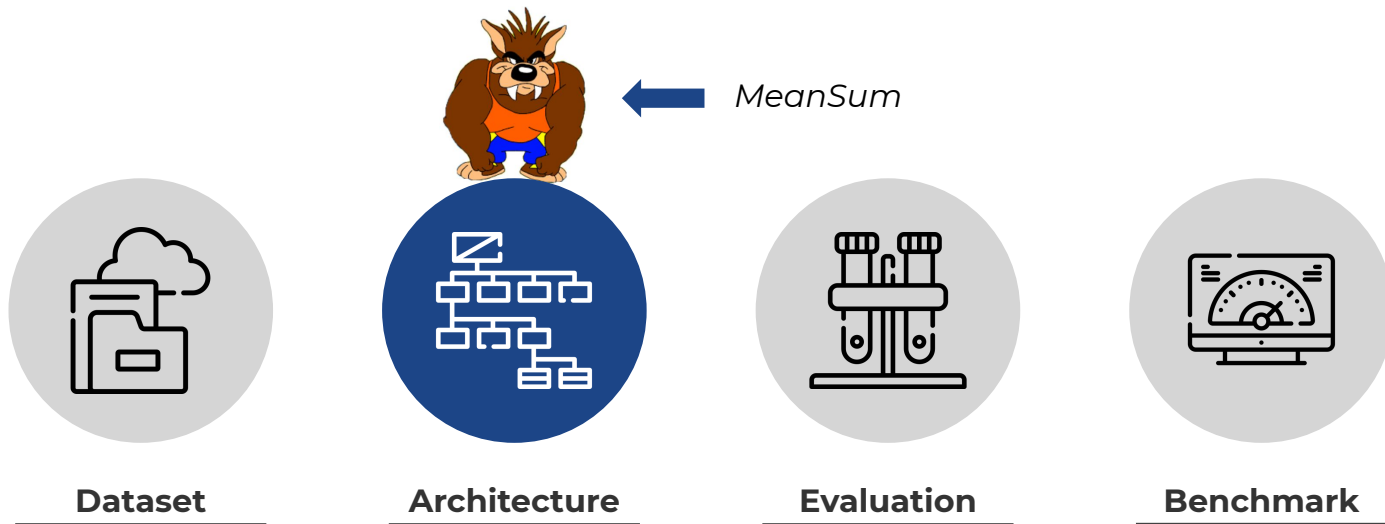


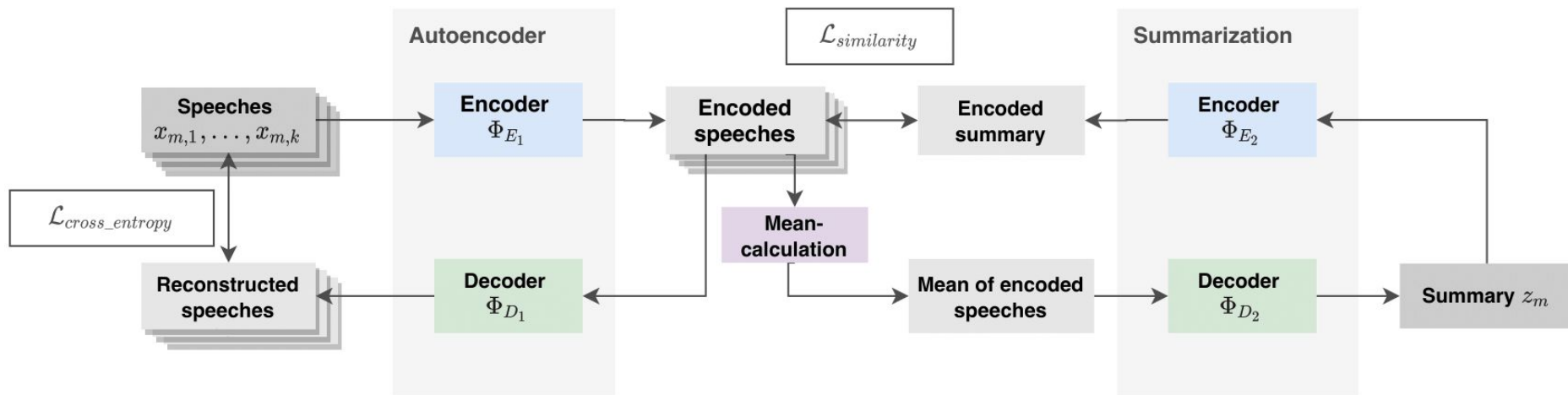
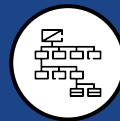
Cluster PP19-speeches by agenda point

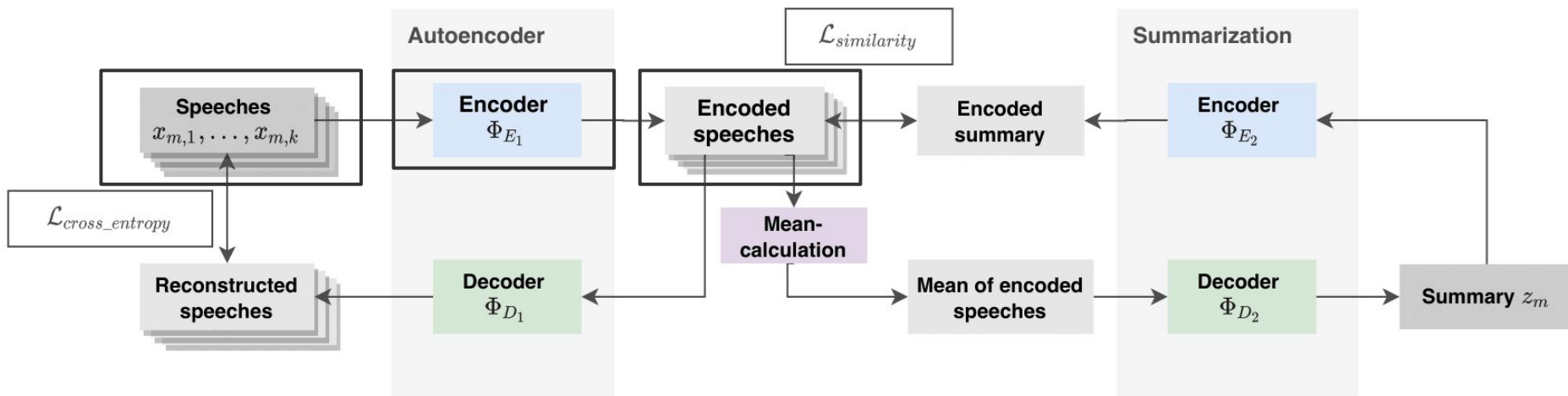
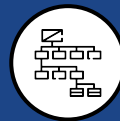


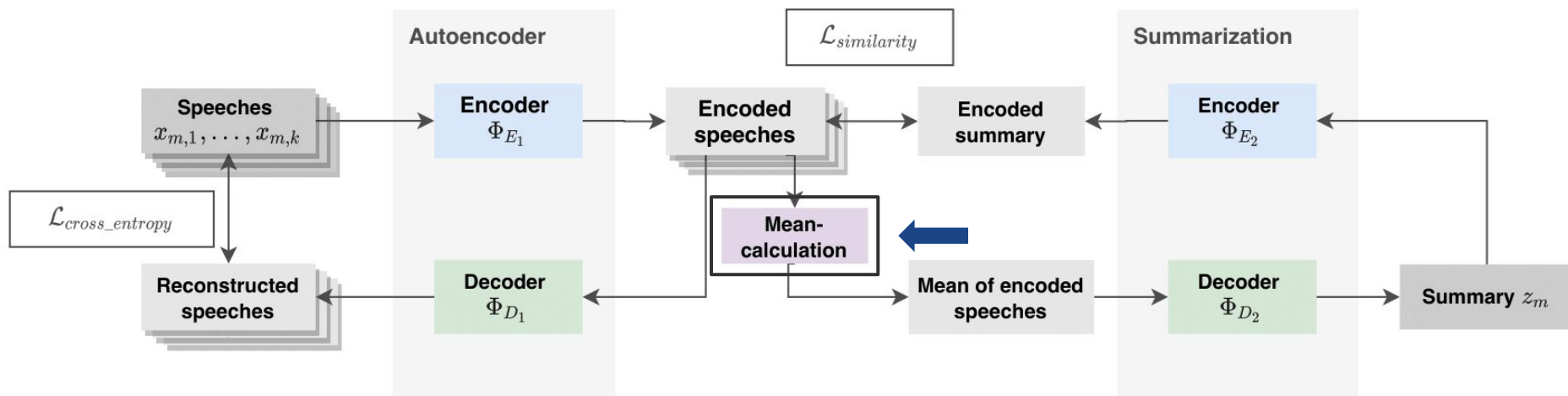
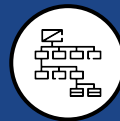
1289 agenda points

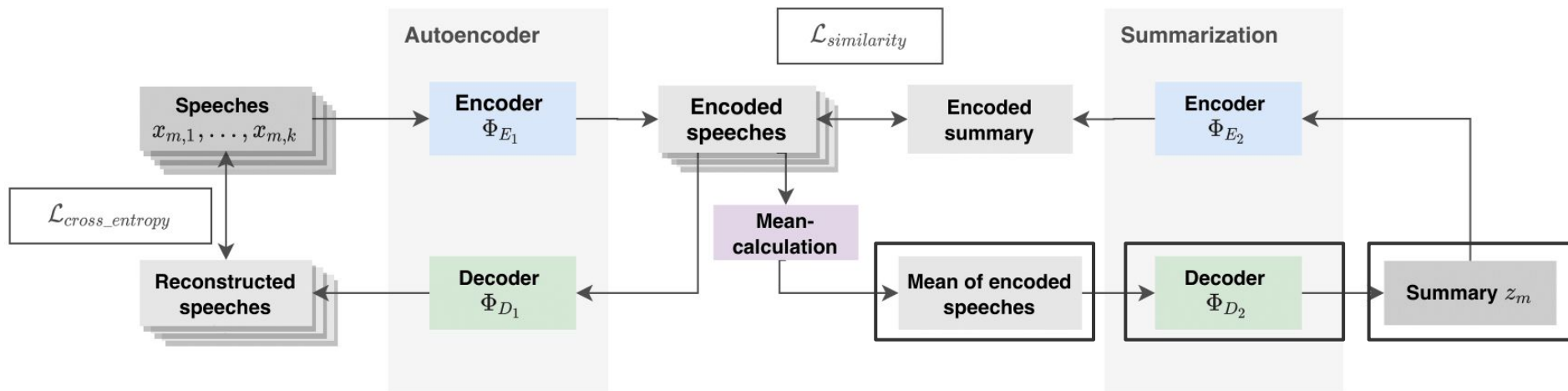
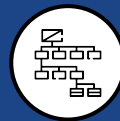
22 speeches on average





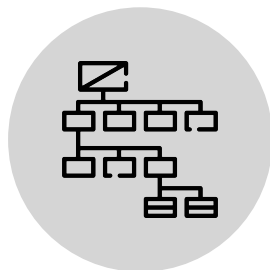




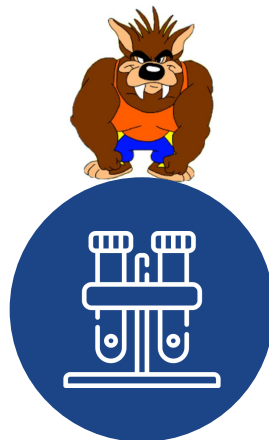




Dataset



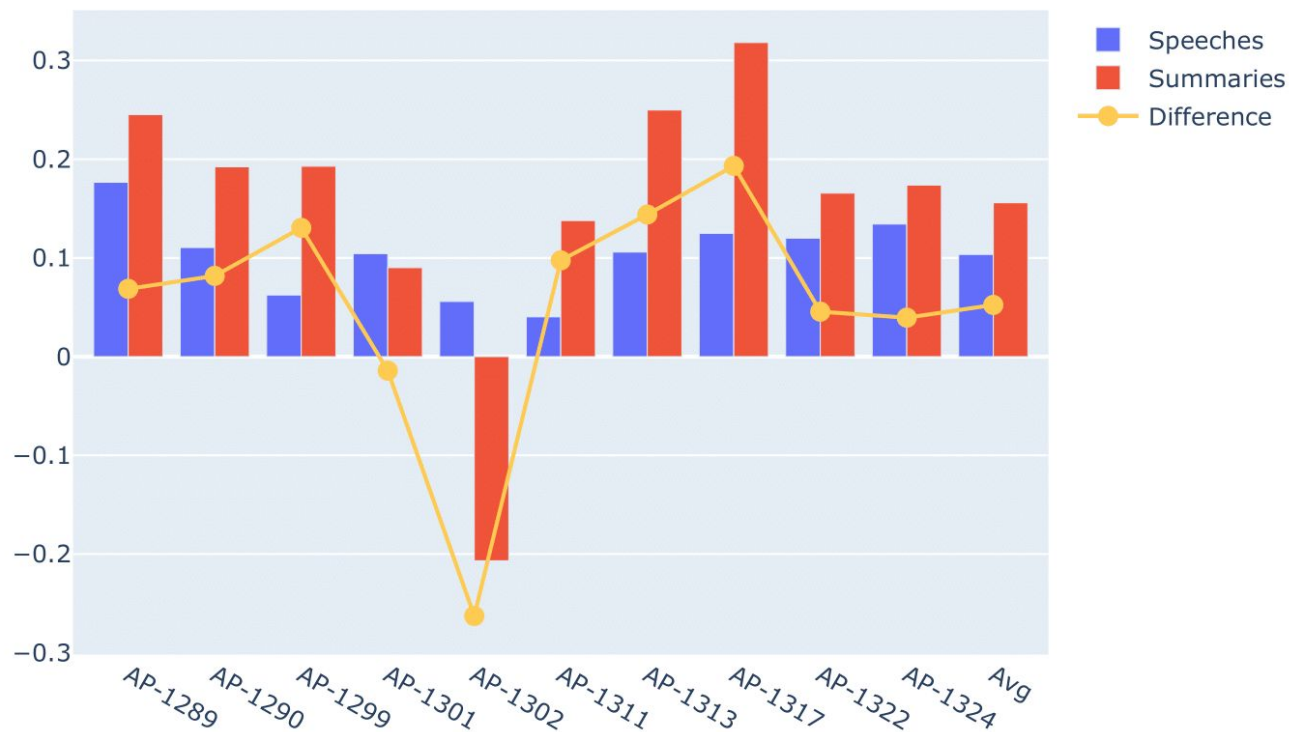
Architecture



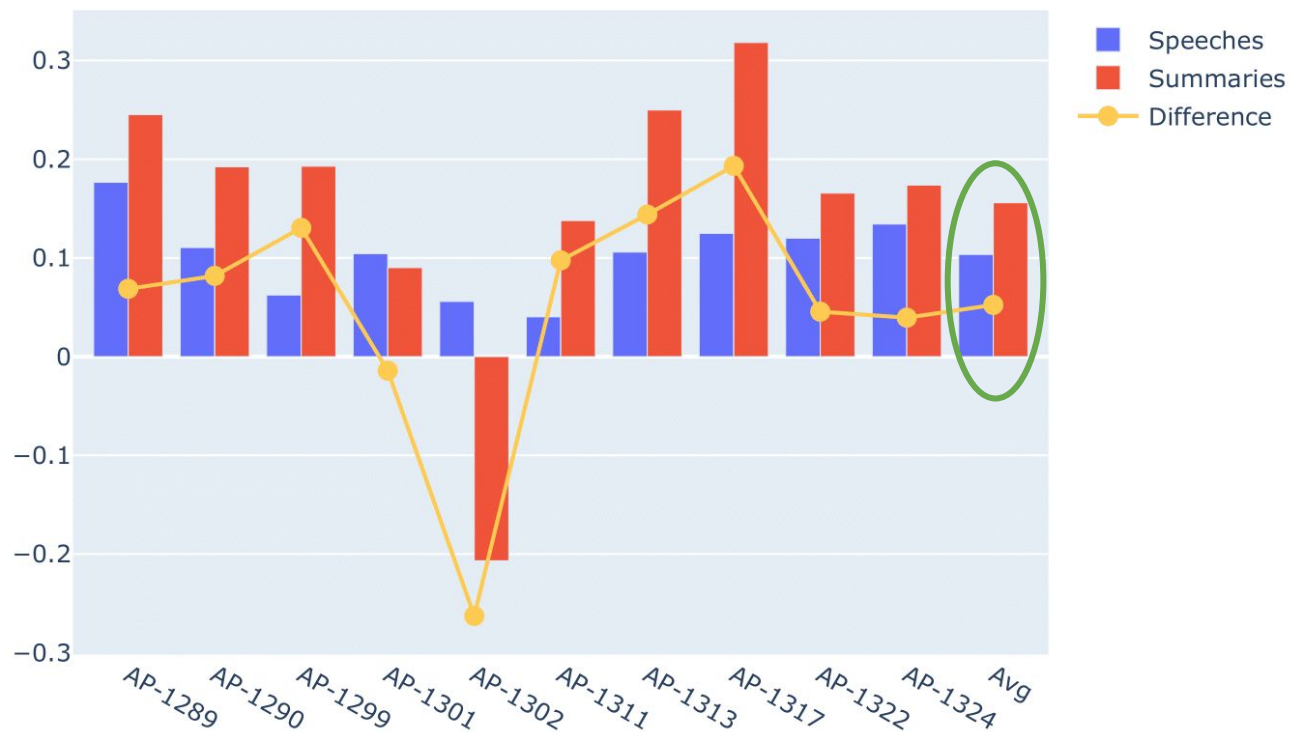
Evaluation



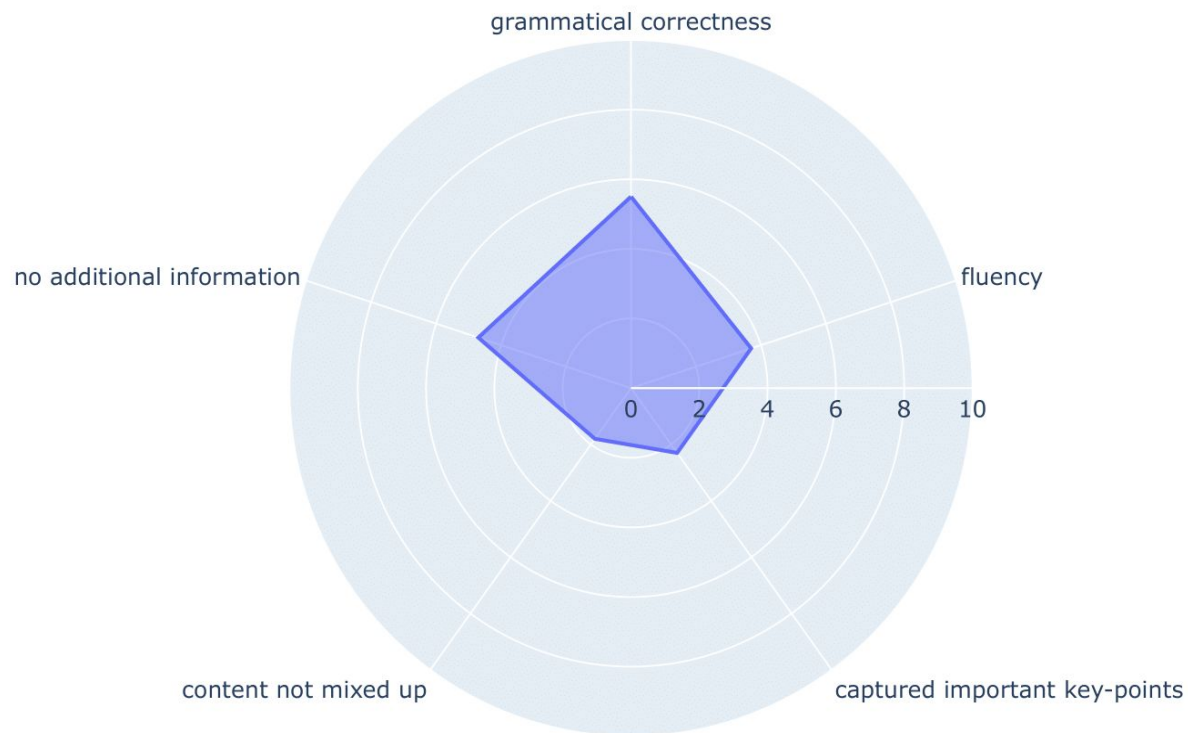
Benchmark



**Sentiment analysis**



Sentiment analysis

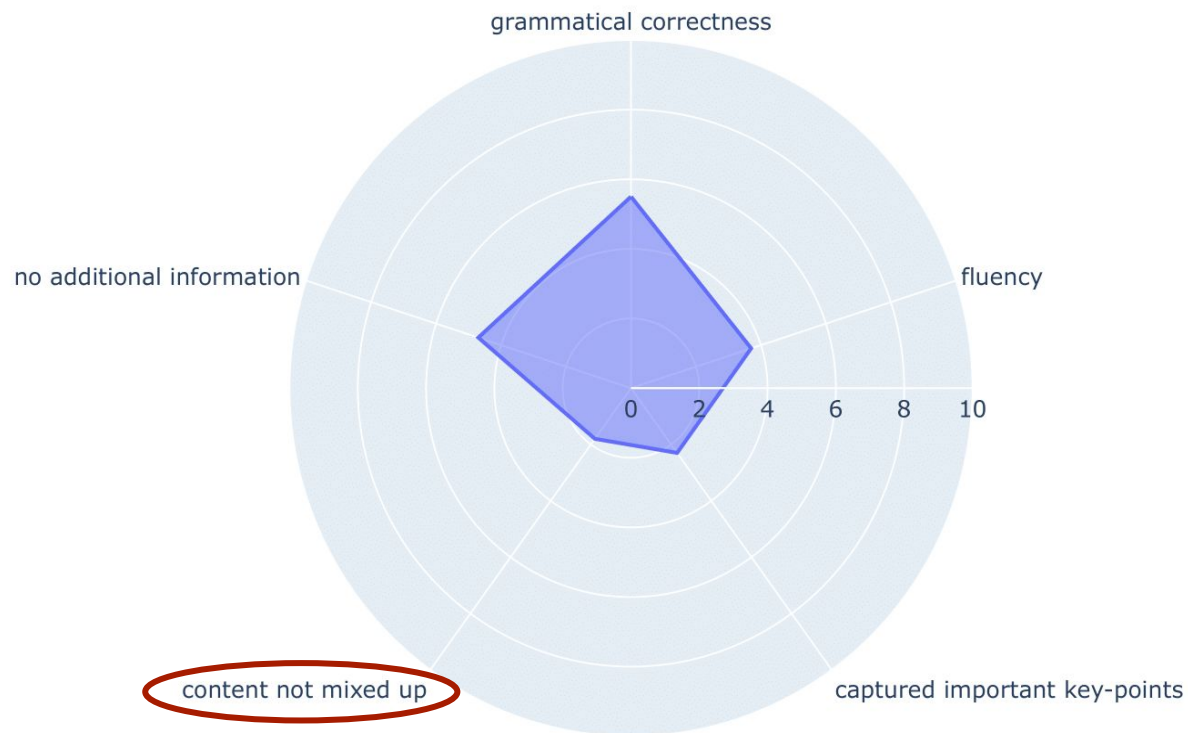


## Questionnaire results

**Model not designed for long input-sequences**



**Attention-mechanism likely to improve model**

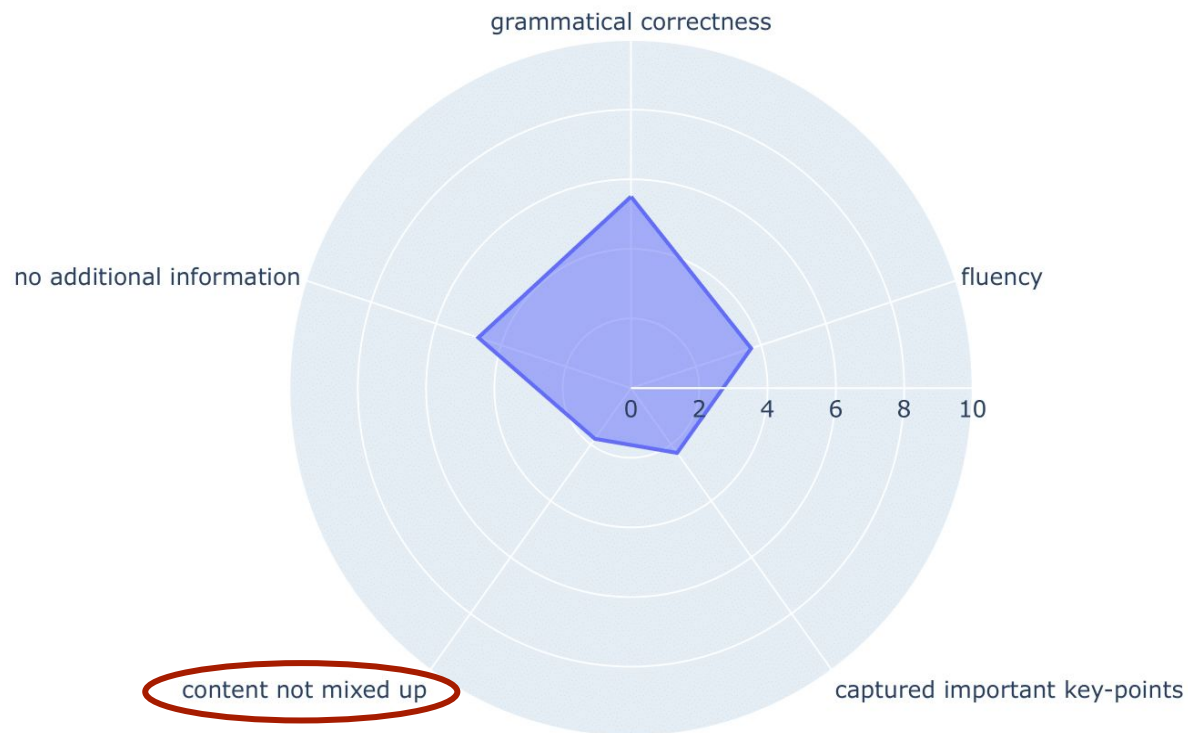


## Questionnaire results

**Model not designed for long input-sequences**



**Attention-mechanism likely to improve model**

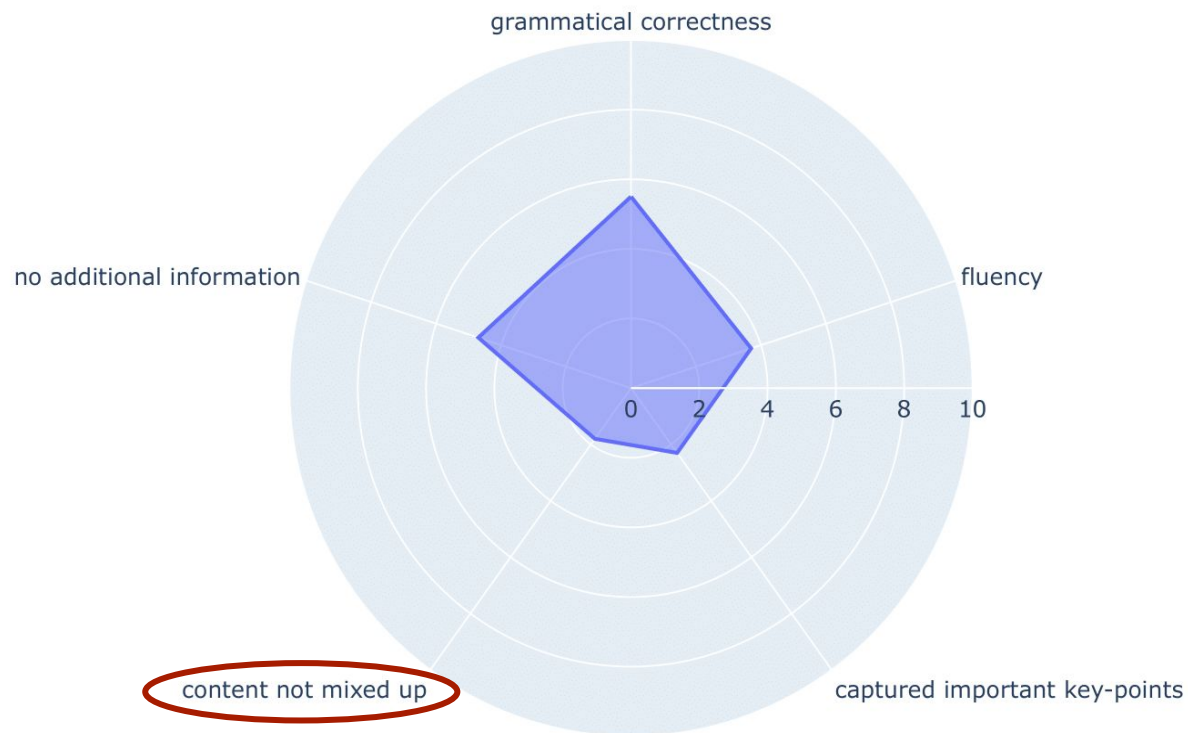


Questionnaire results

**Model not designed for long input-sequences**



**Attention-mechanism likely to improve model**



**Questionnaire results**

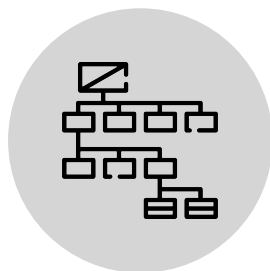
**Model not designed for long input-sequences**



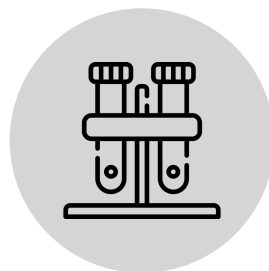
**Attention-mechanism likely to improve model**



Dataset



Architecture



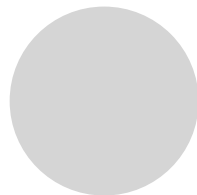
Evaluation



Benchmark



*“Hallo, liebe Kollegen! Liebe Gäste! Es geht darum, die Landwirtschaft zu fördern, auch wenn man sich das Ganze in einem Biergarten ansieht, der jetzt genutzt wird. Dies ist das Ergebnis der Bundesregierung und 2020. Meine Damen und Herren, meine Damen und Herren, lassen Sie uns auf die Erde gehen. Wir sind mitten im Budget. Dies ist ein großer Schritt nach vorne. Das kann man nicht einfach essen. Sie können auch darüber sprechen. Sie fordern das Steuerprivileg des EEG-Zuschlags, zum Beispiel für ein Video von 418 Millionen, das dann natürlich von der Regierung finanziert wird. Dies ist ein sehr wichtiger Schritt in Richtung Weltklima- und Umweltpolitik.”*

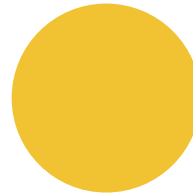


Evaluation





*“Hallo, liebe Kollegen! Liebe Gäste! Es geht darum, die Landwirtschaft zu fördern, auch wenn man sich das Ganze in einem Biergarten ansieht, der jetzt genutzt wird. Dies ist das Ergebnis der Bundesregierung und 2020. Meine Damen und Herren, meine Damen und Herren, lassen Sie uns auf die Erde gehen. Wir sind mitten im Budget. Dies ist ein großer Schritt nach vorne. Das kann man nicht einfach essen. Sie können auch darüber sprechen. Sie fordern das Steuerprivileg des EEG-Zuschlags, zum Beispiel für ein Video von 418 Millionen, das dann natürlich von der Regierung finanziert wird. Dies ist ein sehr wichtiger Schritt in Richtung Weltklima- und Umweltpolitik.”*



Evaluation

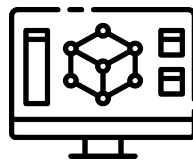




Data



Methodology



Models



Use cases



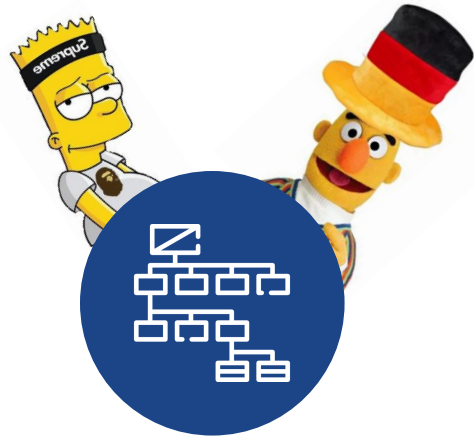
Conclusion

Model from scratch

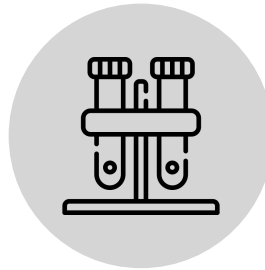
MeanSum

BERT & BART

Overall comparison



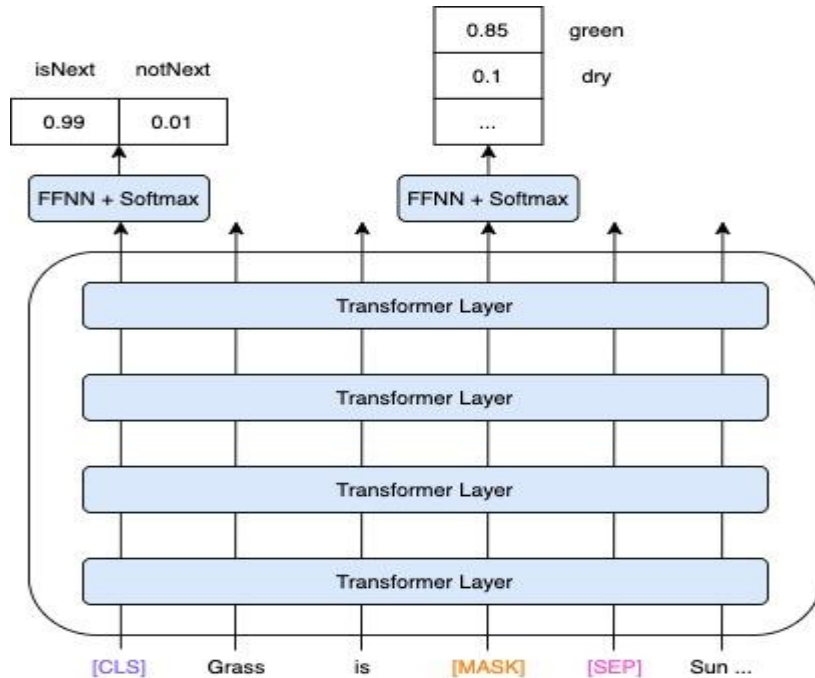
Architecture



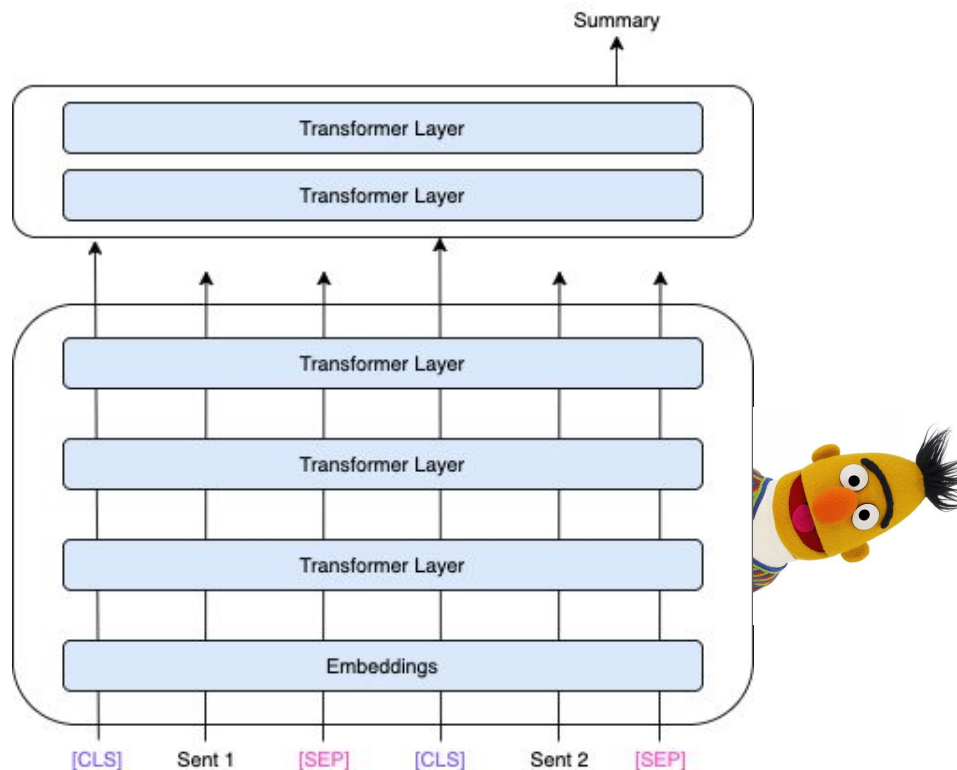
Evaluation



Dataset



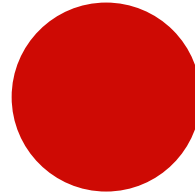
**Bidirectional Encoder Representations from Transformers [2]**



BERTSumAbs architecture [4]



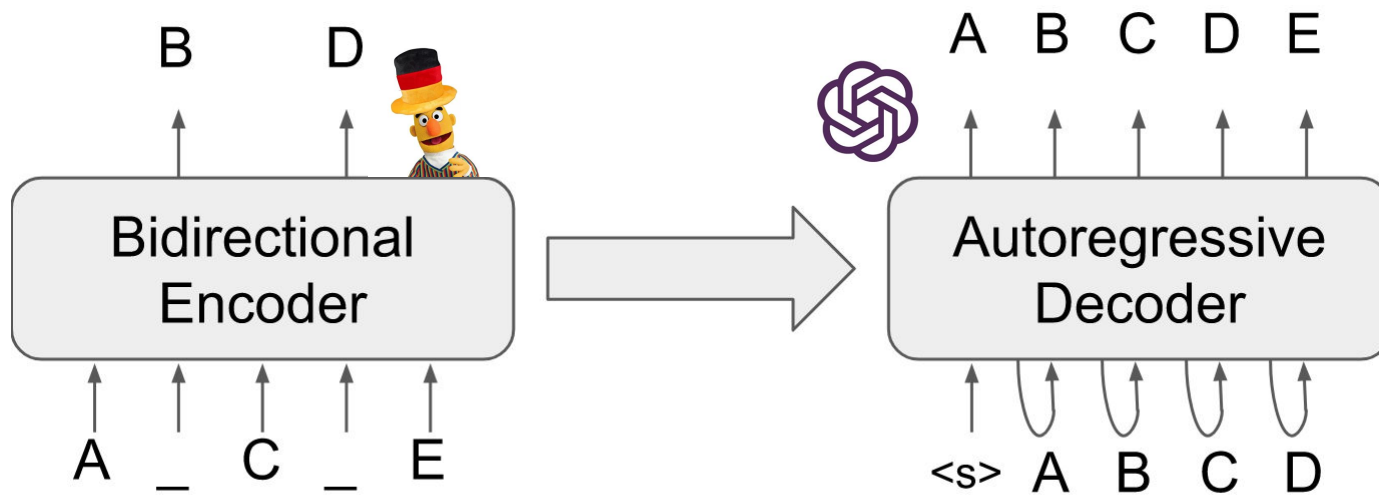
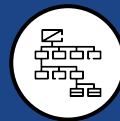
*“Die Simulation ist ein Verfahren zurückragung, bei der für die für die für die führte. Das Verfahren ist in der Regel für das Verfahrenzurückragung der Regel führt der Regel führte das Verfahren zurückragungen von der Regel führte - und das Verfahren zurührt - das Verfahren zur Verfügung der Regel führt der Regel führt - Änderung der Regel führte das Verfahren der Regelführte.”*



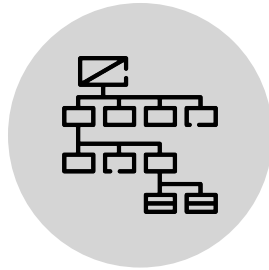
Evaluation



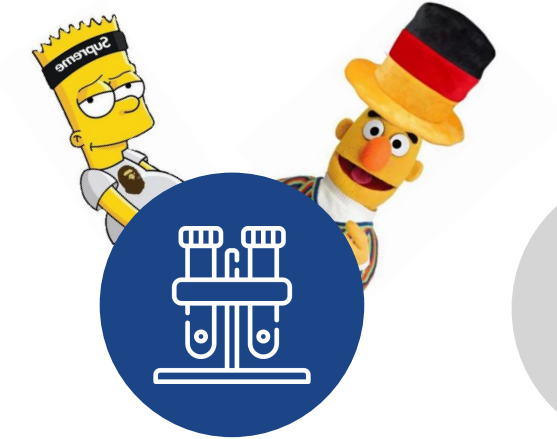




**Bidirectional AutoRegressive Transformer [5]**



Architecture



Evaluation



Dataset



| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision  |
|------------------|----------------------|----------------------|----------------------|----------------------|------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454     |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317     |

*SwissText-Dataset*

# Comparison BART and BERT

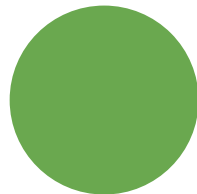


| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision  |
|------------------|----------------------|----------------------|----------------------|----------------------|------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454     |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317     |

*SwissText-Dataset*



*“Die AfD debattiert über das Gesetz zur Umsetzung von Direktzahlungen und die einmalige Erhöhung der Umschichtung um 1,5 Prozentpunkte. Dies entspricht einer Steigerung von 4,50 EUR pro Hektar von der ersten zur zweiten Säule der Agrarumweltprogramme. Die Gruppen Die Linke und Bündnis 90 / DieGrünen erkennen in ihrer Bewerbung zu Recht, dass die Schafzüchter mit ihrer Arbeit uns qualitativ hochwertige Produkte in der Bevölkerung liefern.”*



Evaluation



# Comparison BART and BERT



| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision  |
|------------------|----------------------|----------------------|----------------------|----------------------|------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274 ✓   |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454     |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317 ✓   |

*SwissText-Dataset*

**BART is highly extractive**

# Comparison BART and BERT



| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision  |
|------------------|----------------------|----------------------|----------------------|----------------------|------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274 ✓   |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454     |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317 ✓   |

*SwissText-Dataset*

**BART is highly extractive**

# Comparison BART and BERT



| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision    |
|------------------|----------------------|----------------------|----------------------|----------------------|--------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274 ✓     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454       |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> ✓ |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317 ✓     |

*SwissText-Dataset*

**BART is highly extractive**

# Comparison BART and BERT



| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision    |
|------------------|----------------------|----------------------|----------------------|----------------------|--------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274 ✓     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454       |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> ✓ |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317 ✓     |

SwissText-Dataset

BART is highly extractive



# Comparison BART and BERT



| Model            | R1-score               | R2-score               | RL-score               | BERT-score           | precision    |
|------------------|------------------------|------------------------|------------------------|----------------------|--------------|
| BART             | 0.2024 (0.07)          | 0.0456 (0.04)          | 0.1955 (0.07)          | <b>0.6251 (0.05)</b> | 0.7274 ✓     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> ✓ | <b>0.0801 (0.10)</b> ✓ | <b>0.2479 (0.13)</b> ✓ | 0.6224 (0.08)        | 0.4454       |
| lead-sum         | 0.1882 (0.08)          | 0.0412 (0.05)          | 0.1777 (0.07)          | 0.6112 (0.05)        | <b>1.0</b> ✓ |
| english-lead-sum | 0.1782 (0.07)          | 0.0341 (0.04)          | 0.1682 (0.06)          | 0.6049 (0.04)        | 0.7317 ✓     |

*SwissText-Dataset*

**BART is highly extractive**

**BERT uses patterns**

# Comparison BART and BERT



| Model            | R1-score               | R2-score               | RL-score               | BERT-score           | precision    |
|------------------|------------------------|------------------------|------------------------|----------------------|--------------|
| BART             | 0.2024 (0.07)          | 0.0456 (0.04)          | 0.1955 (0.07)          | <b>0.6251 (0.05)</b> | 0.7274 ✓     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> ✓ | <b>0.0801 (0.10)</b> ✓ | <b>0.2479 (0.13)</b> ✓ | 0.6224 (0.08)        | 0.4454       |
| lead-sum         | 0.1882 (0.08)          | 0.0412 (0.05)          | 0.1777 (0.07)          | 0.6112 (0.05)        | <b>1.0</b> ✓ |
| english-lead-sum | 0.1782 (0.07)          | 0.0341 (0.04)          | 0.1682 (0.06)          | 0.6049 (0.04)        | 0.7317 ✓     |

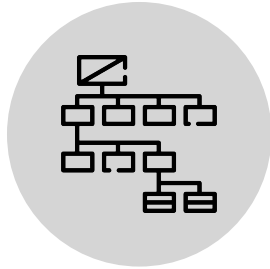
*SwissText-Dataset*

**BART is highly extractive**

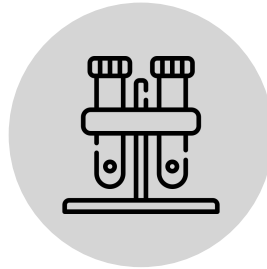
**BERT uses patterns**



**Will BERT generalize?**



Architecture



Evaluation



Dataset



Drucksachen

## Breitband soll Universaldienst werden

Verkehr und digitale Infrastruktur/Antrag - 08.07.2020 (hib 735/2020)

Berlin: (hib/SCR) Die Fraktion Bündnis 90/Die Grünen will einen "Rechtsanspruch auf einen schnellen Breitband-Internetanschluss" festschreiben. Ein solcher Anschluss sei "eine wichtige Voraussetzung für die gleichberechtigte Teilhabe am gesellschaftlichen, wirtschaftlichen und kulturellen Leben", führen die Grünen aus und weisen auf entsprechende Versorgungsprobleme gerade im ländlichen Raum hin. In einem Antrag (  19/20786) fordert die Fraktion die Bundesregierung unter anderem auf, die Bundesnetzagentur zum Handeln zu bewegen. Die Bundesnetzagentur soll nach Vorstellung der Grünen auf Grundlage der Regelungen zum Universaldienst im Telekommunikationsgesetz (§ 78 ff.) "den Bedarf der Breitband-Universaldienstleistung bei den Endnutzerinnen und Endnutzern formal" feststellen, "insbesondere hinsichtlich der geographischen Versorgung". In unterversorgten Gebieten soll die Netzentur die Erbringung des Breitband-Universaldienstes ausschreiben.

Zudem wollen die Grünen die Internetdienstanbieter stärker hinsichtlich der zugesicherten Bandbreite in die Pflicht nehmen. Verbraucher sollen laut Antrag ein Sonderkündigungsrecht beziehungsweise ein Recht auf Tarifierpassung erhalten, wenn die angepriesene und die tatsächlich gemessene Leistung erheblich, kontinuierlich und wiederkehrend voneinander abweichen. Zudem fordern die Grünen erweiterte Sanktionsmöglichkeiten für die Bundesnetzagentur. Sie soll laut Antrag "umsatzbezogene Bußgelder von bis zu vier Prozent des in Deutschland im betreffenden Geschäftsbereich erzielten Jahresumsatzes" des Vorjahres verhängen können.

50 suitable press releases



| Model             | R1-score      | R2-score      | RL-score      | BERT          |
|-------------------|---------------|---------------|---------------|---------------|
| BART              | 0.1275        | 0.0288        | 0.1383        | <b>0.5994</b> |
| distil-BERTSumAbs | 0.0983        | 0.0053        | 0.0784        | 0.4867        |
| rand-sum          | <b>0.2300</b> | <b>0.0338</b> | <b>0.1833</b> | 0.5863        |
| lead-sum          | 0.0123        | 0             | 0.174         | 0.4713        |

*Heute im Bundestag-Dataset*

**BART is highly extractive**

**BERT uses patterns**



**Will BERT generalize?**

# Comparison BART and BERT



| Model             | R1-score      | R2-score      | RL-score      | BERT            |
|-------------------|---------------|---------------|---------------|-----------------|
| BART              | 0.1275        | 0.0288        | 0.1383        | <b>0.5994</b> ✓ |
| distil-BERTSumAbs | 0.0983        | 0.0053        | 0.0784        | 0.4867 ✓        |
| rand-sum          | <b>0.2300</b> | <b>0.0338</b> | <b>0.1833</b> | 0.5863          |
| lead-sum          | 0.0123        | 0             | 0.174         | 0.4713          |

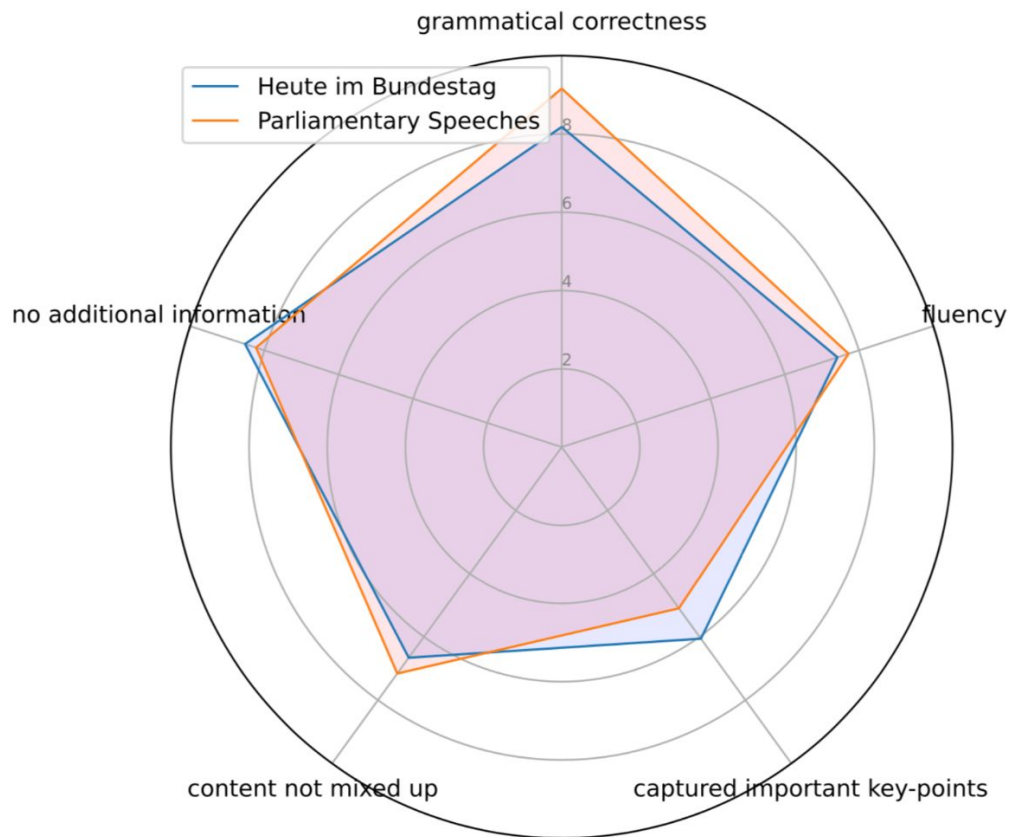
*Heute im Bundestag-Dataset*

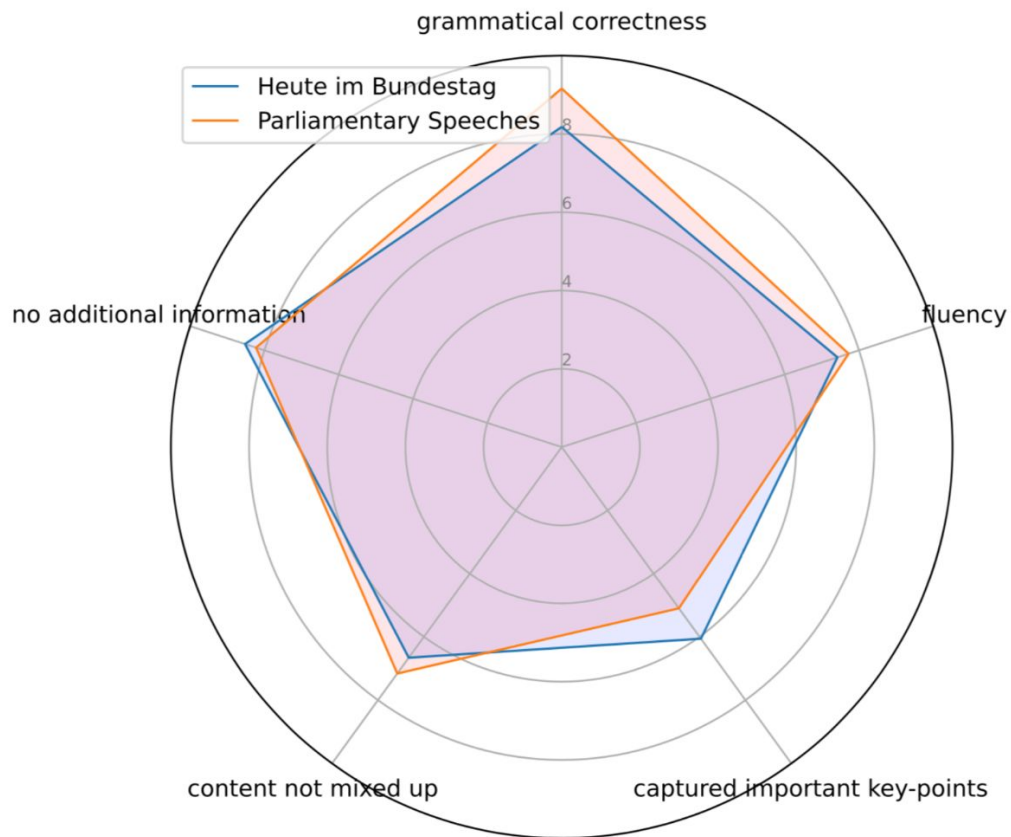
**BART is highly extractive**

**BERT uses patterns**

**BART generalizes**

**BERT does not**



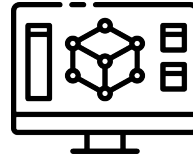




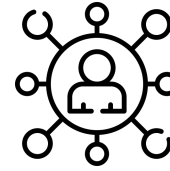
Data



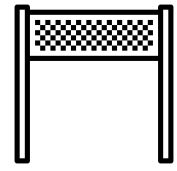
Methodology



Models



Use cases



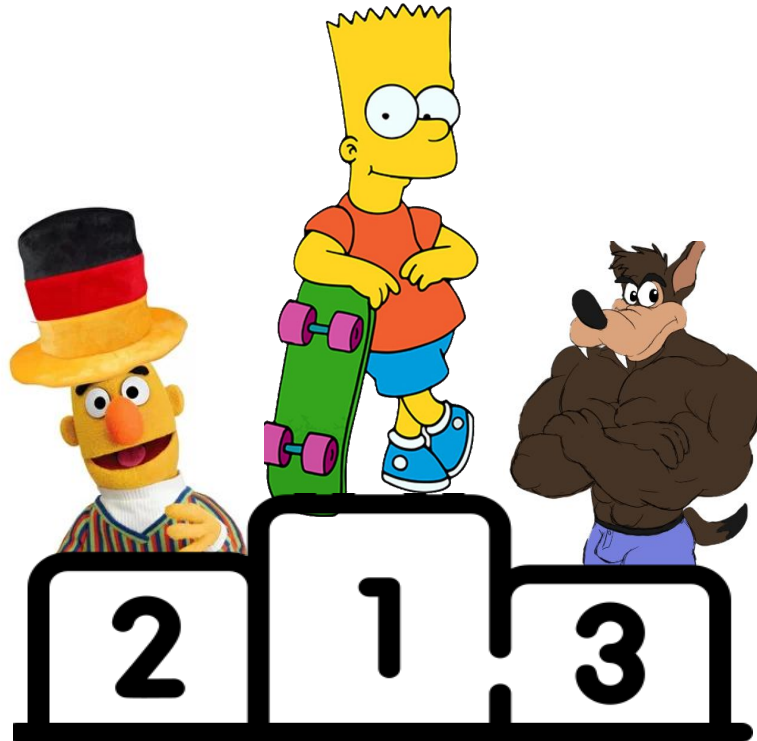
Conclusion

Model from scratch

MeanSum

BERT & BART

Overall comparison

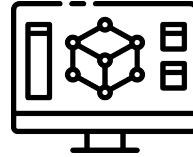




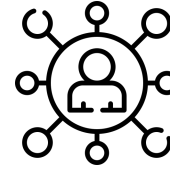
Data



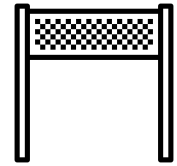
Methodology



Models



Use cases



Conclusion

Journalism

MVP

**When writing articles, journalists face a vast amount of session protocols and Drucksachen**

**An easy way to retrieve key points would facilitate the life of journalists considerably**

**To cover the Bundestag, speed and accuracy are key**



**Meinolf Ellers | CDO @ DPA**

**When writing articles, journalists face a vast amount of session protocols and Drucksachen**

**An easy way to retrieve key points would facilitate the life of journalists considerably**

**To cover the Bundestag, speed and accuracy are key**



**Meinolf Ellers | CDO @ DPA**

**When writing articles, journalists face a vast amount of session protocols and Drucksachen**

**An easy way to retrieve key points would facilitate the life of journalists considerably**

**To cover the Bundestag, speed and accuracy are key**



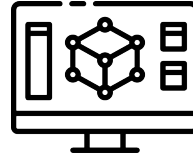
**Meinolf Ellers | CDO @ DPA**



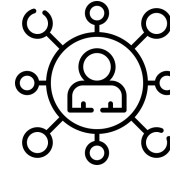
Data



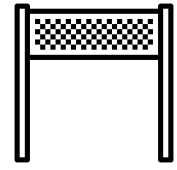
Methodology



Models



Use cases

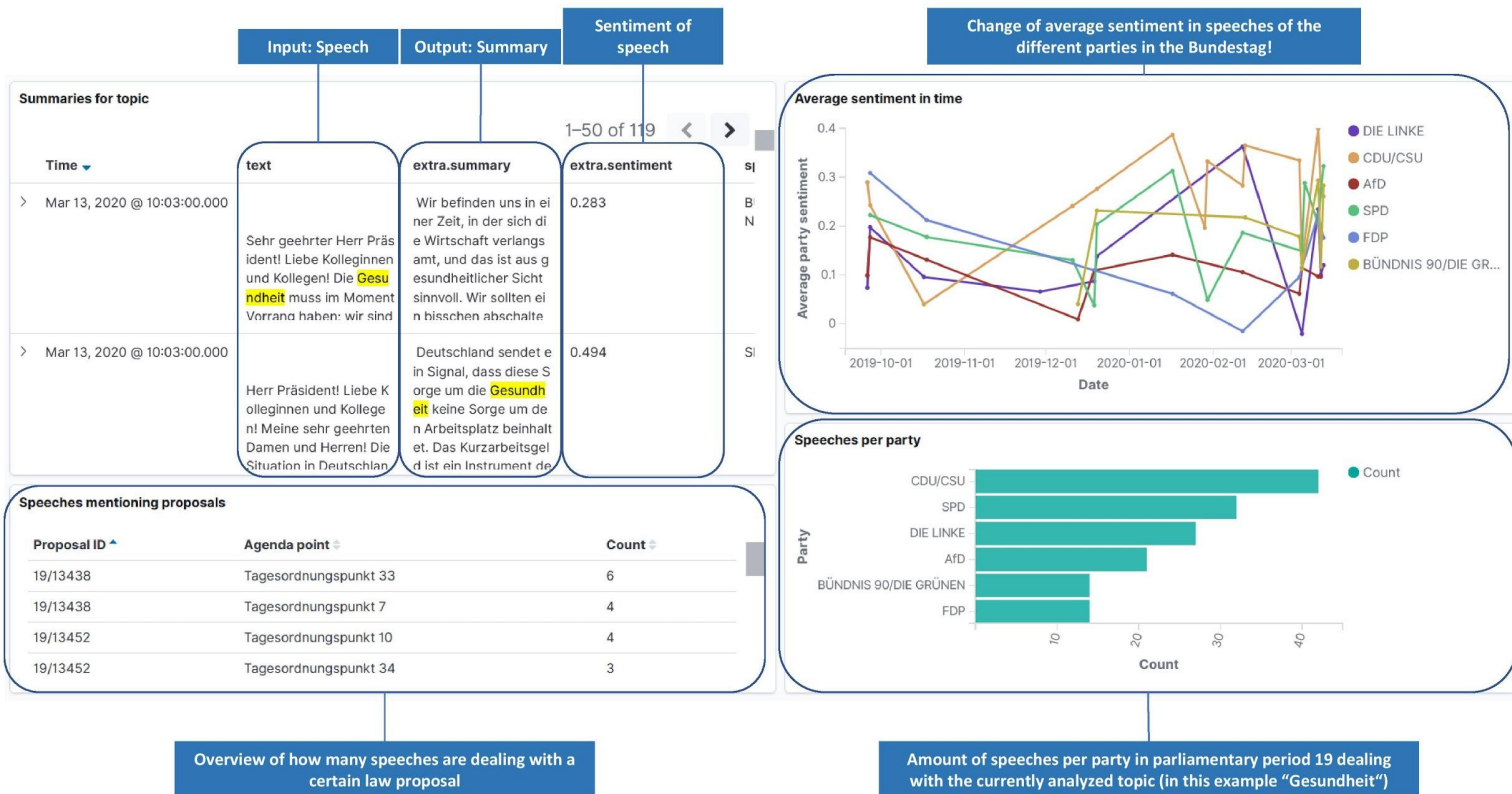


Conclusion

Journalism

MVP

# Our minimum viable product

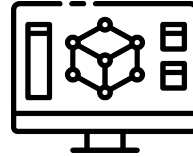




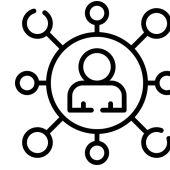
Data



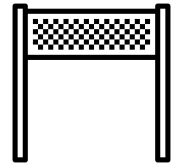
Methodology



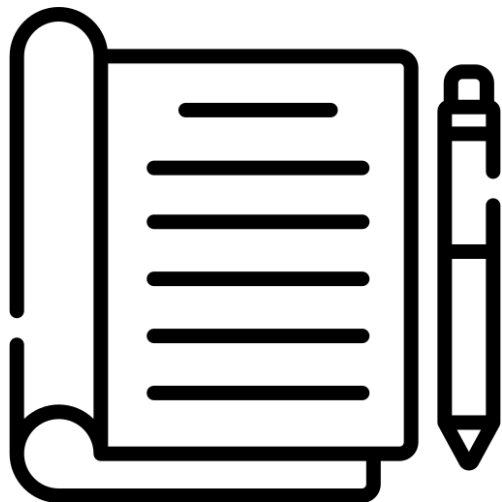
Models



Use cases



Conclusion



1

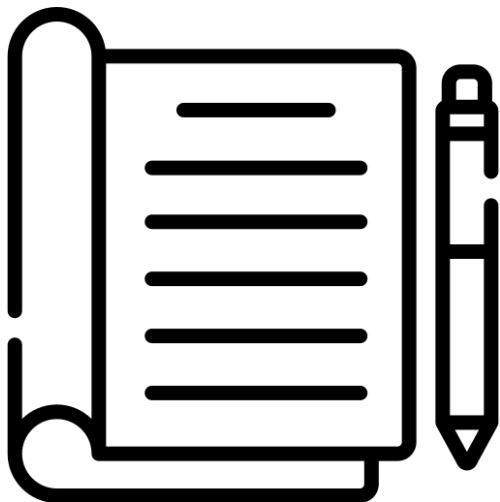
**Use pretrained models**

2

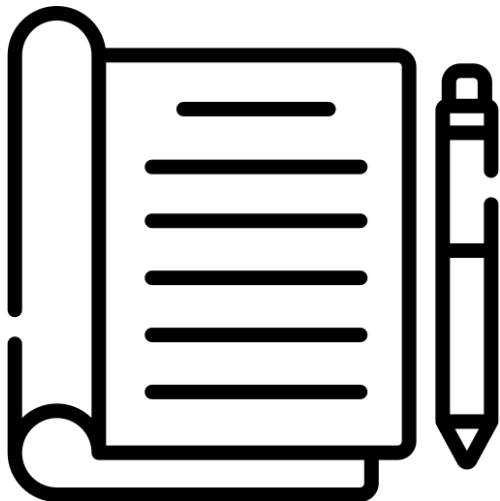
**Use supervised approaches**

3

**Use English-based models**



- 1 Use pretrained models
- 2 Use supervised approaches
- 3 Use English-based models



1

Use pretrained models

2

Use supervised approaches

3

Use English-based models

**Thank you for your  
support throughout our project!**

- [1] Pai, A. (2019). Comprehensive Guide to Text Summarization using Deep Learning in Python. Retrieved from <https://www.analyticsvidhya.com/blog/2019/06/comprehensive-guide-text-summarization-using-deep-learning-python/> .
  
- [2] Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2018). Bert: Pre-training of deep bidirectional transformers for language understanding. arXiv preprint arXiv:1810.04805.
  
- [3] Chu, E., & Liu, P. (2019, May). MeanSum: a neural model for unsupervised multi-document abstractive summarization. In International Conference on Machine Learning (pp. 1223-1232).
  
- [4] Liu, Y., & Lapata, M. (2019). Text summarization with pretrained encoders. arXiv preprint arXiv:1908.08345.
  
- [5] Lewis, M., Liu, Y., Goyal, N., Ghazvininejad, M., Mohamed, A., Levy, O., ... & Zettlemoyer, L. (2019). Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. arXiv preprint arXiv:1910.13461.

All icons made by Freepik from [www.flaticon.com](http://www.flaticon.com)

BART images adapted from

<https://www.goodfon.com/download/multfilm-shou-simpsons-personazh-bart-art-risunok-multseri-2/1920x1080/>

and <https://www.pinterest.de/pin/680676931152976911/>

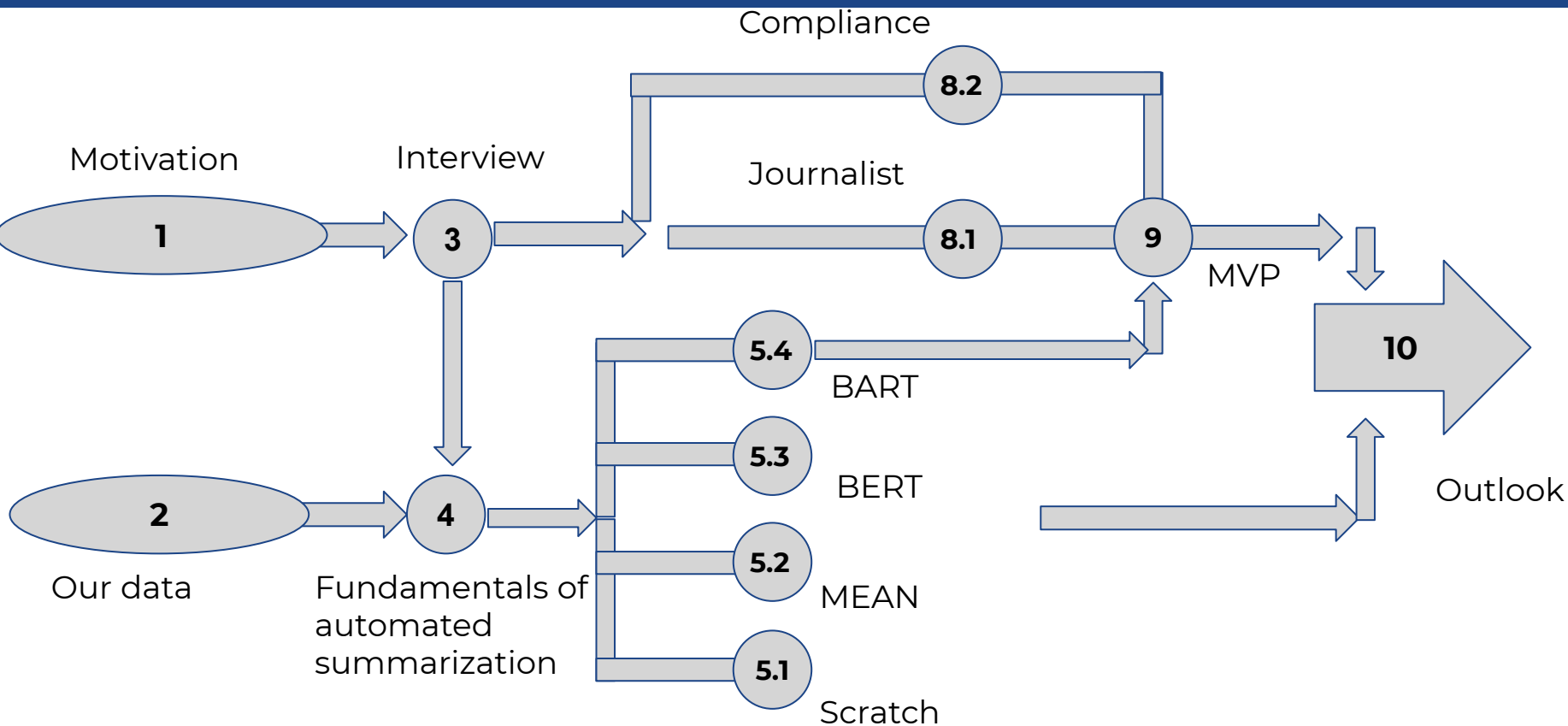
GermanBERT from <https://deepset.ai/german-bert>

HuggingFace logo from <https://huggingface.co/>

OpenAI Logo from <http://allvectorlogo.com/openai-logo/>

Logo “abstraction captioning” <https://creativemarket.com/eucalyp>

**Back-up**



**MongoDB-storage**

**Established text indexing for fast  
searching capabilities**

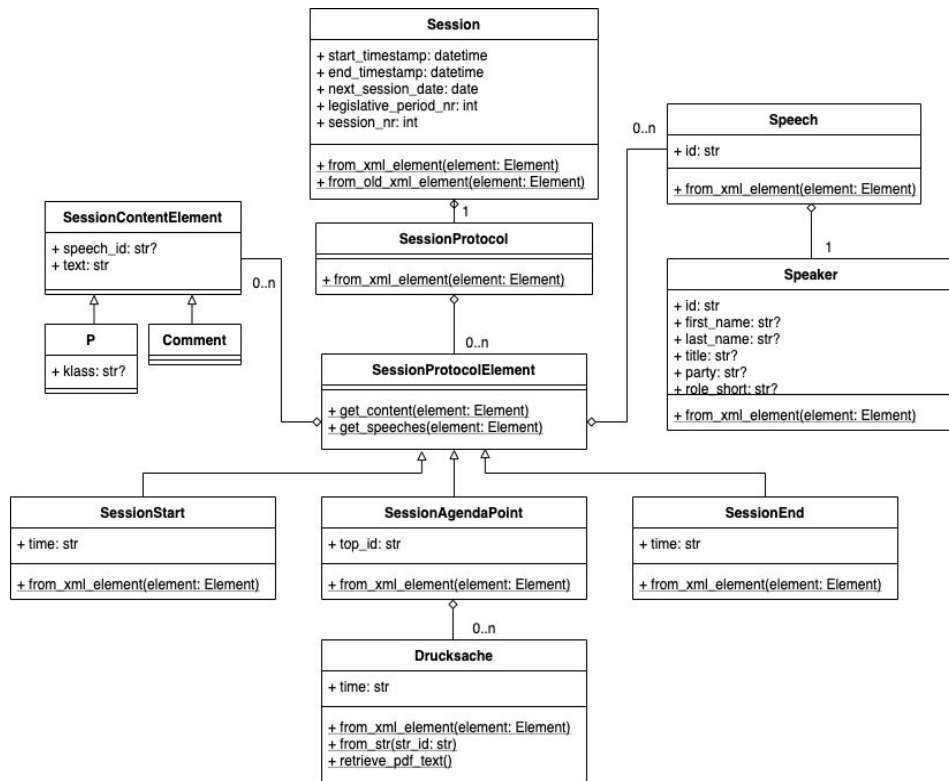
**Able to access additional  
metadata for all Bundestag  
members from 1940 to 2020**

```
_id: "19/1"
start_timestamp: 2017-10-24T11:00:00.000+00:00
end_timestamp: 2017-10-24T17:03:00.000+00:00
next_session_date: 2017-11-22T00:00:00.000+00:00
legislative_period_nr: 19
session_nr: 1
agenda_points: Object
  1: Object
    introduction: Array
    speech: Object
      ID19100100: Object
        _id: "ID19100100"
        speaker: Object
          _id: "11002190"
          first_name: "Alterspräsident Dr. Hermann"
          last_name: "Otto Solms"
          title: null
          role_short: "Alterspräsident"
          role_long: "Alterspräsident"
          location: null
          party: null
        speech_elements: Array
          ID19100200: Object
        2: Object
        3: Object
```

Able to access data from pp14 to pp18 in the same way as the data from pp19

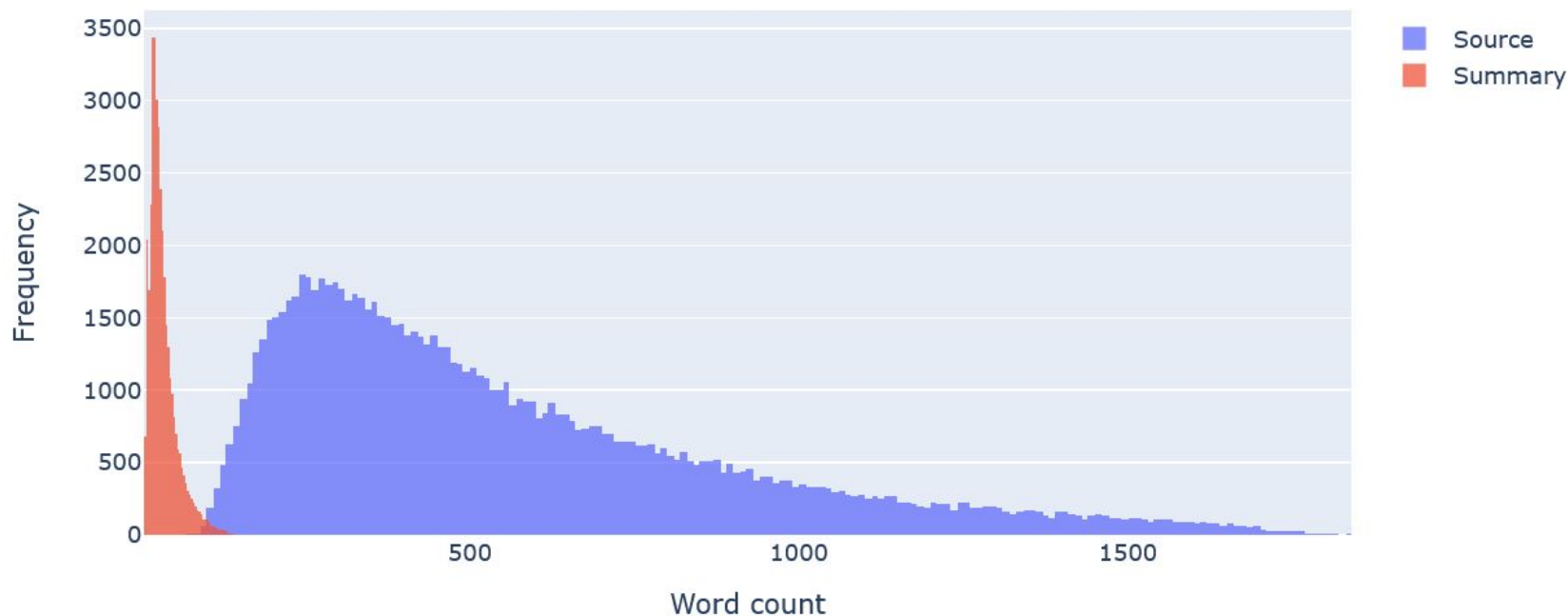
Established automatic extraction of Drucksachen mentioned in protocol text

Created interface for downloading and extracting current Drucksachen from Bundestag website

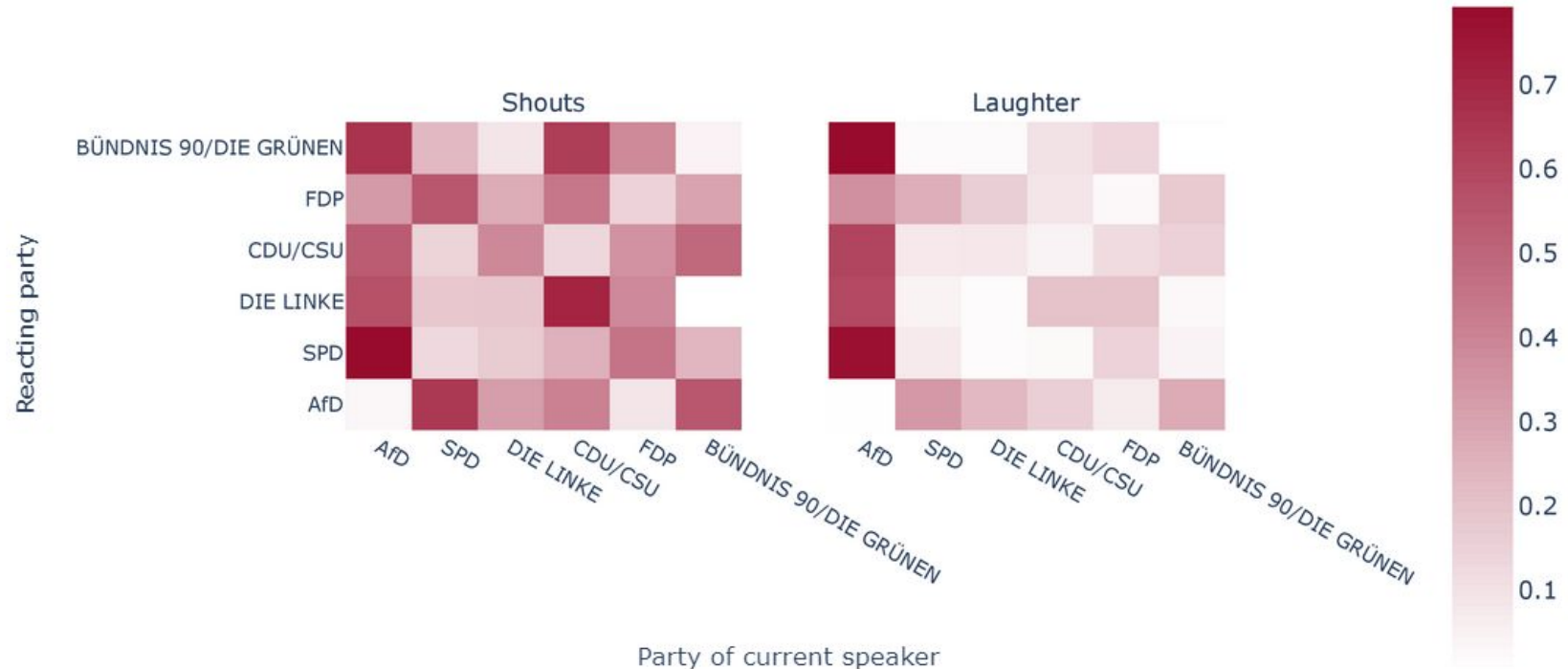




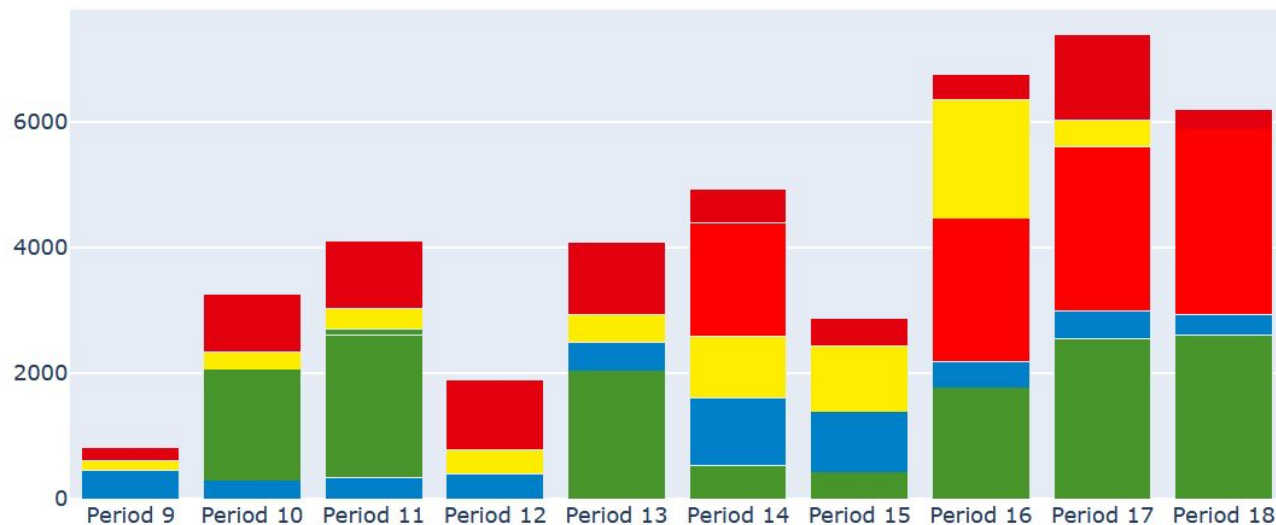
Swiss dataset word count per sample

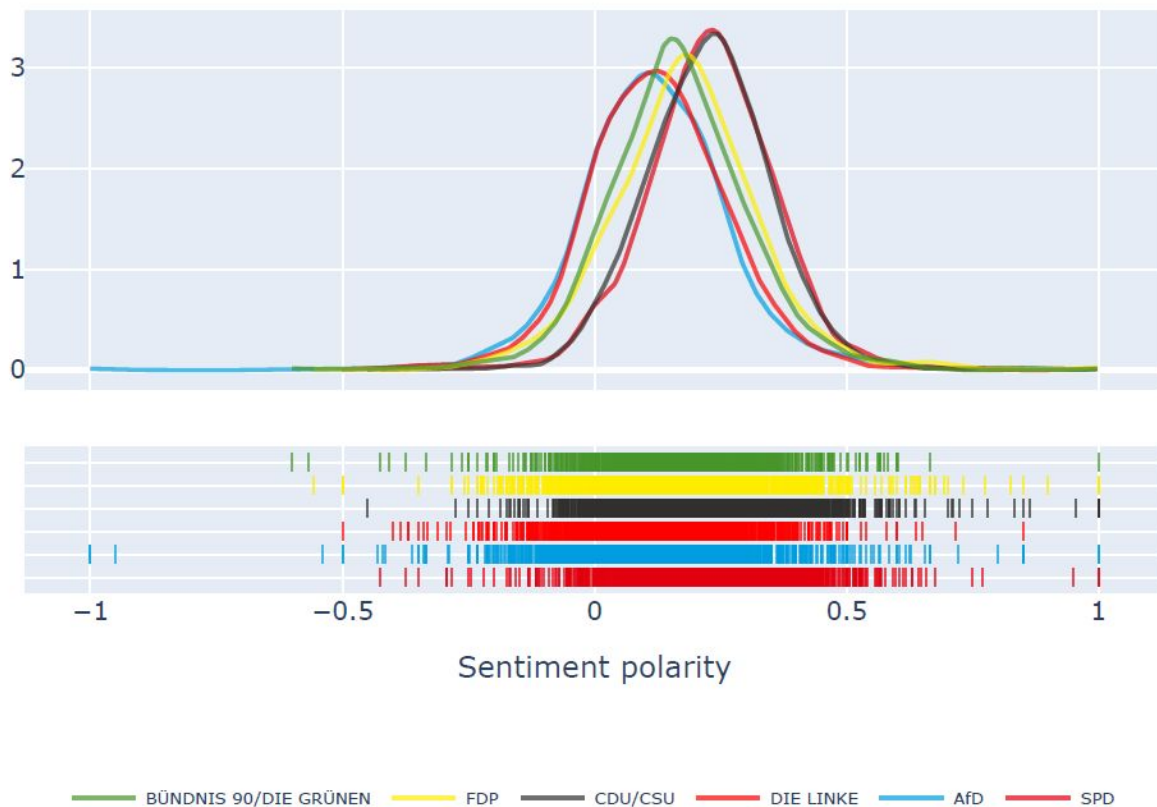


| Legislative Period | Years     | Words                           |
|--------------------|-----------|---------------------------------|
| 14                 | 1998-2002 | Kuba, Stammzellen, Öcalan       |
| 15                 | 2002-2005 | Moldau, HIV, Ukraine, Baukultur |
| 16                 | 2005-2009 | Cannabis, Roma/Sinti, Sri Lanka |
| 17                 | 2009-2013 | Glyphosat, Westsahara, Honig    |
| 18                 | 2013-2017 | Europol, Isis, Substanzen       |
| 19                 | 2017-...  | Gigawatt, Adoption, LKWs        |

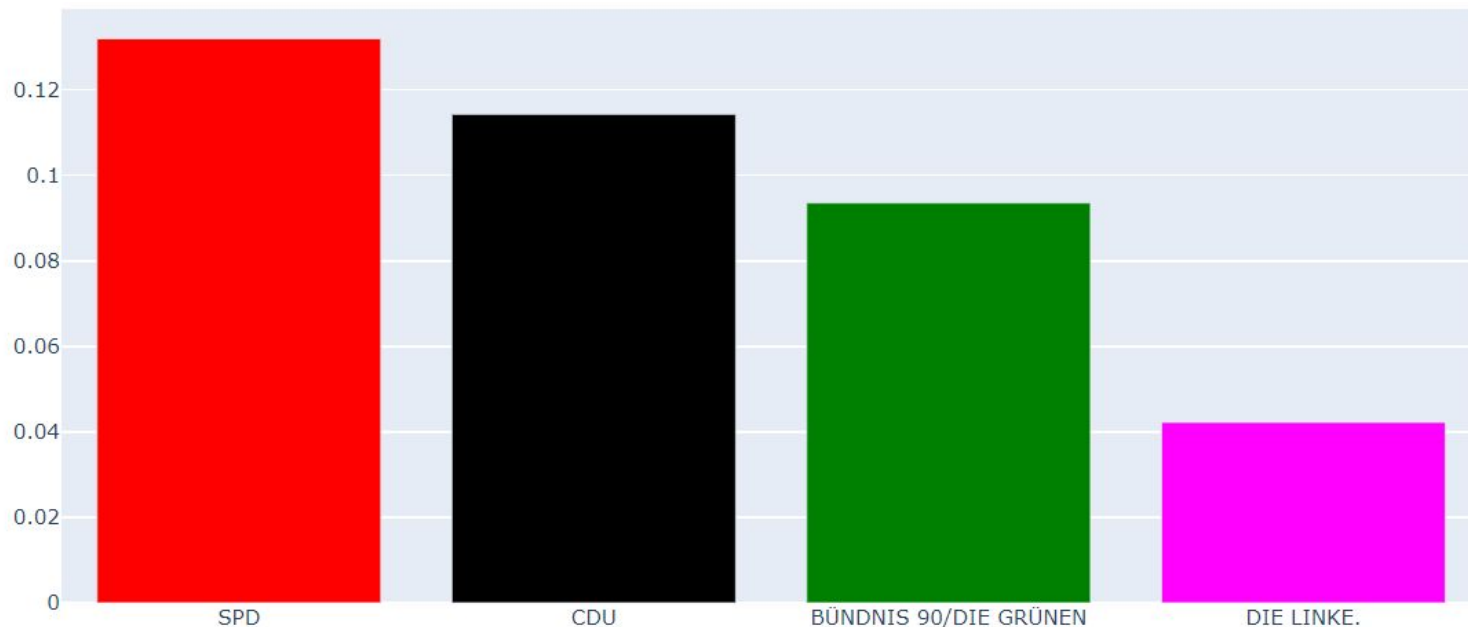


Number of Drucksachen Proposed per Party

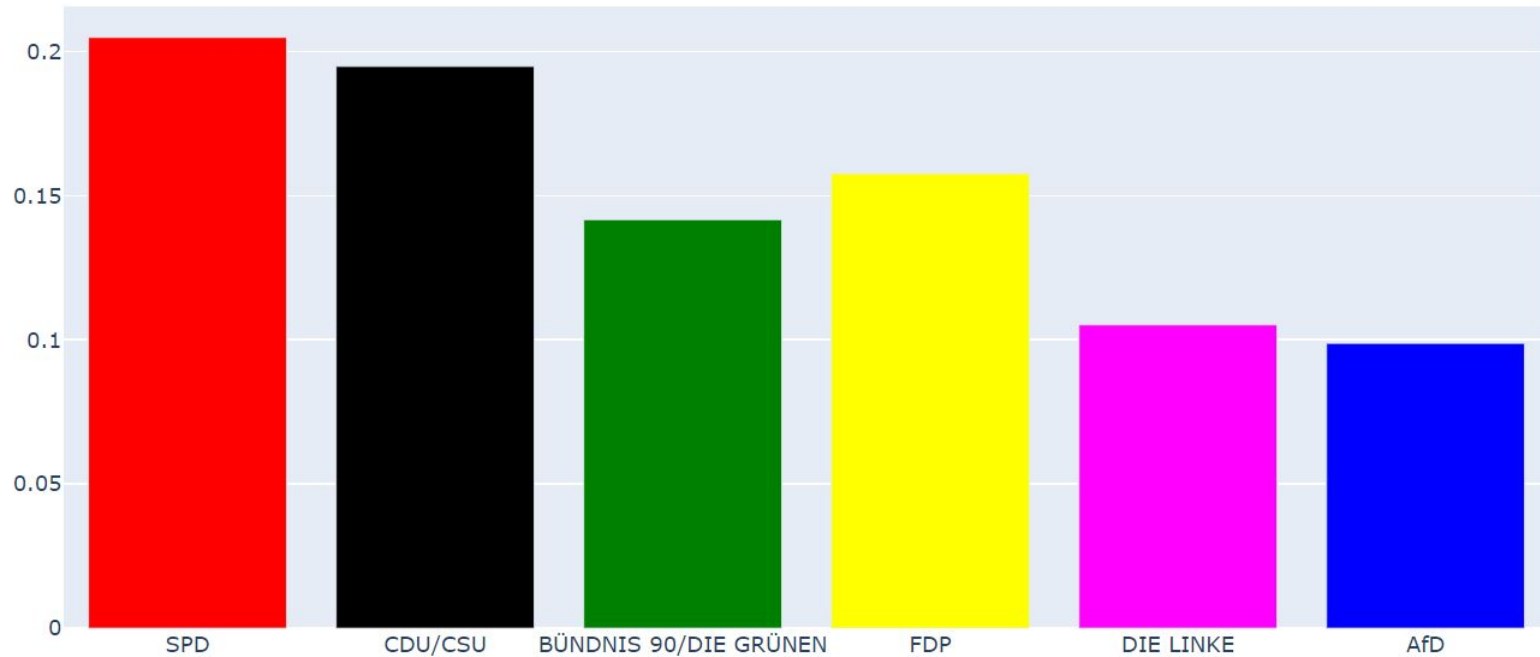




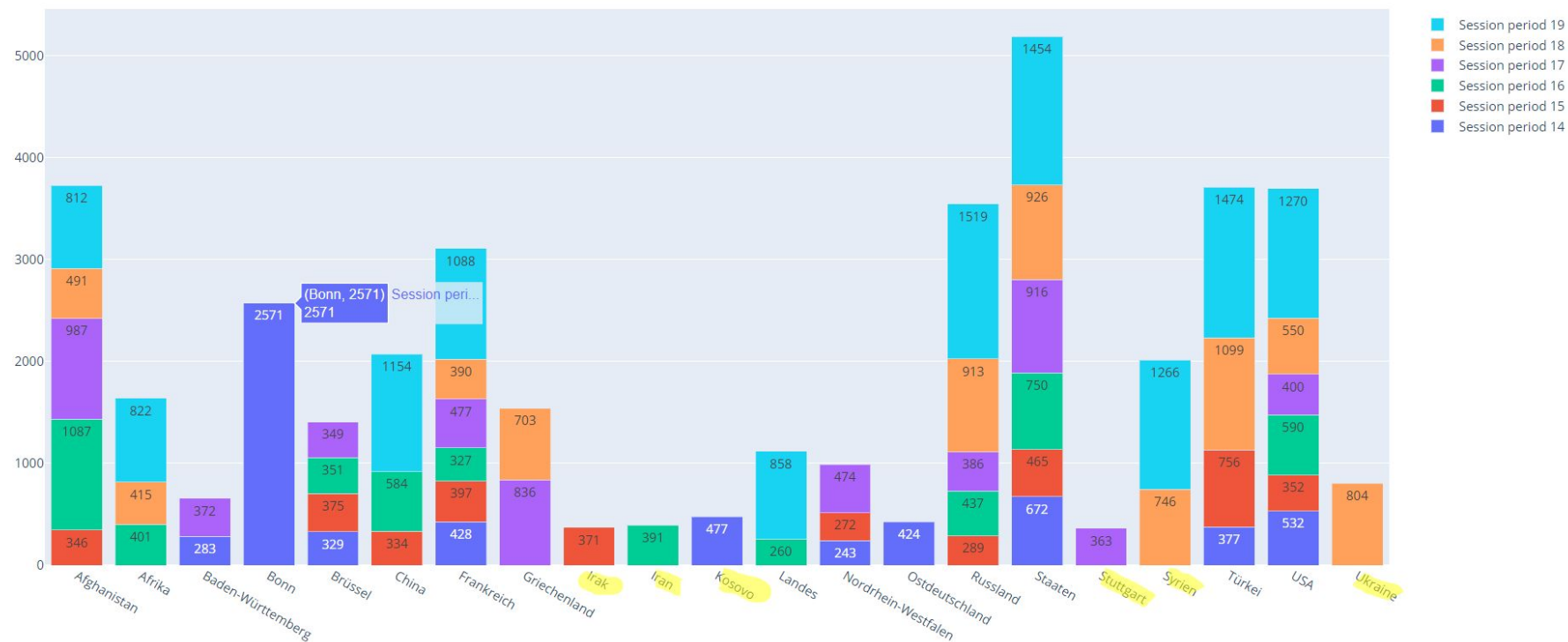
Overall-mean-polarity speeches PP18: 0.10121783467607333



Overall-mean-polarity speeches PP19 0.1590167363025736



Locations mentioned in protocol speeches over session periods



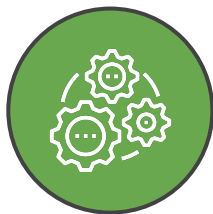
## Intuition

Recall:  
$$\frac{\text{number of overlapping } n\text{-grams}}{\text{total } n\text{-grams in reference summary}}$$

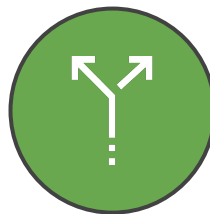
Precision:  
$$\frac{\text{number of overlapping } n\text{-grams}}{\text{total } n\text{-grams in system summary}}$$

F1-Score:  
$$\frac{2}{1/\text{Recall} + 1/\text{Precision}}$$

## Pro & Contra



Computational effort



Number of variants

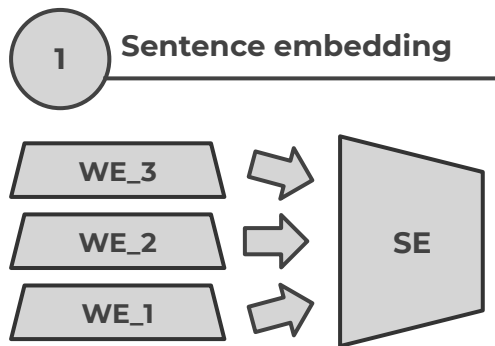


Abstraction capturing

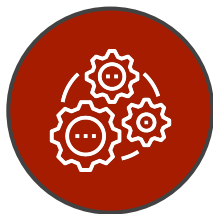
ROUGE



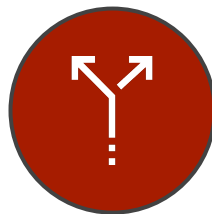
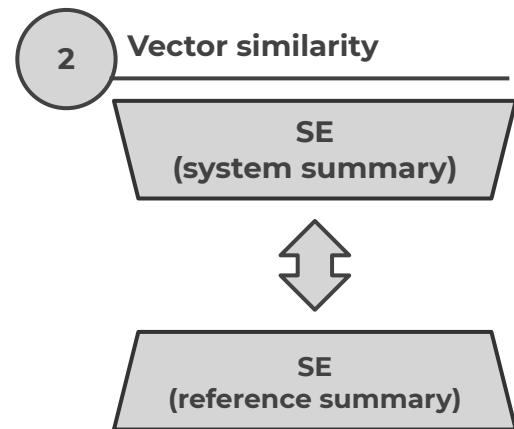
## Intuition



## Pro & Contra



Computational effort



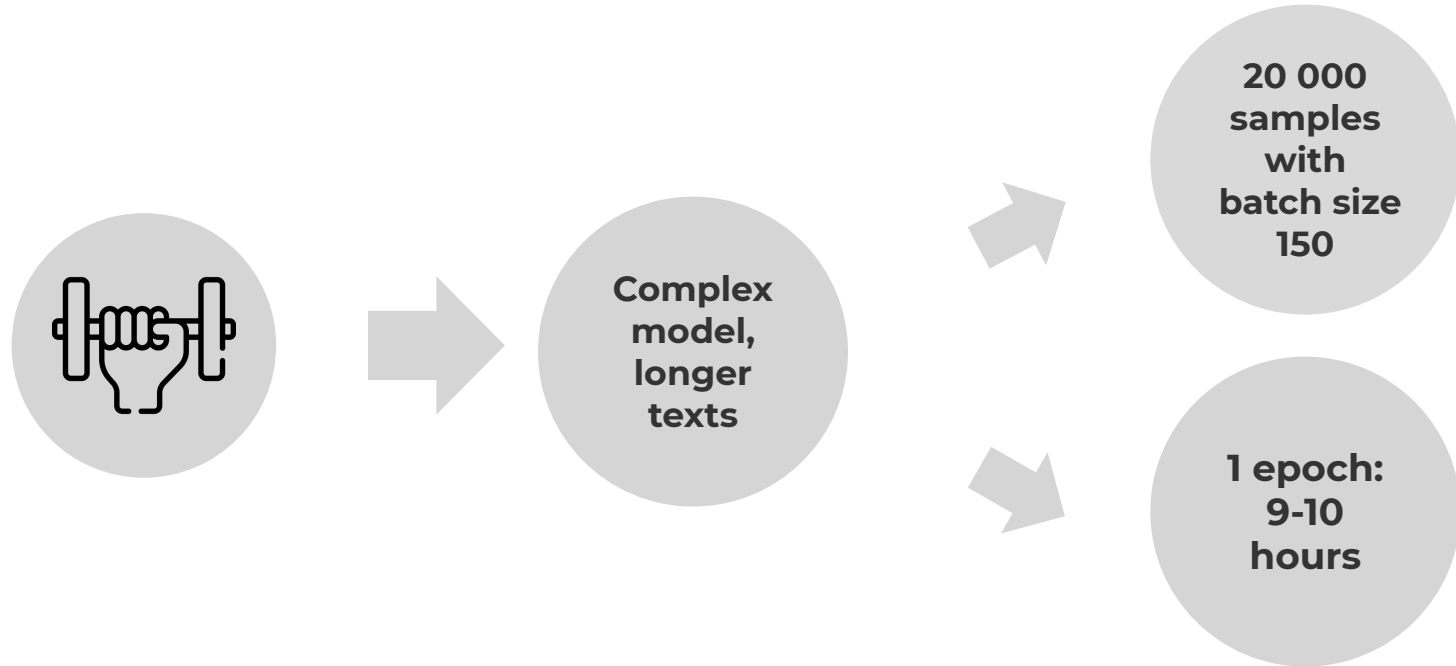
Number of variants



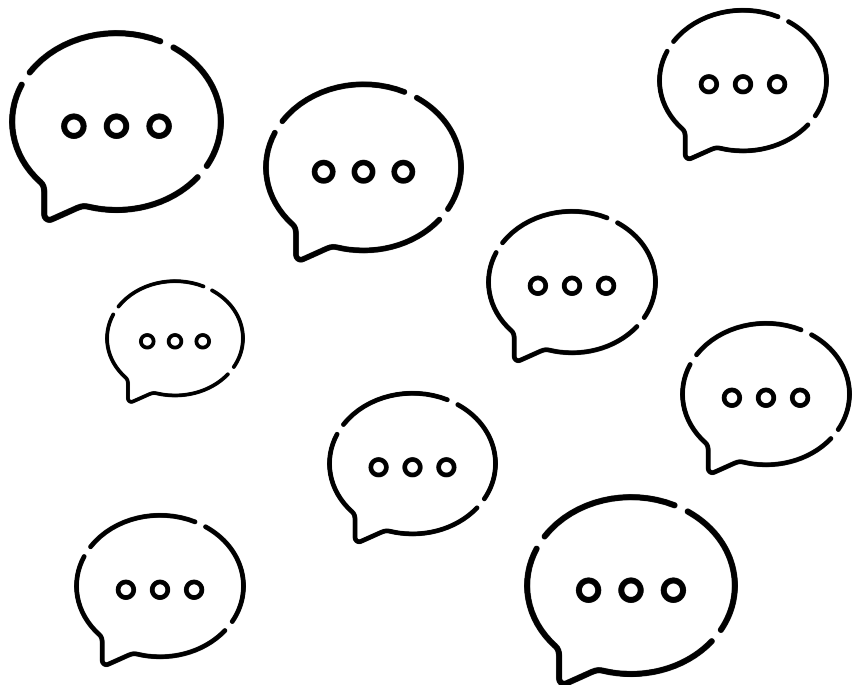
Abstraction capturing

$$\text{similarity}(A,B) = \frac{A \cdot B}{\|A\| \times \|B\|} = \frac{\sum_{i=1}^n A_i \times B_i}{\sqrt{\sum_{i=1}^n A_i^2} \times \sqrt{\sum_{i=1}^n B_i^2}}$$

1. On a scale from 0 (*miserable*) to 10 (*excellent*), evaluate the **grammatical correctness** of the machine-generated summary.
2. On a scale from 0 (*miserable*) to 10 (*excellent*), evaluate the **fluency** of the machine-generated summary.
3. On a scale from 0 (*no key points captured*) to 10 (*all key points captured*), evaluate if the machine-generated summary has **captured all important key-points of the input-text**.
4. On a scale from 0 (*content is so mixed-up, that meaning is changed fundamentally*) to 10 (*content is not mixed up at all*), evaluate if the machine-generated summary has **mixed up the content of the input-text**.
5. On a scale from 0 (*summary talks about smth completely different*) to 10 (*summary didn't add any misleading information*), evaluate if the machine-generated summary has **added misleading information compared with the input-text**.

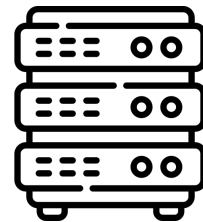


## Multi-document summarization

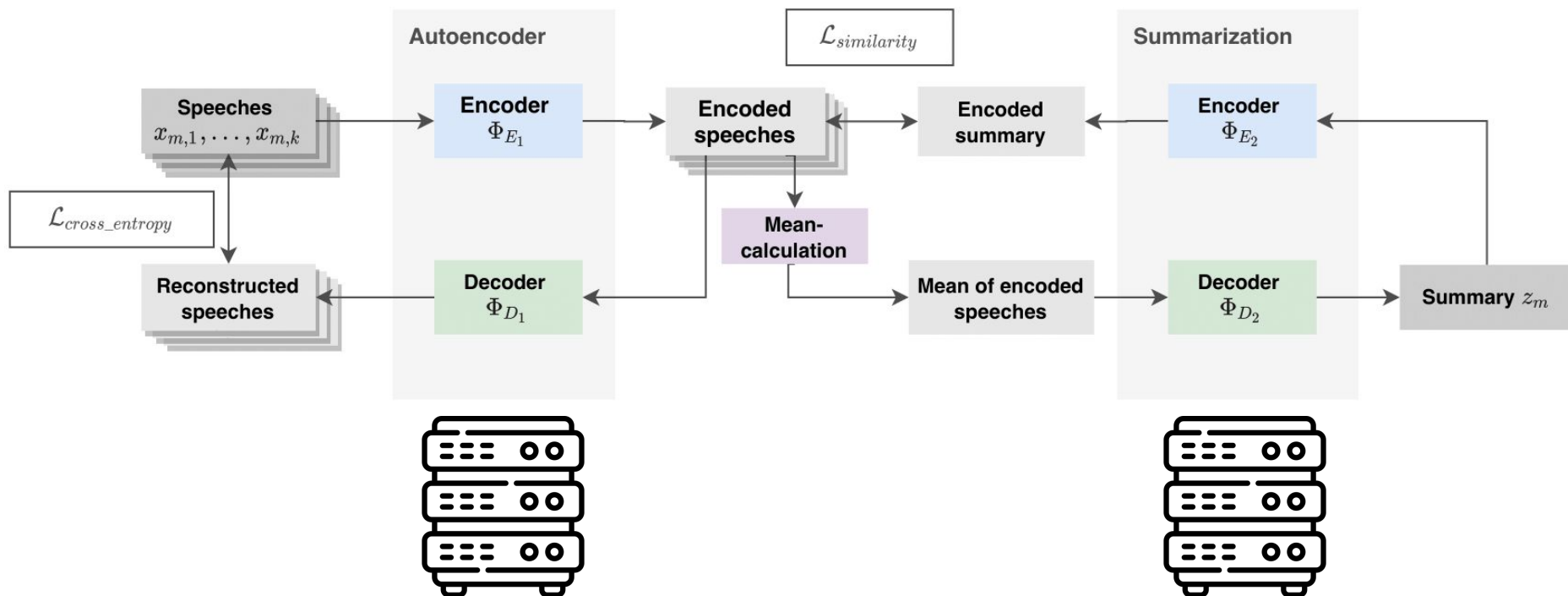


## Translation-approach

MeanSum based on  
pretrained english language  
model

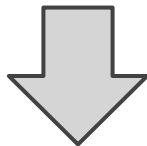


Translate speeches to  
english and translate  
obtained summaries back



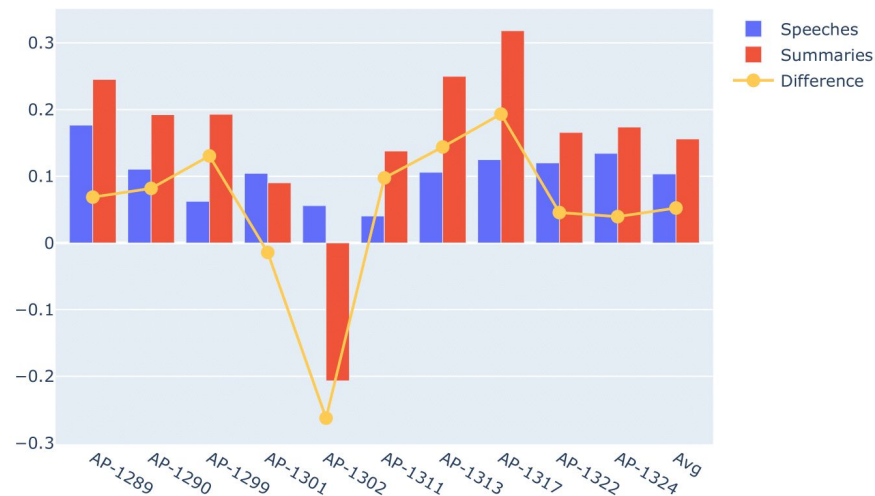
## Rouge-scores

| R1-Score | R2-Score | RL-Score |
|----------|----------|----------|
| 0.201    | 0.043    | 0.108    |



High abstraction

## Sentiment analysis



| Model            | R1-score             | R2-score             | RL-score             | BERT-score           | precision  |
|------------------|----------------------|----------------------|----------------------|----------------------|------------|
| BART             | 0.2024 (0.07)        | 0.0456 (0.04)        | 0.1955 (0.07)        | <b>0.6251 (0.05)</b> | 0.7274     |
| BERTSumAbs       | <b>0.2222 (0.12)</b> | <b>0.0801 (0.10)</b> | <b>0.2479 (0.13)</b> | 0.6224 (0.08)        | 0.4454     |
| lead-sum         | 0.1882 (0.08)        | 0.0412 (0.05)        | 0.1777 (0.07)        | 0.6112 (0.05)        | <b>1.0</b> |
| english-lead-sum | 0.1782 (0.07)        | 0.0341 (0.04)        | 0.1682 (0.06)        | 0.6049 (0.04)        | 0.7317     |

*SwissText-Dataset*

| Model             | R1-score      | R2-score      | RL-score      | BERT          |
|-------------------|---------------|---------------|---------------|---------------|
| BART              | 0.1275        | 0.0288        | 0.1383        | <b>0.5994</b> |
| distil-BERTSumAbs | 0.0983        | 0.0053        | 0.0784        | 0.4867        |
| rand-sum          | <b>0.2300</b> | <b>0.0338</b> | <b>0.1833</b> | 0.5863        |
| lead-sum          | 0.0123        | 0             | 0.174         | 0.4713        |

*Heute im Bundestag-Dataset*

## Takeaways:

- BERT uses patterns
- BART is highly extractive
- HiB contains patterns
- BART generalizes best

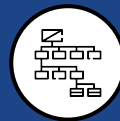


Source: <https://alammar.github.io/illustrated-gpt2/>

Utilized the 345M German GPT-2 model trained on a 27Gb twitter + wikipedia + heise + parole corpus. <http://zamia-speech.org/>

**Fine-tuning for Summarization:**  
**<source> + <tldr> + <summarization>**

**Inference: <source> + <tldr> =>**  
**generated tokens is the summarization**



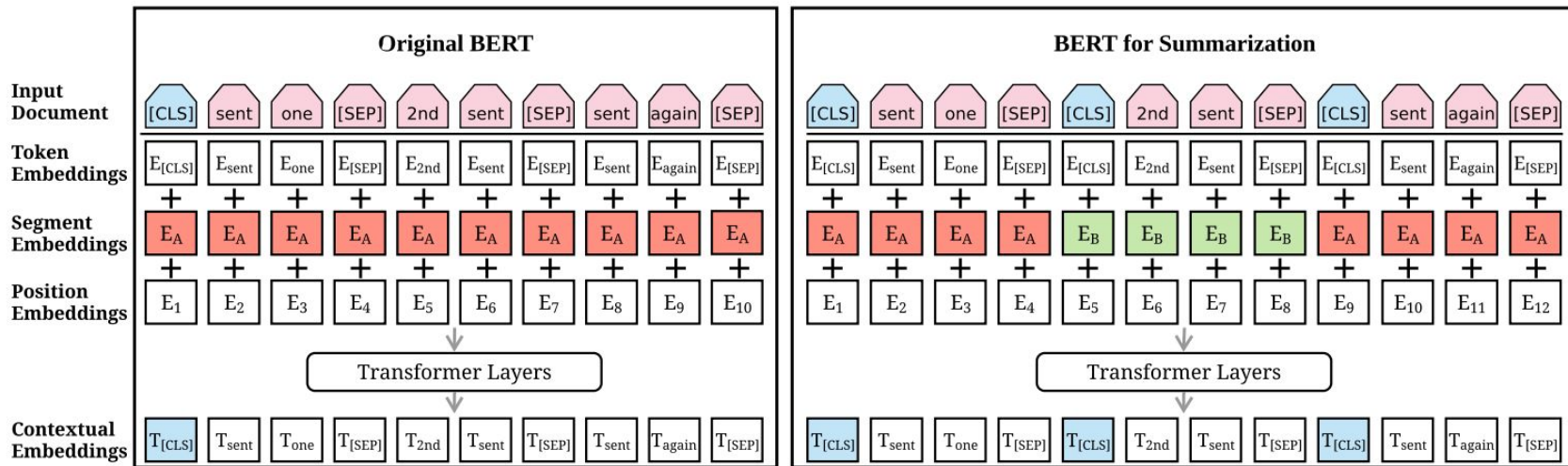
**German BERT**



**Hugging Face**



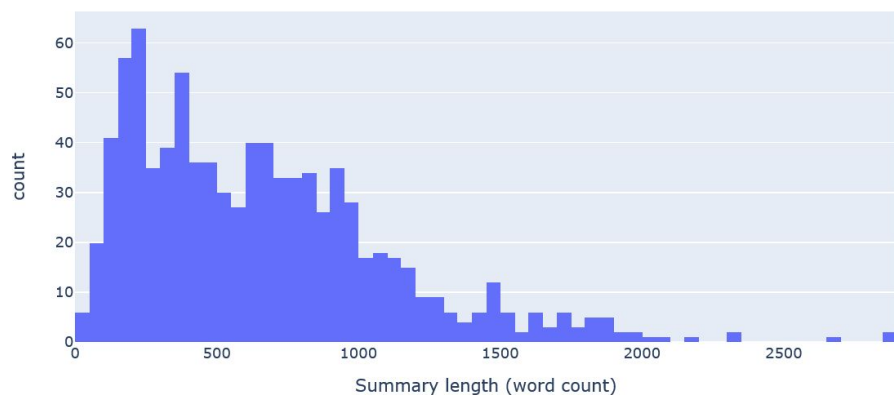
**German Distil-BERT**



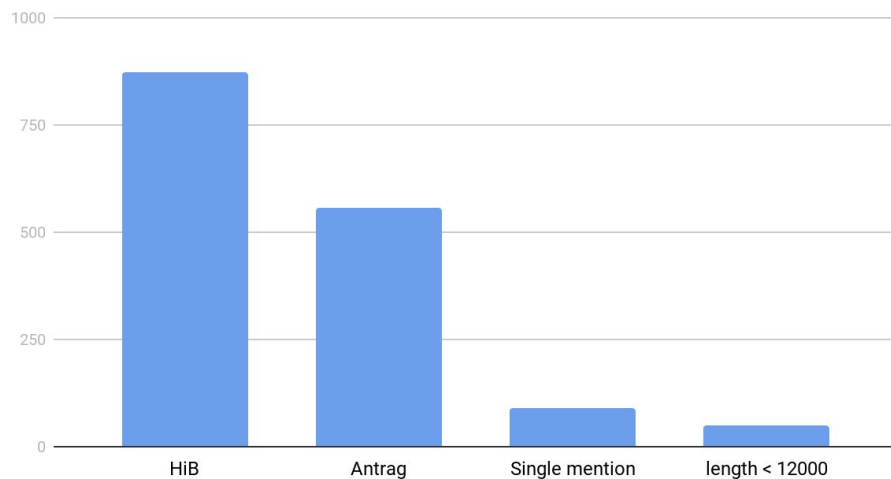
Utilized German Bert and Huggingface Transformers library

Trained on 80k samples from Swiss dataset

Bundestag text archive word count



HiB Dataset



## lead\_sum

Takes first three sentences of source as summary

```
[29] ▶ M4   
def lead_sum(text:str) -> str:  
    first_three = text.split(".")[0:3]  
    return ".".join(first_three)+"."
```

## rand\_sum: How can we mimic overfitting models?

In our case overfitting means that the model

- ignores the input
- returns some summary it has seen during the training process

```
[-] ▶ M4   
def rand_sum(text:str) -> str:  
    return random.choice(summaries_in_trainingset)
```

John Stewart war ein US - amerikanischer Politiker . Zwischen 1823 und 1823 vertrat er den Bundesstaat Pennsylvania im US - Repraesentantenhaus.

William Brown war ein US - amerikanischer Politiker . Zwischen 1847 und 1823 vertrat er den Bundesstaat Virginia im US - Repraesentantenhaus.

Samuel Flander war ein US - amerikanischer Politiker . Zwischen 1847 und 1823 vertrat er den Bundesstaat New York im US - Repraesentantenhaus.

**John Stewart** war ein US-amerikanischer Politiker. Zwischen 1843 und 1845 vertrat er den zweiten Wahlbezirk des Bundesstaates Connecticut im US-Repräsentantenhaus.

- Wikipedia: [https://de.wikipedia.org/wiki/John\\_Stewart\\_\(Politiker,\\_1795\)](https://de.wikipedia.org/wiki/John_Stewart_(Politiker,_1795))



## Context

Speech given by Hermann Färber (CDU/CSU) on 11<sup>th</sup> of November 2019

Bundestag is discussing an agricultural bill

## Key points

Providing direct payments of 4.50€ per hectare grazing land for sheep and goat farmers

Ensuring legal and financial security for farmers participating in agri-environmental programs

Increasing demand for products from sheep and goat farmers



**Hermann Färber (CDU/CSU)**

## From scratch:

*"vereins vereins vereins vereins verbandsfreie verbandsfreie verbandsfreie ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer rat-meyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyererratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer rat-meyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer ratmeyer"*

## MeanSum:

*"Hallo, liebe Kollegen! Liebe Gäste! Es geht darum, die Landwirtschaft zu fördern, auch wenn man sich das Ganze in einem Biergarten ansieht, der jetzt genutzt wird. Dies ist das Ergebnis der Bundesregierung und 2020. Meine Damenund Herren, meine Damen und Herren, lassen Sie uns auf die Erde gehen. Wir sind mitten im Budget. Dies ist ein großer Schritt nach vorne. Das kann man nicht einfach essen. Sie können auch darüber sprechen. Sie fordern das Steuerprivileg des EEG-Zuschlags, zum Beispiel für ein Video von 418 Millionen, das dann natürlich von der Regierung finanziert wird. Dies ist ein sehr wichtiger Schritt in Richtung Weltklima- und Umweltpolitik."*

## BERT:

*"Die Simmulation ist ein Verfahren zurückragung, bei der für die für die für die führte. Das Verfahren ist in der Regel für das Verfahrenzurückragung der Regel führt der Regel führte das Verfahren zurückragungen vonder Regel führte - und das Verfahren zurührt - das Verfahren zur Verfügung derRegel führt der Regel furt - Änderung der Regel führte das Verfahren der Regelführte."*

## lead-sum:

*"Sehr geehrter Herr Präsident! Meine Damen und Herren! Wir debattieren heute das Direktzahlungen-Durchführungsgesetz und damit die einmalige Erhöhung der Umschichtung um 1,5 Prozentpunkte – das entspricht einer Erhöhungum 4,50 Euro pro Hektar – aus der ersten in die zweite Säule."*

## BART:

*"Die AfD debattiert über das Gesetz zur Umsetzung von Direktzahlungen und die einmalige Erhöhung der Umschichtung um 1,5 Prozentpunkte. Dies entspricht einer Steigerung von 4,50 EUR pro Hektar von der ersten zur zweiten Säule der Agrarumweltprogramme. Die Gruppen Die Linke und Bündnis 90 / DieGrünen erkennen in ihrer Bewerbung zu Recht, dass die Schafzüchter mit ihrer Arbeit uns qualitativ hochwertige Produkte in der Bevölkerung liefern."*

## gerBart:

*"Wir debattieren heute das Direktzahlungen-Durchführungsgesetz und damit die einmalige Erhöhung der Umschichtung um 1,5 Prozentpunkte. Wir erwarten jedoch, das entspricht Erhöhung um 4,50 Euro pro Hektar – ausschließlichen die landwirtschaftliche Mittelvergabe."*

## Speech:

Herr Präsident! Liebe Kolleginnen und Kollegen! In dieser Krise zeigt sich eine gewisse Doppelgesichtigkeit: Auf der einen Seite beweisen der Sozialstaat und die Sozialversicherungen in diesen Monaten, dass sie selbst in der jetzigen Extremsituation mit Milliardeneinsatz in der Lage sind, in der Krise ein gewisses Maß an Sicherheit zu geben. Auf der anderen Seite zeigt das umwälzende Ausmaß der Krise auch, dass all diese Gegenmaßnahmen die Erschütterungen durch Arbeitslosigkeit, Existenzverlust und Bildungsmangel nicht vollständig auffangen können. Tiefe Gräben tun sich auf. Es sind gerade die Verwundbarsten und die ärmsten Gruppen, die ohnehin am Rande der Gesellschaft stehen, die nun umso mehr den Anschluss verlieren: Menschen, bei denen das Kurzarbeitergeld nicht reicht, Familien, die Covid-19 vor fast unlösbare Probleme stellt, Frauen, die den größten Teil der Sorgearbeit leisten müssen, Wohnungslose ohne jede Unterstützung oder Studierende...

## Summary bart:

Durch die Pandemie fühlen sich viele Menschen einsam, überwältigt und sogar verzweifelt. Die Coronakrise ist daher auch eine Krise des sozialen und sozialen Zusammenhalts. Materielle Unterstützung allein reicht nicht aus. Was benötigt wird, ist eine Strategie, die alle dabei unterstützt, besser mit bevorstehenden Unsicherheiten umzugehen, die durch all die anderen Krisen verursacht werden, die wir noch haben.

## Summary bertsum:

die französische organisation , die in den vereinigten staaten in der vereinigten staaten , ist . sie ist in der vereingte staaten . ist die grundlage der französischen staaten . sie wurde im jahr 2010 in den usa . wurde sie in den jahren von der vereinigten konigreich gegründet . gehört zu den grossten staaten , die sich im jahre . der grundlage ist die grossten staaten . die entwicklung der grundung der grundlage der grundung des computerspielen . der entwicklung der grundmandate . sie ist sie die sie in der grund die grundung

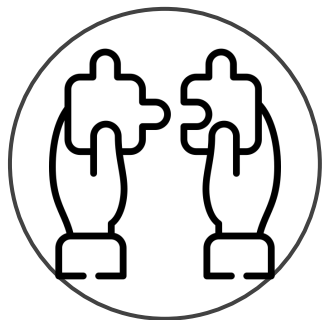
**I need to analyze Drucksachen to understand the intentions of certain laws**

**Unfortunately, there is no convenient way to search for specific law proposals**

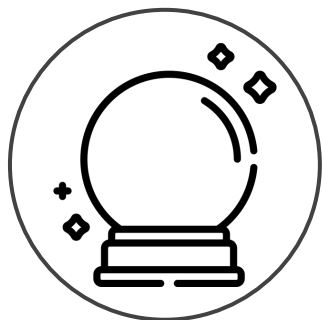
**Therefore, I often rely on Blogposts and other secondary sources of information**



**Tobias Gump | Law student**



- **MODELS:**
  - Model from scratch, BART, BertSUM, MeanSum
- **SIDE-PRODUCTS:**
  - Parsers
  - Translation Pipeline
  - Dataset
- **USE-CASES:**
  - Talked to experienced people in the journalism field



- More sophisticated pretrained language models are released every year
- Our parsing tools can facilitate the creation of German parliamentary data.
- Our translation pipeline can be used to create English summarization models
- Approaches can be extended for query-based or topic-based summarization and generation of summaries of user-defined length

## Challenges

- Remote team communication
- Different schedules
- Different levels of expertise
- Lack of meeting structure

## Improvements

- Regular meetings (2x-3x per week)
- Version control rules
- Team code reviews and mentoring
- Meeting agendas & time keeping