

ArgLegalSumm: Improving Abstractive Summarization of Legal Documents with Argument Mining

Mohamed Elaraby, Diane Litman

University of Pittsburgh

Pittsburgh, PA, USA

{mse30, dlitman}@pitt.edu

Abstract

A challenging task when generating summaries of legal documents is the ability to address their argumentative nature. We introduce a simple technique to capture the argumentative structure of legal documents by integrating *argument role labeling* into the summarization process. Experiments with pre-trained language models show that our proposed approach improves performance over strong baselines.

1 Introduction

Abstractive summarization has made great progress by leveraging large pretrained language models such as BART (Lewis et al., 2020), T5 (Raffel et al., 2020), Pegasus (Zhang et al., 2020), and Longformer (Beltagy et al., 2020). These models leverage large scale datasets such as CNN-DailyMail (Hermann et al., 2015), PubMed (Cohan et al., 2018), and New York Times (Sandhaus, 2008). Unlike news and scientific texts, which contain specific formatting such as topic sentences and abstracts, legal cases contain implicit argument structure spreading across long texts (Xu et al., 2021). Current abstractive summarization models do not take into account the argumentative structure of the text, which poses a challenge towards effective abstractive summarization of legal documents.

In this work, we bridge the gap between prior research focusing on summarizing legal documents through extracting argument roles of legal text (Grover et al., 2003; Xu et al., 2021; Saravanan and Ravindran, 2010), and prior research focused on producing abstractive summaries of legal text (Feijo and Moreira, 2019; Bajaj et al., 2021). Our work proposes a technique that *blends argument role mining and abstractive summarization*, which hasn't been explored extensively in the literature.

Figure 1 describes the main flow of our approach, which decomposes the summarization process into two tasks. First, each sentence in the document is

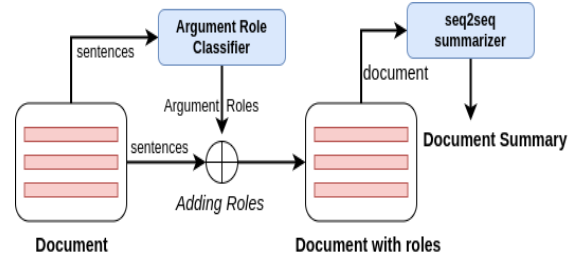


Figure 1: Overview of our approach.

assigned an argument role by using an independent model. Then, the predicted roles are blended with the original document's sentences and fed into a sequence to sequence-based abstractive summarizer.

Our contributions are as follows: **(a)** We propose a simple technique to create an argument-aware neural abstractive summarizer. **(b)** We show the effectiveness of this technique in improving legal document summarization. **(c)** We make our code¹ and argument role annotations freely available².

2 Related Work

Legal Document Summarization. Prior research has mainly focused on extractive techniques (Galgani et al., 2015; Anand and Wagh, 2019; Jain et al., 2021), exploiting features such as the document structure and prior knowledge of legal terms to extract salient sentences that represent the summary of the legal text. Recent research has also shifted gears to abstractive techniques due to their superiority to extractive methods on automatic measures such as ROUGE (Feijo and Moreira, 2019). These abstractive techniques benefited from neural models such as pointer generator networks (See et al., 2017) (utilized in legal public opinion summarization (Huang et al., 2020)) and transformer-based se-

¹<https://github.com/EngSalem/arglegalsumm/>

²The data was obtained through an agreement with the Canadian Legal Information Institute (CanLII) (<https://www.canlii.org/en/>)

quence to sequence encoder-decoder architectures such as BART (Lewis et al., 2020) and Longformer (Beltagy et al., 2020) (employed to summarize long legal documents (Moro and Ragazzi, 2022)). However, the current abstractive approaches ignore the argumentative structure of the legal text. In our work, we combine both the rich argumentative structure of legal documents and state-of-the-art abstractive summarization models.

Argument Mining. Argument mining aims to represent the argumentative structure of a text in a graph structure that contains the argument roles and their relationship to each other. Constructing the graphs usually involves several steps: extracting argument units, classifying the argument roles of the units, and detecting the relationship between different argument roles. Recently, contextualized embeddings were employed to improve argument role labeling (Reimers et al., 2019; Elaraby and Litman, 2021). In many domains, argument roles are classified into *claims*, *major claims*, and *premises* as proposed in Stab and Gurevych (2014). Alternatively, Xu et al. (2021) proposed to classify the argument roles in legal documents to *Issues*, *Reasons*, and *Conclusions* which fits the legal text structure. We use the same set of legal argument role labels in our work, and use contextualized embeddings to automatically predict them.

Argument Mining and Summarization. Prior research integrating argument mining and summarization has mainly focused on extracting chunks of text that contain *key argument units* (Barker et al.; Bar-Haim et al., 2020; Friedman et al., 2021). Combining argument mining and abstractive summarization has received less attention in the literature. Recently, Fabbri et al. (2021) proposed a dialogue summarization dataset with argument information. In their work, the authors included argument information in abstractive summarization by linearizing the argument graph to a textual format, which is used to train an *encoder-decoder* summarization model. However, their proposed approach didn’t improve over *encoder-decoder* baselines. We propose an alternative method that relies on argument roles only, which shows higher improvements over *encoder-decoder* baselines.

3 Dataset and Methods

3.1 Dataset³

Texts. Our dataset is composed of 1049 legal cases

³See Appendix A for more detailed statistics.

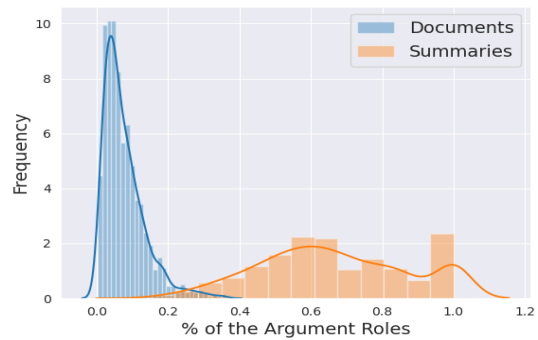


Figure 2: Argumentative sentence % in dataset.

and summary pairs, obtained through an agreement with the Canadian Legal Information Institute. We split these pairs into training (839 pairs, about 80%), validation (106 pairs, about 10%) and testing (104 pairs, about 10%) datasets.

Document Lengths. The maximum length of our input documents is 26k words, which motivates us to include *encoder-decoder* architectures like Longformer that can encode long documents.

Argument Role Annotations. The dataset was annotated prior to our study, using the IRC taxonomy of three legal argument roles described in Xu et al. (2021): **Issues** (legal questions which a court addressed in the document), **Reasons** (pieces of text which indicate why the court reached the specific conclusions), and **Conclusions** (court’s decisions for the corresponding issues). Figure 2 shows the distribution of the percentage of sentences annotated with an argumentative role across the articles and reference human summaries. We can see that while only a small percent of the sentences in the original articles are annotated as argument units, argumentative units dominate the reference summaries. Thus, we hypothesize that augmenting the summarization model with argument roles in the input text should improve the generated summaries.

3.2 Methods

Special Tokens Approach. We designate special marker tokens to distinguish between different argument roles. In prior research, DeYoung et al. (2021) used markers such as `<evidence>`, `</evidence>` to highlight evidence sentences in summarizing medical scientific documents, while Khalifa et al. (2021) used `<neg>`, `</neg>` to mark negation phrases in dialogue summarization. However, we explore the impact of changing token granularity by exper-

Example
<IRC> He also found “on the strong balance of probabilities,” that the late Mrs. Scott intended to make an inter vivos gift to Ms. Akerley. </IRC>
<Issue> [6] Mr. Comeau appeals, arguing that the probate court judge erred: </Issue>

Table 1: Different special tokens for argument roles.

imenting with two sets of special tokens. First, we introduce **<IRC>**, **</IRC>** to distinguish between argumentative and non-argumentative sentences. Second, we broaden the list of the proposed special tokens to differentiate between the three argument roles mentioned in Section 3.1. We assign each argument role two unique tokens to highlight its boundaries in the text, e.g., we use **<Reason>**, **</Reason>** to highlight the **reason** roles. Table 1 shows examples using tokens to highlight an argumentative sentence (top) versus a specific argumentative role (bottom).

Sentence-level Argument Role Mining. Our data’s argument role annotation is at the sentence level, thus, we perform sentence-level classification experiments using the same data splits employed in summarization to detect argument roles.⁴ We experiment with several contextualized embedding-based techniques, namely *BERT* (Devlin et al., 2019), *RoBERTa* (Liu et al., 2019), and *legalBERT* (Zheng et al., 2021). We employ these models to predict sentences’ argument roles (issues, reasons, conclusions, or non-argumentative). Figure 3 shows that *legalBERT* achieved the best classification results. We achieved a macro average F1 of 63.4% in argument role classification and 71% in binary classification using *legalBERT*. Thus, we rely on its predictions to integrate argument roles into summarization below.

4 Experiments and Results

Our experiments are conducted in two settings: assuming argument roles are manually labeled (which we refer to as *oracle*) versus predicting argument role labels (referred to as *predicted*).

4.1 Baselines

We compare our proposed argument-aware summarization method to two sets of baselines⁵:

⁴See Appendix B for argument mining training details.

⁵See Appendix B for summarization training details.

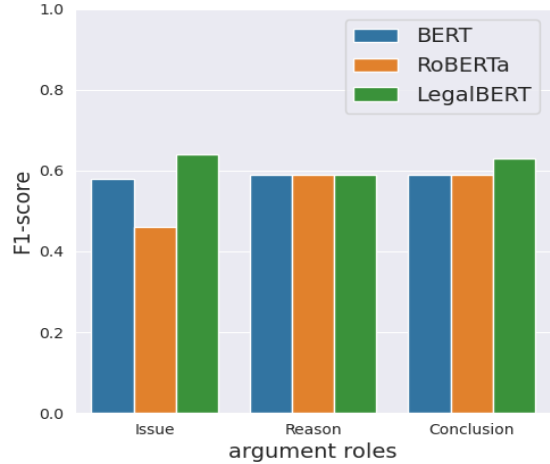


Figure 3: Argument role detection results.

Extractive Baseline. We employ the unsupervised method of Miller (2019). The model leverages BERT embeddings and k-means to extract salient sentences based on their proximity to cluster centroids.

Abstractive Baselines. **Vanilla BART-Large** refers to finetuning BART-large on our dataset. For **Vanilla LED-base**, similarly to BART, the model is finetuned using Longformer-base checkpoint.

4.2 Results and Discussion

Table 2 evaluates the results of the different summarization models using Rouge-1, Rouge-2, and Rouge-L scores.⁶ We refer to BART and Longformer augmented with argument roles as **arg-BART-Large** and **arg-LED-base**, respectively. We use 2 *markers* to denote the use of binary special tokens (i.e; **<IRC>**, **</IRC>**) and 6 *markers* to refer to the full set of argument role tokens. We include two markers sets to examine whether it’s necessary to include explicit argument roles or if it’s sufficient to highlight argumentative text only.

We first evaluate the models augmented with the *manually labeled argument roles* to examine whether adding argument information has the potential to improve over the baselines. The *oracle* results in Table 2 show that **arg-LED-base** improves performance in terms of Rouge-1, Rouge-2, and Rouge-L (Lin, 2004) by approximately 1, 4, and 1.5 points, respectively, over the vanilla LED-base baseline when using the 6 *markers*. The 2 *markers* set showed marginal improvements on Rouge-1 and Rouge-L, but showed 4 Rouge-2 points improvement over the baseline. These results indi-

⁶See Appendix C for example generated summaries.

Setting	Experiment	Model	Rouge-1	Rouge-2	Rouge-L
Oracle	Baselines	Unsupervised Extractive BERT	37.71	14.77	36.41
		Vanilla BART-Large	47.93	22.34	44.74
		Vanilla LED-base	49.56	22.75	46.48
	arg-BART-Large	BART-Large + 2 markers	47.11	21.77	43.12
		BART-Large + 6 markers	46.80	22.14	44.14
	arg-LED-base	LED-base + 2 markers	49.64	26.81	46.70
		LED-base + 6 markers	50.53	26.31	47.90
		LED-base + 2 markers	49.50	<i>26.46</i>	46.60
Predicted	arg-LED-base	LED-base + 6 markers	<i>50.23</i>	<i>26.29</i>	<i>47.49</i>

Table 2: Summarization results on the test set. Best results **bolded**. Best results using predicted roles *italicized*.

Model	Rouge-1	Rouge-2	Rouge-L	Mean Summary Length
Vanilla LED-base	48.25	21.60	44.88	267
arg-LED-base + 2 markers	50.43	27.76	47.05	156
arg-LED-base + 6 markers	50.73	27.29	47.30	174

Table 3: Comparing Longformer (LED) summaries with sentences labeled as argumentative in reference summary.

cate that representing argument roles using fine-grained labels is the most effective in improving LED model output. In contrast, **arg-BART-Large** didn’t show improvements over the vanilla BART-Large baseline. We hypothesize that this is due to the sparsity of the argumentative sentences in the input documents (recall Figure 2). Since Longformer can encode more words, it was likely able to capture more argument markers added to the input, increasing the model’s ability to grasp the argument structure of the legal text. To validate this hypothesis, we analyze the positions of each argument role across the input articles. Figure 4 shows that the argument roles are distributed across the article and not centered around a unique position. This aligns with our hypothesis that the Longformer’s encoding limit (*blue dashed line*) can cover significantly more roles when compared to the BART’s encoding limit (*red dashed line*).

Next, we evaluate the summarization using *predicted argument roles* obtained from our classifier (Section 3.2). We evaluate the Longformer summarization model only, since BART didn’t show oracle improvements. The last two rows of Table 2 (the *predicted* results) show that including predicted argument roles showed consistent improvements with the manually labeled ones (oracle). The results showed a minimal drop in Rouge scores ranging from 0.02 – 0.41 points when using the

predicted argument roles both in the *6 markers* and *2 markers* cases, which indicates the effectiveness of our approach in practical scenarios.

Finally, to estimate the argumentativeness of the LED-based (oracle) summaries, we evaluate them against a summary containing only the sentences manually annotated as an IRC sentence in the reference summary.

Table 3 shows that adding argument role markers increases the overlap between the generated summaries and the argumentative sentence subset from the reference summaries, suggesting that our proposed model’s gains are mainly obtained from an increase in argumentativeness of the generated summaries. The generated summaries from our **arg-LED-base** are lower in length compared to the baseline. We hypothesize that this is due to the focus on argument roles mainly, discarding some of the non-argumentative content.⁷

5 Conclusion and Future Work

We proposed to utilize argument roles in the abstractive summarization of legal documents to accommodate their argumentative structure. Our experiments with state-of-the-art encoder-decoder models showed that including argument role information can improve the ROUGE scores of summarization models capable of handling long

⁷See Appendix C for an illustrative example.

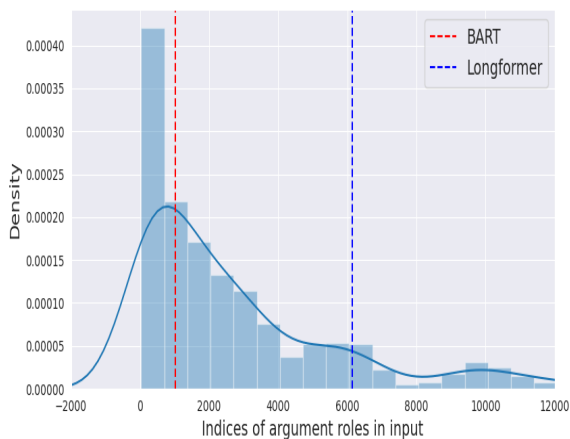


Figure 4: Argument position distribution in the input.

documents. Specifically, improved results were achieved using Longformer with input documents augmented with argument roles (highlighted using special marker tokens), and this finding was robust across two special token schemes. We also showed that using predicted argument roles showed consistent improvements to using the manually labeled ones. In future work, we plan to explore methods to unify argument mining and summarization to reduce the computational resources necessary to host two models.

Acknowledgements

This material is based upon work supported by the National Science Foundation under Grant No. 2040490 and by Amazon. We would like to thank the members of both the Pitt AI Fairness and Law Project and the Pitt PETAL group, as well as the anonymous reviewers, for valuable comments in improving this work.

References

- Deepa Anand and Rupali Wagh. 2019. Effective deep learning approaches for summarization of legal texts. *Journal of King Saud University-Computer and Information Sciences*.
- Ahsaas Bajaj, Pavitra Dangati, Kalpesh Krishna, Pradhiksha Ashok Kumar, Rheeya Uppaal, Bradford Windsor, Eliot Brenner, Dominic Dotterer, Rajarshi Das, and Andrew McCallum. 2021. Long document summarization in a low resource setting using pretrained language models. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing: Student Research Workshop*, pages 71–80.
- Roy Bar-Haim, Yoav Kantor, Lilach Eden, Roni Friedman, Dan Lahav, and Noam Slonim. 2020. Quantitative argument summarization and beyond: Cross-domain key point analysis. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 39–49.
- Emma Barker, Monica Paramita, Adam Funk, Emina Kurtic, Ahmet Aker, Jonathan Foster, Mark Happle, and Robert Gaizauskas. What’s the issue here?: Task-based evaluation of reader comment summarization systems.
- Iz Beltagy, Matthew E Peters, and Arman Cohan. 2020. Longformer: The long-document transformer. *arXiv preprint arXiv:2004.05150*.
- Arman Cohan, Franck Dernoncourt, Doo Soon Kim, Trung Bui, Seokhwan Kim, Walter Chang, and Nazli Goharian. 2018. A discourse-aware attention model for abstractive summarization of long documents. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*, pages 615–621.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186.
- Jay DeYoung, Iz Beltagy, Madeleine van Zuylen, Bailey Kuehl, and Lucy Wang. 2021. Ms²: Multi-document summarization of medical studies. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 7494–7513.
- Mohamed Elaraby and Diane Litman. 2021. Self-trained pretrained language models for evidence detection. In *Proceedings of the 8th Workshop on Argument Mining*, pages 142–147.
- Alexander Richard Fabbri, Faiaz Rahman, Imad Rizvi, Borui Wang, Haoran Li, Yashar Mehdad, and Dragomir Radev. 2021. Convosumm: Conversation summarization benchmark and improved abstractive summarization with argument mining. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 6866–6880.
- Diego Feijo and Viviane Moreira. 2019. Summarizing legal rulings: Comparative experiments. In *proceedings of the international conference on recent advances in natural language processing (RANLP 2019)*, pages 313–322.

- Roni Friedman, Lena Dankin, Yufang Hou, Ranit Aharonov, Yoav Katz, and Noam Slonim. 2021. Overview of the 2021 key point analysis shared task. In *Proceedings of the 8th Workshop on Argument Mining*, pages 154–164.
- Filippo Galgani, Paul Compton, and Achim Hoffmann. 2015. Summarization based on bi-directional citation analysis. *Information processing & management*, 51(1):1–24.
- Claire Grover, Ben Hachey, Ian Hughson, and Chris Korycinski. 2003. Automatic summarisation of legal documents. In *Proceedings of the 9th international conference on Artificial intelligence and law*, pages 243–251.
- Karl Moritz Hermann, Tomas Kocisky, Edward Grefenstette, Lasse Espeholt, Will Kay, Mustafa Suleyman, and Phil Blunsom. 2015. Teaching machines to read and comprehend. *Advances in neural information processing systems*, 28:1693–1701.
- Yuxin Huang, Zhengtao Yu, Junjun Guo, Zhiqiang Yu, and Yantuan Xian. 2020. Legal public opinion news abstractive summarization by incorporating topic information. *International Journal of Machine Learning and Cybernetics*, 11(9):2039–2050.
- Deepali Jain, Malaya Dutta Borah, and Anupam Biswas. 2021. Automatic summarization of legal bills: A comparative analysis of classical extractive approaches. In *2021 International Conference on Computing, Communication, and Intelligent Systems (ICCCIS)*, pages 394–400. IEEE.
- Muhammad Khalifa, Miguel Ballesteros, and Kathleen Mckeown. 2021. A bag of tricks for dialogue summarization. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 8014–8022.
- Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Veselin Stoyanov, and Luke Zettlemoyer. 2020. Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 7871–7880.
- Chin-Yew Lin. 2004. Rouge: A package for automatic evaluation of summaries. In *Text summarization branches out*, pages 74–81.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. Roberta: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*.
- Derek Miller. 2019. Leveraging bert for extractive text summarization on lectures. *arXiv preprint arXiv:1906.04165*.
- Gianluca Moro and Luca Ragazzi. 2022. Semantic self-segmentation for abstractive summarization of long legal documents in low-resource regimes. In *Proceedings of the Thirty-Six AAAI Conference on Artificial Intelligence, Virtual*, volume 22.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21:1–67.
- Nils Reimers, Benjamin Schiller, Tilman Beck, Johannes Daxenberger, Christian Stab, and Iryna Gurevych. 2019. Classification and clustering of arguments with contextualized word embeddings. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 567–578.
- Evan Sandhaus. 2008. The new york times annotated corpus. *Linguistic Data Consortium, Philadelphia*, 6(12):e26752.
- M Saravanan and Balaraman Ravindran. 2010. Identification of rhetorical roles for segmentation and summarization of a legal judgment. *Artificial Intelligence and Law*, 18(1):45–76.
- Abigail See, Peter J Liu, and Christopher D Manning. 2017. Get to the point: Summarization with pointer-generator networks. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1073–1083.
- Christian Stab and Iryna Gurevych. 2014. Identifying argumentative discourse structures in persuasive essays. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 46–56.
- Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, et al. 2019. Huggingface’s transformers: State-of-the-art natural language processing. *arXiv preprint arXiv:1910.03771*.
- Huihui Xu, Jaromir Savelka, and Kevin D Ashley. 2021. Toward summarizing case decisions via extracting argument issues, reasons, and conclusions. In *Proceedings of the eighteenth international conference on artificial intelligence and law*, pages 250–254.
- Jingqing Zhang, Yao Zhao, Mohammad Saleh, and Peter Liu. 2020. Pegasus: Pre-training with extracted gap-sentences for abstractive summarization. In *International Conference on Machine Learning*, pages 11328–11339. PMLR.
- Lucia Zheng, Neel Guha, Brandon R. Anderson, Peter Henderson, and Daniel E. Ho. 2021. [When does pretraining help? assessing self-supervised learning for law and the casehold dataset](#). In *Proceedings*

A Data statistics

A.1 Length statistics

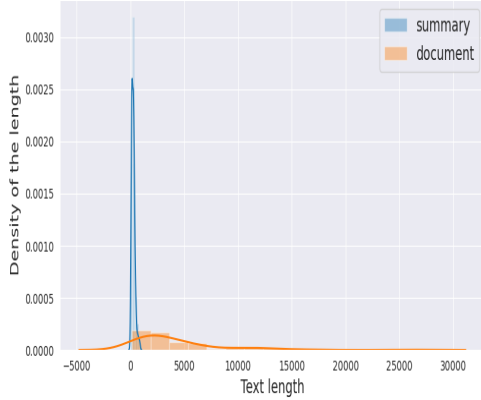


Figure 5: Distribution of article and summary length.

Figure 5 shows the distribution of document and summary lengths. The summaries’ lengths are centered around a mean of 255 words, with a maximum length of 850 words. The 90th percentile of summary length is 490 words. Thus we chose the maximum length of generated summary in our hyperparameters to be set to 512 words. Unlike the summaries, the documents are more spread across the distribution. In our analysis, we found that the mean document length is 4180 words, while the maximum document length is 26235 words.

A.2 Argument role distribution

While they are essential to legal cases, argument roles represent a small percentage of the document. Figure 6 shows the high imbalance of the manually annotated argumentative versus non-argumentative sentences in our training set, which poses a challenge in building a sentence level classifier of argument roles. In our analysis we found that the non-argumentative sentences count is approximately $1000\times$ the argumentative sentences, which we use to adjust class weights in our learning objective.

B Training details and hyperparameters

All experiments use the model implementations provided in the *Huggingface library* (Wolf et al., 2019). We initialize all our models with the same

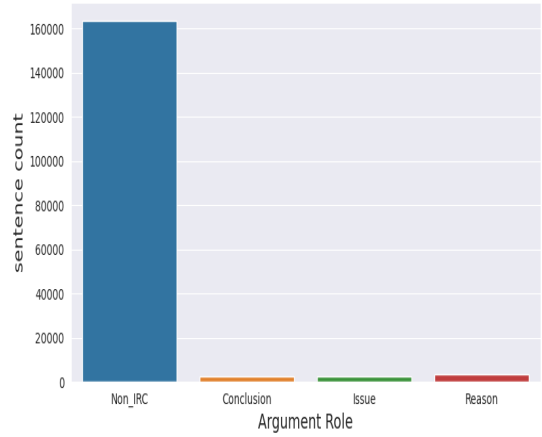


Figure 6: Count of argument roles across training set. *non_IRC* refers to non-argumentative sentences.

learning rate of $2e^{-5}$. We train both our summarization and argument role classification models for 10 epochs with early stopping with 3 epoch patience. For training summarization models, we set the maximum summary length to 512 words. We truncate the input length to 1024 words for the *BART* model while truncating the input length to 6144 words for the *Longformer* due to our GPU limitation⁸. We pick our best model based on its *ROUGE-2* (Lin, 2004) score on the validation set. For the classification models introduced in Section 3.2, due to the high imbalance of our argumentative labels, we introduce fixed class weights to our cross-entropy loss. For argumentative sentences, we modify the *cross-entropy* weight to be 1000 compared to 1 for non-argumentative sentences. We chose these weights based on label distribution in our training set described in Figure 6. Our best model is chosen based on the F1 score on the validation set.

C Effect of argument roles on generated summaries

Table 4 shows an example of generated summaries with adding special tokens and without the special tokens.

⁸We use Quadro RTX 5000 which has 16 GB RAM.

Reference summary	Vanilla LED-base	arg-LED-base "6 markers"
The plaintiffs claim for crop damage caused by the defendant's cattle entering the plaintiff's canola field. The defendant denies he is responsible or negligent, but admits his cattle were in the field. HELD: Damages were awarded in the amount of \$2,533.45. The court valued the loss at \$3,052.36. It then deducted 2% attributed to wild animal damage, 5% for the plaintiffs failure to calculate actual yield from the rest of the crop, 5% that the plaintiffs would have paid in dockage and 5% for the cost of production. Civil liability for crop damage caused by cattle flows from the Stray Animals Act. A cattle owner is strictly liable for damages caused by his straying cattle. This liability may possibly even be absolute and only an act of god may serve as a defence. The issue here is the quantum of the damage. The fact that the defendant was refused access to the plaintiff's property to repair the fence has no bearing on liability in this case. An independent adjuster assessed the damage at 557 bushels. The court found on the facts that the damage was caused exclusively by the defendant's cattle on several occasions.	The plaintiff's claim against the defendant is for crop damage sustained as a result of the defendant's cattle getting into the plaintiff's canola field in late July 1999. The plaintiff admits that his cattle were in the Plaintiff's crop on one occasion, but says that there is no evidence that their cattle are in the canola on any other occasion, and that such loss as was caused by his cattle is minimal. HELD: Judgment in the amount of \$2,533.45 was granted in favour of the plaintiffs, together with pre-judgment interest from January 2000 and costs. The present s.27 (1) has been in the Act since 1977, when the Act was substantially re-written, apparently in response to the increasing urbanization of the province. Under both the present and pre-1977 statutes, the legal situation was the same – the cattle owner is strictly liable for damage caused by the straying cattle. The fact remains that the cattle did move onto the field unbeknownst to the defendant, his spouse, his brother, or his hired help. The Carrot River did not impede the cattle from moving into the field on more than one occasion and other cattle were not in the field.	The plaintiff's claim against the defendant was for crop damage sustained as a result of his cattle getting into the plaintiff's canola field in late July 1999. The defendant admits that his cattle were in the Plaintiff's crop on one occasion, but says that there is no evidence that any other cattle were there. HELD: Judgment in the amount of \$2,533.45 was granted in favour of the plaintiffs, together with pre-judgment interest from January 2000 and costs. Under both the present s.27 (1) and the pre-1977 statutes, the legal situation was the same – the cattle owner is strictly liable for damage caused by his straying cattle.

Table 4: Example of generated summaries with **Vanilla LED-base** and **arg-LED-base** versus reference summary.