

FedCSIS_2023_paper_6516.pdf

Classification of US Supreme Court Cases using BERT-Based Techniques

Shubham Vatsal

New York University, CIMS

New York, USA

Email: sv2128@nyu.edu

Adam Meyers

New York University, CIMS

New York, USA

Email: meyers@cs.nyu.edu

John E. Ortega

New York University, CIMS

New York, USA

Email: jortega@cs.nyu.edu

1

Abstract—Models based on bidirectional encoder representations from transformers (BERT) produce state of the art (SOTA) results on many natural language processing (NLP) tasks such as named entity recognition (NER), part-of-speech (POS) tagging etc. An interesting phenomenon occurs when classifying long documents such as those from the US supreme court where BERT-based models can be considered difficult to use on a first-pass or out-of-the-box basis. In this paper, we experiment with several BERT-based classification techniques for US supreme court decisions or supreme court database (SCDB) and compare them with the previous SOTA results. We then compare our results specifically with SOTA models for long documents. We compare our results for two classification tasks: (1) a broad classification task with 15 categories and (2) a fine-grained classification task with 279 categories. Our best result produces an accuracy of 80% on the 15 broad categories and 60% on the fine-grained 279 categories which marks an improvement of 8% and 28% respectively from previously reported SOTA results.

I. INTRODUCTION

Every October, the US supreme court begins a new term on the first Monday. Each term includes a number of significant and complex cases that address a variety of issues, including environmental protection law, free speech, equal opportunity, tax law etc. Legal court decisions are lawful statements acknowledged by a judge providing the justification and reasoning for a court ruling. The court's decisions not only have repercussions on the people involved in the case but also on the society as a whole. During a regularly scheduled court session, the justice who wrote the main decision sums it up from the bench. A copy of the opinion is then quickly posted on this website. Legal experts must search through and classify these court decisions to support their research. This can require a lot of manual effort. Artificial intelligence (AI) and machine learning (ML) can significantly reduce the burden of this type of manual work. Washington University's collection of 8419 manually labeled supreme court documents (SCDB) [1] provides the basis for bench marking this type of AI task.

There are multiple challenges involved when dealing with SCDB or legal documents of similar kinds. First, these documents are exceptionally long which causes numerous difficulties. For example, feature vectors can be too large to fit in the memory or contextual variance can be too high for models to handle. Next, the decision of a case requires identification of other relevant cases that can support the decision of the

current case, which usually involve similar circumstances. Therefore, it is of utmost importance to identify the similarities in terms of legal aspects while classifying the cases in the same category. Finally, understanding of these legal documents requires supervision of legal expert. Legal statements, like other specialized domains, has unique characteristics when compared to other generic corpora, such as unique vocabulary, exclusive syntax, idiosyncratic semantics etc. Vocabulary, syntax, semantics and other linguistic characteristics can be specific to the legal domain or even the subdomain of court decisions, thus necessitating tools that are trained specifically for such domains. Hence, an automatic system to categorize and process these documents is extremely useful.

In this paper, we explore many novel techniques to counter the problem faced by models similar to BERT [2] when dealing with documents of length more than 512 tokens. Some of the techniques include: analysing which 512 token chunks of documents makes the best contribution towards classification, using summarized version of documents for classification, a voting-based ensemble approach, classification based on concatenation of 512 token chunks. We compare our results with previous non-BERT like models as well as models which can take inputs of length more than 512 tokens.

The rest of the paper is organized in the following way. Section 2 talks about related works. Section 3 talks about the SCDB. We describe the working of BERT-based techniques in Section 4. Section 5 concentrates on the experiments we conducted and the corresponding results we achieved. The final section discusses future work.

II. RELATED WORK

Recently, there has been a surge in research in the domain of NLP associated with the legal documents. [3] talks about combining linguistic information from WordNet with a syntax-based retrieval of rules from legal text along with logic-based retrieval of dependencies from chunks of such texts leading to extraction of machine-readable rules from legal documents. [4] illustrates several embedding-based and symbol-based approaches for judgement prediction, legal question answering and similar case matching. [5] discusses five areas of legal activity where NLP is playing an important role. These areas include legal research for finding information relevant to a legal decision, electronic discovery determining the relevance

of documents in an information request, contract review to check that a contract is complete, document automation to generate routine legal documents and legal advice using question-answering dialogues. [6] discusses different available approaches for summarization of legal texts and compares their performances on various datasets. [7] showcases a scalable and flexible information extraction method, aimed at extraction of information from legal documents regardless of format, layout or structure, by considering the context. [8] comes up with a German legal BERT model and evaluates its performance on downstream NLP tasks including classification, regression and similarity.

BERT-based models have shown some ground breaking performance in many NLP tasks. Nowadays, they are being widely used in the legal domain as well. [9] proposes BERT-PLI to capture the semantic relationships at the paragraph-level and then goes on to infer the relevance between two legal cases by aggregating paragraph-level interactions. [10] releases Legal-BERT, a family of BERT models for the legal domain intended to assist legal NLP research. [11] studies a case in the context of legal professional search and presents how BERT-based approach outperforms other traditional approaches. [12] proposes an answer selection approach by fine-tuning BERT on their Vietnamese legal question-answer pair corpus. They further pre-train BERT on a Vietnamese legal domain-specific corpus and show that this new BERT performs better than the fine-tuned BERT.

Classification of legal documents is an important NLP task which can automate alignment of legal documents with human-defined categories. [13] classifies a proprietary corpus consisting of hundreds of thousands of legal agreements using BERT. [14] compares multi-label and binary classification of legal documents using variances of pre-trained BERT-based models and other approaches to handle long documents. Our work falls somewhat along similar lines but we use a different dataset of legal documents with considerably different properties. Moreover, many of our BERT-based techniques are significantly different from [14]. [15] presents baseline results for document type classification and theme assignment, a multi-label problem using their newly built Brazil’s supreme court digitalized legal documents dataset. [16] proposes a method for learning Chinese legal document classification using graph long short-term memory (LSTM) combined with domain knowledge extraction. [17] addresses multi-label classification of Croatian legal documents using EuroVoc thesaurus. [18] experimented classification of Singapore supreme court judgments using topic models, word embedding feature models and pre-trained language models. [19] presents results on predicting the rulings of the Turkish Constitutional Court and Courts of Appeal using fact descriptions. [20] investigates various text classification techniques to predict French Supreme Court decisions whereas [21] does similar work for Philippine Supreme Court.

SCDB has been used in various prominent NLP tasks. [22] uses supreme court data and performs topic modelling using domain-specific embeddings. These embeddings are obtained

TABLE I
DATA STATISTICS

Metric	Value
Dataset Size	8419
Min # Tokens	0
Max # Tokens	87246
Median # Tokens	5552
Mean # Tokens	6960.60
Min # Tokens After Pre-Processing	0
Max # Tokens After Pre-Processing	87246
Median # Tokens After Pre-Processing	3420
Mean # Tokens After Pre-Processing	4458.15

1 from pre-trained Legal-BERT. [23] constructs a model to predict the voting behavior of US supreme court and its justices in a generalized, out-of-sample context by using SCDB along with some other derived features. [24] talks about classification performance of different BERT-based as well as non-neural architectures on many datasets of legal domain including SCDB. Our experiments differ from their work in multiple ways. First, they don’t analyze the performance of these models across fine-grained 279 categories which is a harder classification task. Second, they do not experiment different techniques to tackle the problem of restricted input sequence length of 512 tokens in BERT-based models which is one of the key points of our work. The analysis of these techniques helps us in achieving SOTA across both broad as well as fine-grained classification tasks. Finally, the version of SCDB used by them differs from what we have used in our experiments. There is a significant overlap but the versions are not exactly the same. [25] presents classification of SCDB across broad 15 categories and fine-grained 279 categories. Our work exactly aligns with this work but the usage of BERT-based techniques helps us in achieving better results.

III. DATA

Our paper is primarily based on classification of US supreme court decisions dataset or supreme court database from Washington University School of Law [1]. Documents are classified by topic in a 2 level ontology, providing the basis for two different classification tasks: one using 15 broad category labels and another using 279 fine-grained category labels. The general statistics associated with this dataset can be found in Table I. As we can see from Fig. 1 and Fig. 2 of [25], SCDB is highly imbalanced in terms of number of data points per label in both classification tasks. Before we conduct our experiments, we apply a pre-processing step to remove footnotes¹. The SCDB tasks pose some difficulties for BERT-based classification because the average length of SCDB documents is much larger than most other legal document datasets used for classification task. For example, European Court of Human Rights (ECHR) [26] and Overruling [27] datasets have only 1662.08 and 21.94 mean length in comparison to SCDB’s mean length of 6960.60 tokens.

¹We provisionally assume that footnotes constitute noise.

TABLE II
RESULTS FOR BEST-512

Chunk I	Labels	Accuracy	Precision	F1
Chunk 1	15	0.747	0.752	0.744
	279	0.545	0.498	0.500
Chunk 2	15	0.688	0.692	0.683
	279	0.464	0.417	0.419
Chunk 3	15	0.683	0.682	0.673
	279	0.457	0.408	0.409
Chunk 4	15	0.704	0.702	0.696
	279	0.459	0.405	0.412
Chunk 5	15	0.687	0.685	0.682
	279	0.452	0.404	0.406
Chunk 6	15	0.679	0.689	0.676
	279	0.454	0.411	0.409

TABLE III
RESULTS FOR STRIDE-64, 128

Stride	Labels	Accuracy	Precision	F1
64	15	0.774	0.777	0.771
	279	0.563	0.514	0.519
128	15	0.762	0.779	0.763
	279	0.557	0.505	0.510

IV. PROPOSED TECHNIQUES

We explore various BERT-based techniques to classify SCDB in this section. We apply these individual techniques to both the classification tasks i.e one with 15 categories and the other one with 279 categories. We use either BERT or different versions of BERT for our classification tasks. Particularly, we use BERT [2], RoBERTa [28] and Legal-BERT [29] which have the restriction of maximum input sequence length of 512 tokens. We first try out our different techniques using BERT, RoBERTa and Legal-BERT. Later, we compare the results of these techniques with other transformer-based models like LongFormer [30] and Legal-Longformer² which accept longer sequences.

A. Best-512

Let c_i represent the i^{th} 512-length chunk of a given document in SCDB. We have taken the length of each chunk to be 512 because that is maximum length of sequence accepted by a BERT-based model. We calculate our evaluation metrics on these chunks and thus analyse which chunk of documents contribute the most towards their correct classification. Since the median length of documents in SCDB is around 3000, we compute the performance of our three BERT-based models on c_1, c_2, c_3, c_4, c_5 and c_6 where c_1 represents 1st 512 length chunk of documents, c_2 represents the 2nd 512 length chunk of documents and so on. When the length of a document is less than $i \times 512$, we take the last 512 or less than 512 (when $i=1$) tokens of the document. Initially, we use BERT to find the best chunk c_i for all the documents and later use the same c_i to experiment with different BERT-based models discussed in Section V. The result showing the best chunk using BERT can be seen in Table II. The final result comparing the performance of different BERT-based models using the best chunk obtained

from II for both the classification tasks can be seen in Table IV and Table V.

B. Summarization-512

In this technique, we summarize the documents of SCDB in 512 tokens. We used the summarization pipeline from Hugging Face with default parameters. The maximum sequence length that this summarization model can accept is 1024. So, we first convert a document to some splits based on the length of that document. The number of splits n_i for a given document d_i was defined as $l_i/1024$ where l_i was the length of the document. Now, since the total length of the summarized version of a document can only be 512 tokens long, we further calculate the number of tokens per split nw_i for all the splits of a given document d_i . Finally, we concatenate all nw_i 's of a given document d_i to get the final summarized version. These summarized versions are then used for both the classification tasks. The results are shown in Table IV and Table V.

C. Concat-512

Let c_i represent the i^{th} 512-length chunk of a given document in SCDB. In this technique, we accept i parallel inputs of 512 sequence length. Corresponding to i parallel inputs we have i BERT-based models which are trained simultaneously and their outputs are concatenated. In a sense, we concatenate i CLS tokens from i BERT-based models. This concatenated output is then fed into a dense layer with softmax activation and number of units being equal to 15 or 279 based on the classification task. Again, because the median length of documents is around 3000, for this experiment we only take c_1, c_3, c_4, c_5 and c_6 into consideration where c_1 represents 1st 512 length chunk of documents, c_2 represents the 2nd 512 length chunk of documents and so on. The results comparing different BERT-based models can be seen in Table IV and Table V.

²<https://huggingface.co/saibo/legal-longformer-base-4096>

D. Ensemble

Let c_i represent the i^{th} 512-length chunk of a given document in SCDB. In our ensemble approach, we train a model m_i for each c_i . During testing, we predict the final label of a document using a maximum voting mechanism where the final prediction is what the majority of m_i end up choosing which is given by equation 1. The $m_{i,t}$ term in equation 1 can take either the value of 0 or 1 based on its prediction. If i^{th} classifier chooses class t , then $m_{i,t} = 1$, and 0, otherwise. The $\#nc$ term in equation 1 refers to the number of classes for the corresponding classification task. In the case where the length of a document is less than $i*512$, we ignore that document for the training of that m_i . Since the median length of documents in SCDB is around 3000, for this experiment we only take c_1, c_2, c_3, c_4, c_5 and c_6 into consideration where c_1 represents 1st 512 length chunk of documents, c_2 represents the 2nd 512 length chunk of documents and so on. As stated previously, we run this experiment on different BERT-based models and note down the results for both the classification tasks. The results can be seen in Table IV and Table V.

$$label = \underset{t \in \{1, 2, \dots, \#nc\}}{\operatorname{argmax}} \sum_{i=1}^6 m_{i,t} \quad (1)$$

E. Stride-64, 128

Let c_i represent the i^{th} 512-length chunk of a given document in SCDB. Stride basically represents the chunk of tokens which are shared amongst any two c_i and c_{i+1} . Let's take an example with stride value being 64 to develop more clarity about how strides are used here. Let $c_i[0 : 512]$ represent the 512 tokens present in c_i . Then, $c_1[448 : 512]$ and $c_2[0 : 64]$ tokens of c_1 and c_2 respectively are exactly same. So, if we have a document d_i of length 1024 tokens, we will have three c_i 's with $c_1[448 : 512] = c_2[0 : 64]$ and $c_2[448 : 512] = c_3[0 : 64]$. Also, to elaborate $d_i[0 : 512] = c_1[0 : 512]$, $d_i[448 : 512] = c_2[0 : 64]$, $d_i[512 : 960] = c_2[64 : 512]$, $d_i[896 : 960] = c_3[0 : 64]$ and $d_i[960 : 1024] = c_3[64 : 128]$. We have pad tokens in $c_3[128 : 512]$. Taking the median length of documents of SCDB into consideration, we again experiment up till length around 3000 tokens as explained for other techniques previously. Initially, we use BERT model to find the best stride for all the documents and later used the same stride to experiment with different BERT-based models discussed in Section V. The result showing the best stride using BERT can be seen in Table III. The final result comparing the performance of different BERT-based models using the best stride obtained from III for both the classification tasks can be seen in Table IV and Table V.

F. Longer Sequence Model (LSM)

These are the models which can accept input sequence larger than 512 tokens. We ran our experiments with two such models, LongFormer and Legal-Longformer. Apart from these two models, we also used the results reported by [25] where the best performing model used convolutional neural network (CNN) architecture.

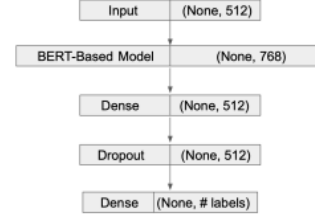


Fig. 1. Best-512, Summarization-512, Ensemble, LSMs General Architecture

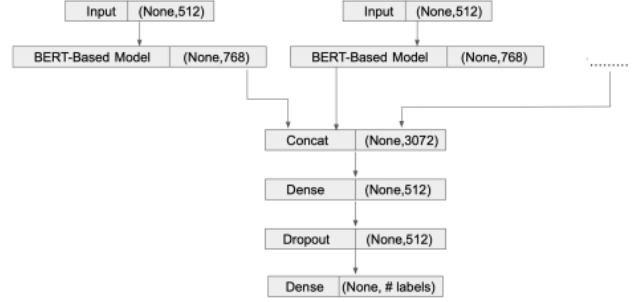


Fig. 2. Concat-512, Stride-64 General Architecture

V. EXPERIMENTS & RESULTS

The code³ related to all the experiments discussed below have been made public. We use weighted F1, accuracy and weighted precision as our evaluation metrics for both the classification tasks. The hyper-parameters used for all the above techniques except RoBERTa include batch size to be 8, number of epochs as 5, learning rate to be $3e-5$ and loss to be Categorical Cross Entropy. For RoBERTa, we keep all the hyper-parameters to be the same except the learning rate which is changed to $1e-5$. We split 8419 data points in 90 to 10 ratio of train and test sets. We ran each experiment 5 times and took the average of the best epoch score to get the final score. We chose Adam [31] as our optimizer and Mean Squared Error (MSE) as our loss function. The accuracy is calculated by rounding off the regressional output to the nearest integer. The labels for an app and its negative data point is given as -1 whereas an app and its positive data point is given as 1 during the training process. We used bert-base-uncased version of BERT, roberta-base version of RoBERTa, legal-bert-base-uncased version of Legal-BERT, longformer-base-4096 version of Longformer and legal-longformer-base-4096 version of Legal-Longformer from Hugging Face⁴. All the BERT-based models accept a sequence of maximum of 512 tokens whereas Longformer and Legal-Longformer accept a sequence of maximum of 4096 tokens. Figure 1 shows the gen-

³Web-Of-Law-Experiments
⁴<https://github.com/huggingface/transformers>

TABLE IV
RESULTS FOR 15 CATEGORIES

Technique	Model	Accuracy	Precision	F1
Best-512	BERT	0.747	0.752	0.744
	RoBERTa	0.768	0.773	0.766
	Legal-BERT	0.785	0.795	0.785
Summarization-512	BERT	0.736	0.748	0.734
	RoBERTa	0.745	0.761	0.747
	Legal-BERT	0.789	0.801	0.790
Concat-512	BERT	0.772	0.775	0.769
	RoBERTa	0.772	0.782	0.773
	Legal-BERT	0.791	0.799	0.791
Ensemble	BERT	0.755	0.755	0.752
	RoBERTa	0.766	0.770	0.763
	Legal-BERT	0.782	0.792	0.782
Stride-64	BERT	0.774	0.777	0.771
	RoBERTa	0.779	0.785	0.778
	Legal-BERT	0.801	0.805	0.800
LSMs	LongFormer	0.742	0.753	0.739
	Legal-LongFormer	0.775	0.785	0.775
	CNN [25]	0.724	-	-

TABLE V
RESULTS FOR 279 CATEGORIES

Technique	Model	Accuracy	Precision	F1
Best-512	BERT	0.545	0.498	0.500
	RoBERTa	0.533	0.474	0.480
	Legal-BERT	0.586	0.554	0.547
Summarization-512	BERT	0.529	0.486	0.483
	RoBERTa	0.522	0.456	0.466
	Legal-BERT	0.585	0.553	0.549
Concat-512	BERT	0.554	0.511	0.511
	RoBERTa	0.534	0.460	0.475
	Legal-BERT	0.596	0.560	0.559
Ensemble	BERT	0.520	0.464	0.471
	RoBERTa	0.520	0.455	0.467
	Legal-BERT	0.553	0.529	0.520
Stride-64	BERT	0.563	0.514	0.519
	RoBERTa	0.536	0.465	0.479
	Legal-BERT	0.609	0.584	0.575
LSMs	LongFormer	0.534	0.481	0.487
	Legal-LongFormer	0.562	0.515	0.519
	CNN [25]	0.319	-	-

eral architecture of Best-512, Summarization-512, Ensemble and LSMs with some differences in corresponding dimensions and input type. Similarly, the general architecture of Concat-512 and Stride-64 with some differences in corresponding dimensions and input type can be seen in Figure 2. The architecture of both the types of models is mostly similar except for the form in which they accept their inputs.

As we can see from Table IV and V a BERT-based model which has been trained on legal data like Legal-BERT or other transformer-based model with it's training data coming from legal domain like Legal-Longformer always outperforms other models within a given technique. When comparing the best performing model across different BERT-based techniques for 15 categories, the techniques can be ranked as Stride-64 giving the best result, followed by Concat-512, followed by Summarization-512, followed by Best-512, followed by Ensemble and finally we have LSMs. Similarly, when comparing the best performing model across different BERT-based techniques for 279 categories, the techniques can be

ranked as Stride-64 giving the best result, followed by Concat-512, followed by Summarization-512, followed by Best-512, followed by Ensemble and finally we have LSMs. There could be multiple reasons for LSMs to not have a better performance than other models. One, LSMs are designed in a way to allow multi-head attentions to adhere to a restricted window contrary to BERT-based models where these multi-head attentions are free to concentrate on any of the tokens. Second, with more number of tokens, more variance is created which can lead to poor performance. Also we have already seen from Table II it is just the first 512 tokens which contribute the most towards the classification of the corresponding documents. The reason for the poor performance of Best-512 can be attributed to the fact that even though it is the first 512 token chunk which contributes the most for these classification tasks but still context beyond these 512 tokens does make an impact. Summarization-512 tries to capture the context across the entire document but due to its limitation to express this context in just 512 tokens, it is not as efficient as Concat-512,

Stride-64 or Ensemble. Ensemble outperforms Best-512 and Summarization-512 because it tries to exploit joint learning across multiple 512 token chunks through its maximum voting mechanism rule. Concat-512 on the other hand captures better context across multiple 512 token chunks as it learns this knowledge during back propagation of the model. Finally, Stride-64 outperforms Concat-512 because when we take disjoint chunks, the continuity of context goes missing whereas if there is an overlapping portion of text between two consecutive chunks, it gives better context understanding.

VI. CONCLUSION & FUTURE WORK

In this paper, we experimented with various BERT-based techniques on SCDB and presented the corresponding results comparing them with other state of the art models. We further did an analysis on how even with the given restriction of the input sequence length of 512 tokens for BERT-based models, we can leverage these techniques to get some improvement.

As a part of future work, we can leverage the knowledge embedded in references of a given SCDB document. A reference in an SCDB document basically refers to some other SCDB document that has been cited to legally justify the decision taken on the former SCDB document. The raw text of these references may not be very helpful in improving the classification tasks. We can use a graph structure to denote the relations between a given SCDB document with other documents cited in it. The final classification result of a given SCDB document can be calculated based on some form of aggregation incorporating classification results of its references weighted by the graph structure representing quantified relations with the SCDB document at hand.

Another area of future work can be applying greedy approaches to the techniques discussed in this work. For example, we can have a greedy summarization technique where in the final 512 token summary, we can include more number of tokens from the best performing 512 token chunk as inferred from Best-512 technique. Similarly, we can have greedy ensemble technique where during the voting phase, we can give more weight to the best performing 512 token chunk as the results show from Best-512 technique.

REFERENCES

- [1] H. J. Spaeth, L. Epstein, A. D. Martin, J. A. Segal, T. J. Ruger, and S. C. Benesh, "Supreme court database, version 2013 release 01," *Database at <http://supremecourtdatabase.org>*, 2013.
- [2] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," *arXiv preprint arXiv:1810.04805*, 2018.
- [3] M. Dragoni, S. Villata, W. Rizzi, and G. Governatori, "Combining nlp approaches for rule extraction from legal documents," in *1st Workshop on Mining and Reasoning with Legal texts (MIREL 2016)*, 2016.
- [4] H. Zhong, C. Xiao, C. Tu, T. Zhang, Z. Liu, and M. Sun, "How does nlp benefit legal system: A summary of legal artificial intelligence," *arXiv preprint arXiv:2004.12158*, 2020.
- [5] R. Dale, "Law and word order: Nlp in legal tech," *Natural Language Engineering*, vol. 25, no. 1, pp. 211–217, 2019.
- [6] A. Kanapala, S. Pal, and R. Pamula, "Text summarization from legal documents: a survey," *Artificial Intelligence Review*, vol. 51, no. 3, pp. 371–402, 2019.
- [7] M. García-Constantino, K. Atkinson, D. Bollegala, K. Chapman, F. Coenen, C. Roberts, and K. Robson, "Ciel: context-based information extraction from commercial law documents," in *Proceedings of the 16th edition of the International Conference on Artificial Intelligence and Law*, 2017, pp. 79–87.
- [8] C. M. Yeung, "Effects of inserting domain vocabulary and fine-tuning bert for german legal language," Master's thesis, University of Twente, 2019.
- [9] Y. Shao, J. Mao, Y. Liu, W. Ma, K. Satoh, M. Zhang, and S. Ma, "Bert-pli: Modeling paragraph-level interactions for legal case retrieval," in *IJCAI*, 2020, pp. 3501–3507.
- [10] I. Chalkidis, M. Fergadiotis, P. Malakasiotis, N. Aletras, and I. Androutsopoulos, "Legal-bert: The muppets straight out of law school," *arXiv preprint arXiv:2010.02559*, 2020.
- [11] L. Sanchez, J. He, J. Manotumruksa, D. Albakour, M. Martinez, and A. Lipani, "Easing legal news monitoring with learning to rank and bert," in *European Conference on Information Retrieval*. Springer, 2020, pp. 336–343.
- [12] C.-N. Chau, T.-S. Nguyen, and L.-M. Nguyen, "Vnlawbert: A vietnamese legal answer selection approach using bert language model," in *2020 7th NAFOSTED Conference on Information and Computer Science (NICS)*. IEEE, 2020, pp. 298–301.
- [13] E. Elwany, D. Moore, and G. Oberoi, "Bert goes to law school: Quantifying the competitive advantage of access to large legal corpora in contract understanding," *arXiv preprint arXiv:1911.00473*, 2019.
- [14] N. Limsopatham, "Effectively leveraging bert for legal document classification," in *Proceedings of the Natural Language Processing Workshop 2021*, 2021, pp. 210–216.
- [15] P. H. L. De Araujo, T. E. de Campos, F. A. Braz, and N. C. da Silva, "Victor: a dataset for brazilian legal documents classification," in *Proceedings of the 12th Language Resources and Evaluation Conference*, 2020, pp. 1449–1458.
- [16] G. Li, Z. Wang, and Y. Ma, "Combining domain knowledge extraction with graph long short-term memory for learning classification of chinese legal documents," *IEEE Access*, vol. 7, pp. 139 616–139 627, 2019.
- [17] F. Šarić, B. Dalbelo Bašić, M.-F. Moens, and J. Šnajder, "Multi-label classification of croatian legal documents using eurovoc thesaurus," in *Proceedings of SPLet-Semantic processing of legal texts: Legal resources and access to law workshop*. ELRA, 2014.
- [18] J. S. T. Howe, L. H. Khang, and I. E. Chai, "Legal area classification: A comparative study of text classifiers on singapore supreme court judgments," *arXiv preprint arXiv:1904.06470*, 2019.
- [19] E. Mumcuoglu, C. E. Öztürk, H. M. Ozaktas, and A. Koç, "Natural language processing in law: Prediction of outcomes in the higher courts of turkey," *Information Processing & Management*, vol. 58, no. 5, p. 102684, 2021.
- [20] O.-M. Sulea, M. Zampieri, M. Vela, and J. Van Genabith, "Predicting the law area and decisions of french supreme court cases," *arXiv preprint arXiv:1708.01681*, 2017.
- [21] M. B. L. Virtucio, J. A. Aborot, J. K. C. Abonita, R. S. Avinante, R. J. B. Copino, M. P. Neverida, V. O. Osiana, E. C. Peramo, J. G. Syjuco, and G. B. A. Tan, "Predicting decisions of the philippine supreme court using natural language processing and machine learning," in *2018 IEEE 42nd annual computer software and applications conference (COMPSAC)*, vol. 2. IEEE, 2018, pp. 130–135.
- [22] R. Silveira, C. Fernandes, J. A. M. Neto, V. Furtado, and J. E. Pimentel Filho, "Topic modelling of legal documents via legal-bert," *Proceedings <http://ceur-ws.org> ISSN*, vol. 1613, p. 0073, 2021.
- [23] D. M. Katz, M. J. Bommarito, and J. Blackman, "A general approach for predicting the behavior of the supreme court of the united states," *PloS one*, vol. 12, no. 4, p. e0174698, 2017.
- [24] I. Chalkidis, A. Jana, D. Hartung, M. Bommarito, I. Androutsopoulos, D. M. Katz, and N. Aletras, "Lexglue: A benchmark dataset for legal language understanding in english," *arXiv preprint arXiv:2110.00976*, 2021.
- [25] S. Undavia, A. Meyers, and J. E. Ortega, "A comparative study of classifying legal documents with neural networks," in *2018 Federated Conference on Computer Science and Information Systems (FedCSIS)*. IEEE, 2018, pp. 515–522.
- [26] I. Chalkidis, M. Fergadiotis, D. Tsarapatsanis, N. Aletras, I. Androutsopoulos, and P. Malakasiotis, "Paragraph-level rationale extraction through regularization: A case study on european court of human rights cases," *arXiv preprint arXiv:2103.13084*, 2021.

- [27] L. Zheng, N. Guha, B. R. Anderson, P. Henderson, and D. E. Ho, "When does pretraining help? assessing self-supervised learning for law and the casehold dataset of 53,000+ legal holdings," in *Proceedings of the Eighteenth International Conference on Artificial Intelligence and Law*, 2021, pp. 159–168.
- [28] Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, and V. Stoyanov, "Roberta: A robustly optimized bert pretraining approach," *arXiv preprint arXiv:1907.11692*, 2019.
- [29] I. Chalkidis, M. Fergadiotis, P. Malakasiotis, N. Aletras, and I. Androutsopoulos, "LEGAL-BERT: The muppets straight out of law school," in *Findings of the Association for Computational Linguistics: EMNLP 2020*. Online: Association for Computational Linguistics, Nov. 2020, pp. 2898–2904. [Online]. Available: <https://aclanthology.org/2020.findings-emnlp.261>
- [30] I. Beltagy, M. E. Peters, and A. Cohan, "Longformer: The long-document transformer," *arXiv preprint arXiv:2004.05150*, 2020.
- [31] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," *arXiv preprint arXiv:1412.6980*, 2014.

FedCSIS_2023_paper_6516.pdf

ORIGINALITY REPORT

91%

SIMILARITY INDEX

PRIMARY SOURCES

1

arxiv.org

Internet

5002 words — 91%

EXCLUDE QUOTES OFF
EXCLUDE BIBLIOGRAPHY OFF

EXCLUDE SOURCES OFF
EXCLUDE MATCHES OFF