

Topic modelling for risk identification in Data Protection Act Judgements

Aaron Ceross¹[0000–0001–8863–6772] and Andrew Simpson¹[0000–0003–3597–2232]

Department of Computer Science, University of Oxford,
Wolfson Building, Parks Road, Oxford OX1 3QD
United Kingdom
{firstname.lastname}@cs.ox.ac.uk

Abstract. Data protection legislation, such as the EU’s General Data Protection Regulation (GDPR), obliges data controllers to address risks to personal data. Risk assessment rules for data protection stipulate taking into account instances where the processing of personal data may affect other rights of the individual. It is acknowledged that engineering systems in order to address all risks is challenging and there is a need for prioritisation of risks. Previously decided decisions regarding personal data may provide insight to facilitate this. To this end, we ask: (i) in what context has data protection legislation been invoked in courts? and (ii) what other rights were affected by these cases? To answer these questions, we use structural topic modelling (STM) to extract topics from the case judgements related to the United Kingdom’s Data Protection Act (DPA), incorporating covariate information related to the case outcomes, such as court level and year. The outputs of the model can be utilised to provide topics which relate to context; they can also examine how the other associated variables relate to the resultant topics. We demonstrate the utility of unsupervised text clustering for context and risk identification in legal texts. In our application, we find that STM provides clear topics and allows for the analysis of trends regarding the topics, clearly showing where data protection issues succeed and fail in courts.

Keywords: natural language processing · machine learning · data protection.

1 Introduction

Data controllers are obliged to protect the personal data of data subjects held within their information systems. These obligations are derived from laws which are designed to hold the controllers of information systems accountable for the misuse of the data. The European Union’s General Data Protection Regulation (GDPR) [2] provides a framework by which to govern the management of the personal data of individuals. The legislation provides rights for the data subjects (those whose personal data is collected, stored, and processed) and obligations for the data controllers (those that are managing the personal information). The GDPR contains, within Article 25, a duty on data controllers to incorporate a

“data protection-by-design” approach (DPbD) to information systems. DPbD is based on the wider concept of “Privacy-by-Design” [10], which argues that seven foundational principles of privacy protection ought to motivate system design choices. Aside from the aspirational aim of the principles, the literature has recognised the challenge in transforming the abstract nature of the principles into actionable system requirements [32, 26].

There are challenges regarding the combination of: (i) a duty to adopt a “data protection-by-design” with (ii) an objective, metrics-driven risk-based approach to personal data management. The lack of specified methodology and established risk metrics makes adoption difficult for practitioners [12], especially as this field has few empirical studies [11]. While the emphasis on risk-based assessment features prominently within the GDPR, the nature, scope, and operation of this risk-based approach is not adequately defined within the regulation itself. The questions addressed by this work are therefore:

1. In what contexts has the Data Protection Act been utilised?
2. What other rights were invoked or affected in these contexts?

We approach these questions utilising natural language processing and unsupervised clustering. We aim to extract the context of judicial decisions through topic modelling and assess the resultant topics in respect of the case outcomes. Furthermore, we parse the cases and identify cited legislation. The cited legislation acts as a proxy for other rights aside from those found in the data protection legislation. We then group these cited legislative instruments according to the resultant topics.

We provide a background to data protection obligations and the challenges these pose to the design of information systems in Section 2. Section 3 details the collection and description of the data used in this analysis. We describe the methodology in Section 4. Section 5 presents the results. A discussion is provided in Section 6. We outline future work in the final section.

2 Background

Article 33(1) of the GDPR obliges data controllers to undertake a ‘data protection impact assessment’ (DPIA) where the processing of personal data has a high risk to the “rights and freedoms of natural persons”.¹ The provided guidance on this provision in order to assist those responsible for managing personal data states that the ‘rights and freedoms’ contemplated in Article 33(1) go beyond privacy and may include other rights related to expression and conscience [6].

¹ “Where a type of processing in particular using new technologies, and taking into account the **nature, scope, context and purposes of the processing**, is likely to result in a **high risk to the rights and freedoms of natural persons**, the controller shall, prior to the processing, carry out an assessment of the impact of the envisaged processing operations on the protection of personal data. A single assessment may address a set of similar processing operations that present similar high risks.” (emphasis added)

However, the guidance provides no indication as to the precise methodology which might be employed to determine and evaluate risk, only that it ought to be done in the form of a DPIA — a derivative of ‘privacy impact assessments’ (PIAs). Both DPIAs and PIAs are qualitative assessments [35, 33]. However, there have been increased efforts to incorporate more formal methods into the risk assessment methodologies so as to maximise utility to information system designers [22, 4]. In examining the nature and function of DPIAs, Binns [7] argues that the mandatory nature of the DPIA is a form of ‘meta-regulation’, which functions more accurately as a legal compliance assessment rather than any purported holistic means for system design and management.

Gellert [18] argues that the notion of risk described in Article 33(1) is more about “compliance risk” — with less compliance meaning a higher likelihood of infringing on rights. This definition of risk by the GDPR and supplementary material on this topic by the Article 29 Working Party [5, 6] leaves essential elements of risk assessment unaddressed, including the criteria for harm, calculation of likelihood, and selection of appropriate methodologies. It is argued therein that the effect of this is to render the notion of risk in this article “irrelevant”, concluding that it only considers severity rather than likelihood, has no means to calculate any notion of risk, and does not identify appropriate sources from data.

Thus, it may be argued that a quantitative approach is necessary for legal analysis to derive these necessary contextual and risk informatics. There has been some work regarding prediction of classification of judicial outcomes. For example, Katz *et al.* [19] utilised random forests to train a classifier for US Supreme Court cases. Aletras *et al.* [3] trained an SVM classifier using n-grams and topic in order to predict judgements from the European Court of Human Rights. Šadl and Olsen [29] evaluated the judgements of the Court of Justice of the European Union and the European Court of Human Rights through corpus linguistics and citation network analysis. In their work, the authors focused on determining: (i) the structure of the examined case law; (ii) the language utilised by the court (re-occurring and co-occurring expressions); and (iii) the identification of legal arguments. Ceross and Zhu [13] examined machine learning models for the prediction of data protection case outcomes based on case briefs translated in three languages. The work found that the identified predictive features were different in the three languages evaluated, raising issues regarding the generalisability of data protection risk predictors in a multilingual setting. Nevertheless, these contributions have shown that there is promise in using natural language processing in legal analysis and may suggest a means by which to derive metrics for legal analysis which may inform data protection risk analysis.

3 Data

The GDPR, which replaced the EU’s Personal Data Directive [1], is applicable to all member states of the EU. Despite renouncing its membership of the EU, the United Kingdom (UK) mapped the GDPR into its law in 2018 to ensure that the

GDPR would remain applicable. While there have been indications from the UK government that there may be changes to data protection legislation (potentially resulting in deviation from the GDPR) [15], the current GDPR rules remain in force in the UK and justifies utilising judgements from the UK courts for our experiment.

In total, there have been three Data Protection Acts in British law: in 1984, 1998, and 2018. A dataset of all judgements related to all of these was constructed using the texts provided by the British and Irish Legal Information Institute (BAILII) website [9]. We input the search string “data protection act” as an exact phrase into the case law search. We limited the search to the (i) Supreme Court / House of Lords; (ii) Court of Appeal; and (iii) the High Court (excluding the Family Division). We excluded the courts of Northern Ireland and Scotland, given the different legal systems in those jurisdictions. This resulted in total of 408 cases. After removing duplicate results, cases that mentioned the Act fewer than five times were removed in order ensure that the cases were truly discussing the Act and not merely mentioning it in passing.

The resulting corpus is comprised of 238 cases, dating from 1999 to 2019. There are also 11 variables in addition to the text of the judgements. These include name of case, court level, citation, judges, the date of the judgement, which party was relying on the DPA, the relevant Data Protection Act, and outcome of the case. The vast majority of the cases (227 (95%)) involve the 1998 version of the DPA, with the older 1984 version having only 6 cases, and only 5 cases invoking the newer 2018 version. This provides uniformity and consistency in the analysis of Section 5.

4 Methods

This section describes the methodology used. We describe the theory and construction of STM, as well as methods for increasing coherence in topic models. Additionally, we detail the pre-processing steps of the data in order to input the text into the model.

4.1 Structural topic modelling

To address the challenges set out in Sections 1 and 2, we require a means by which to identify context and quantify the relationship between the different contexts. In this study, we utilise *structural topic modelling* (STM), an unsupervised machine learning approach to deriving topics from a text corpora, incorporating the document’s covariate information. This allows for the identification of topics within a text corpora and the ability to link this to the document’s metadata. In this present work, the variable of interest is decision outcomes.

We briefly describe the details of the STM model. This model is combination of three existing models: (i) the correlated topic model [8]; (ii) the Dirichlet-Multinomial Regression topic model [24]; and (iii) the Sparse Additive Generative model [17].

To summarise the model as proposed by and detailed by Roberts *et al.* [27], a text corpus is composed of documents, $d_1, \dots, d_n$, where each document is a mixture of k topics $z_1, \dots, z_k$. Topics are composed of any words, w . The distribution of topics $\theta_{d,d}$ is drawn from a global distribution. For each w in d (indexed by n), a topic is assigned based on the specific-document distribution over topics $z_{d,n} | \theta_d \sim \text{Multinomial}(\theta)$

Based on the topic assigned, the observed word $w_{d,n}$ is drawn over the vocabulary

$$w_{d,n} | z_{d,n}, \beta_d, k = z \sim \text{Multinomial}(\beta_{d,k=z})$$

Here, $\beta_{d,k=z}$ is the probability of selecting the v -th word in the vocabulary for topic k . However, unlike LDA, topics may be correlated and the prevalence of those topics can be influenced by a set of covariates X through a logistic-normal generalised linear model (GLM) based on document covariates X_d :

$$\theta_d \gamma, \Sigma \sim \text{LogisticNormal}(X_d \gamma, \Sigma)$$

This forms the document-specific distribution over words representing each topic (k) using the baseline word distribution (m), the topic specific deviation (κ_k), the covariate group deviation κ_g , and the interaction between the two κ_i .

$$\beta_{d,k} \propto \exp(m + \kappa_k + \kappa_{g_d} + \kappa_{i=\kappa_{g_d}})$$

The model utilises a semi-collapsed variational expectation-maximisation (EM) algorithm, whereby the E-step provides the joint optimum of θ while the global parameters, κ , γ and Σ , are inferred during the M-step. Roberts *et al.* [28] note that the theoretical guarantees associated with mean-field variational inference do not hold as the STM prior is not conjugate to the likelihood.

4.2 Determining topic numbers

The determination of number of topics for topic modelling is a challenge. There is no set means of selecting an appropriate number of topics, but methods have been proposed. Wallach *et al.* [34] demonstrated that the heldout likelihood may help in determining the number of topics. However this is contrasted with the findings in Chang *et al.* [14], which showed a negative relationship between perplexity and human based evaluations.

In our study we use the method proposed by Mimno *et al.* [23], which focuses on semantic coherence of the model. In that work, the authors found that their proposed metric significantly correlates with human annotator judgement of topic quality. The method determines the number of times words v and v' occur simultaneously in a document, given as $D(v, v')$. For a list of the M most probable words in topic k , the semantic coherence for topic k is given as

$$C_k = \sum_{i=2}^M \sum_{j=1}^{i-1} \log \left(\frac{D(v, v') + 1}{D(v_j)} \right)$$

Informed by measures of semantic coherence and exclusivity, we set k to 6 for this study based on a search for topic numbers based on these parameters. In one of the introductory works on STM, Roberts *et al.* [28] combined semantic coherence with exclusivity in order to reduce the possibility of few, yet highly occurring, words dominating topics. In that study, the authors used FREX, a weighted harmonic mean of the word’s rank in terms of exclusivity and frequency.

$$FREX_{k,v} = \left(\frac{\omega}{ECDF(\beta_k, v / \sum_{j=1}^K \beta_j, v)} + \frac{1 - \omega}{ECDF(\beta_k, v)} \right)^{-1}$$

Here, ECDF is the empirical cumulative distribution function and ω is the weight.

4.3 Pre-processing of text

Schofield *et al.* [30] found that aggressive removal of stop-words may heuristically improve the topics returned. In this study the stop words list included common legal terms (e.g. “court”, “evidence”, “hearing”). As we are interested in the themes of the cases and features pertaining to data protection issues, the names of the parties to the case, the legal representation, and the sitting judges or panel were identified from the header notes of each case document and added to the stop word list. In this work we lemmatised the words rather than stemming as the latter has been shown to be less effective than lemmatisation when returning topics that are coherent [31]. We also tagged the word tokens with parts-of-speech in order better refine the topic models as Martin and Johnson [21] found that use of noun-only topic models were efficient. We included both nouns and verbs in order to widen the scope of what may constitute “risk” in the judgements.

5 Results

The results of the model are provided in two parts. In the first part, we describe the outcome of STM and detail the topics. Further information using covariate data is also examined in relation to the provided topics. In the second part, we evaluate the court outcomes in respect of the topics from the model and identify associated cited legislation in these cases.

5.1 Data protection case topics

A topic model provides clusters of documents based on words which are closely associated within the corpus. Figure 1 shows the outcome of the model on the DPA judgements, listing the top three terms of that cluster. Topic 2 is the largest topic proportion, which involves the complaints and the police. Topic 1 groups issues related to damages for breaches. The smallest topic proportion relates to governmental information, particularly correspondence and exemptions.

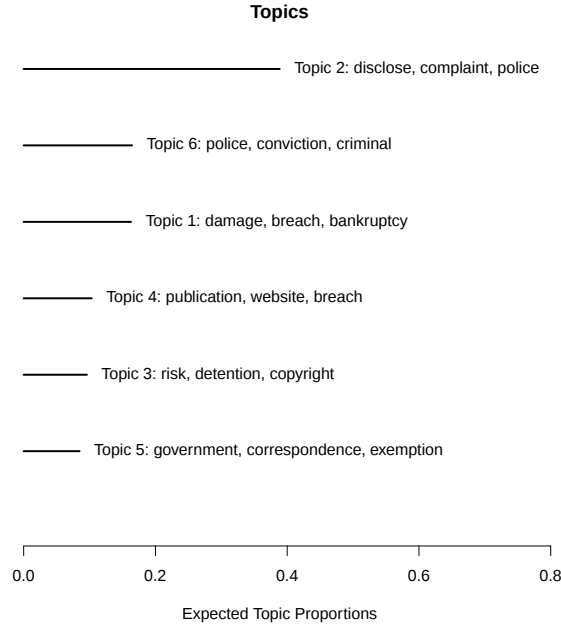


Fig. 1: Topics and the proportions

One of the features of STM is the ability to link covariate information with the topics. We are interested in the change of contexts for the topics over time. Figure 2 illustrates the changes in expected proportion of topics between 1999 and 2019. Topic 2 (disclose, complaint, police), the largest topic, has a decline from 2010. Topic 3 (risk, detention, copyright), which was one of the smallest topics, has a consistent increase. It starts in 1999 with virtually no reference, rising to its modest 20% of topics nearly 20 years later. Topic 1 (damage, breach, bankruptcy), is the most variable, rising and falling in more pronounced manner than the other topics. Topics 6 (police, conviction, criminal) has a constant trend, with no meaningful changes.

5.2 Court outcomes by topic

The 238 cases are divided with 100 (39%) cases overall having an acceptance of the data protection arguments and 138 (65%) dismissing such argument. The acceptance rate by court is shown in Table 1, which shows that the Court of Appeal significantly has the lowest acceptance rate ($p < 0.05$), with the High Court and Supreme Court somewhat higher. In our analysis we do not consider the nature of the appeals or the final outcome of the case, as not all cases related directly to the Data Protection Act; cases for inclusion may have merely raised a point of data protection.

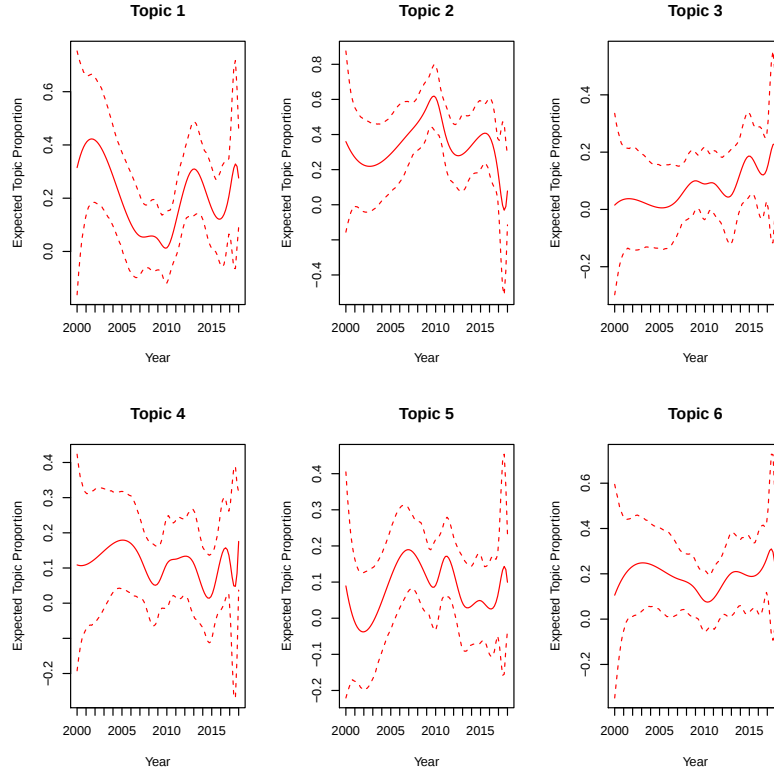


Fig. 2: Topic trends by year

Using the resultant STM for the data protection cases, the association of outcomes with topics may be determined. Figure 3 shows that Topic 1 (damage, breach, bankruptcy) is most associated with arguments for data protection issues being dismissed. Topic 5 (government, correspondence, exemption) is also more strongly related to being dismissed as well. Topic 6 (police, conviction, criminal) is most correlated with being allowed. Topic 2 (disclose, complaint, police) and Topic 3 (risk, detention, copyright) are also more allowed. Topic 4 (publication, website, breach) is only slightly more associated with the ‘dismissed’ label.

In addition, we identified there are 161 legislative instruments that are considered when arguing cases that involves data protection. For this analysis, we consider the top five most cited legislation, listed in Table 2. The Human Rights Act 1998 (HRA 1998) incorporated the European Convention of Human Rights, which includes the right to privacy (Article 8), something that had been absent in English and Welsh law until the Act. The HRA 1998 accounts for 20% of the cited legislation, greatly outnumbering all other legislation. This is likely due to data protection issues deeply impinging on privacy concerns and thus necessi-

Table 1: Number of cases by court.

Court	#Cases	Allowed	Dismissed	Allowed Rate
High Court	160	73	87	0.46
Court of Appeal	69	25	46	0.33
Supreme Court	9	4	5	0.44

Table 2: Legislation cited other than the DPA.

Legislation	n	pct
Human Rights Act 1998	90	0.204
Protection from Harassment Act 1997	23	0.052
Freedom of Information Act 2000	20	0.045
Contempt of Court Act 1981	16	0.036
Police and Criminal Evidence Act 1984	12	0.027
Regulation of Investigatory Powers Act 2000	12	0.027

tating a consideration of the law in this area. A group of legislation addresses issues in justice, related to investigative powers and rules of evidence: (i) the Contempt of Court Act 1981, (ii) the Police and Criminal Evidence Act 1984, and (iii) the Regulation of Investigatory Powers Act 2000. The Protection from Harassment Act 1997 deals with criminal offences related to the physical and mental harassment against individuals.

In Figure 4 the division between the case outcome becomes more apparent. The HRA is cited nearly equally for both case outcomes. In cases where the Freedom of Information Act 2000 is cited, there are slightly more dismissals of data protection arguments than not. However, where investigatory powers by the state are invoked, arguments for data protection issues are dismissed.

6 Discussion

The context for a decision is fundamental to the application of the law. In this work, there are clear associations with topics to court outcomes, which broadens knowledge about how the Data Protection Act has been utilised in the courts of England and Wales. In relation to Question 1 of Section 1, we find that the analysis has provided six discrete topics in which data protection cases have fallen within the last 20 years. Through STM, we also have identified the trajectory for each of those topics over time. From the analysis, disclosures, complaints about disclosures and the police (Topics 2 and 6) dominate the subject matter for DPA cases. However, these topics are often have negative outcomes (Figure 3). This means that arguments provided for data protection issues are dismissed. Regarding Question 2 of Section 1, we find that the cases tend to be in tension with the use of investigatory powers for criminal matters. This suggests that the courts are less likely to prioritise data protection issues over criminal investigations.

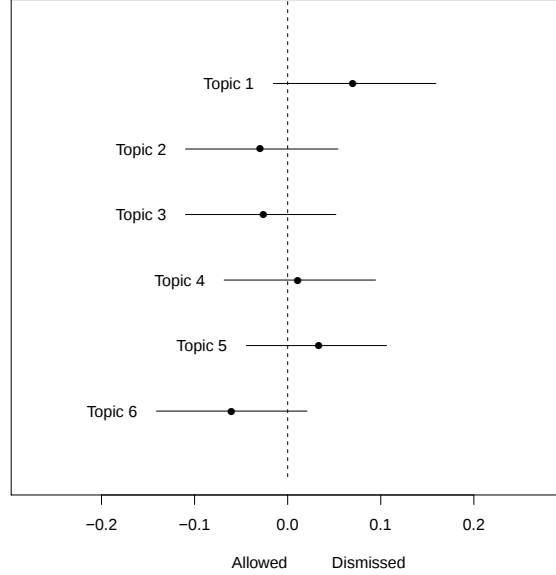


Fig. 3: The point estimate and 80% confidence interval of the mean difference in topic proportions for court decisions allowing or dismissing arguments for data protection issues.

From this preliminary experiment, we observe that system designers ought to consider how manage personal data that likely to be requested by authorities. This has implications for not only the types of data that these systems collect, but also how the data is conveyed when requested.

Considering the results more widely, this work provides continuing evidence for the use of computational methods, such as natural language processing and wider artificial intelligence, in order to derive novel insights into the law. A quantitative approach to legal scholarship is arguably necessary in order to address the acknowledged limited nature of legal theory [25]. Quantitative approaches allow for more experimentation and may complement more traditional doctrinal analysis, providing a more holistic view of legal informatics. Automated methods, such as the unsupervised clustering used in this work, can be applied across large corpora of texts, providing novel insights on emergent patterns.

There are several limitations that ought to be borne in mind when considering the outcome of this work. The first is that we have analysed the cases without a view of precedent. In many jurisdictions, the lower courts follow the decisions provided by higher courts: cases may be decided in one manner only to be reversed later by a higher court, as a legal doctrine develops. Furthermore, the complexity of issues in court cases may not be captured adequately by topic

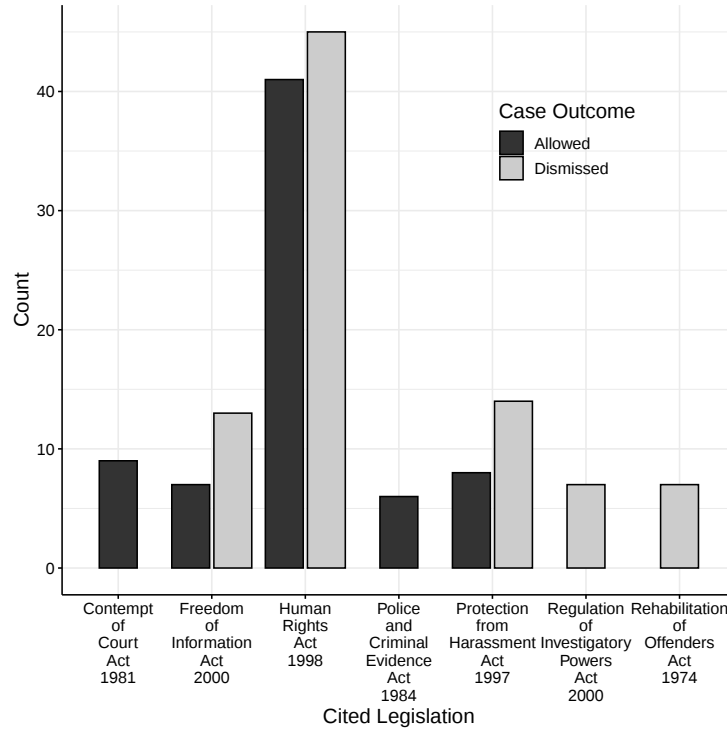


Fig. 4: Count of legislation by case outcome.

modelling. The model relies on the use of n-grams which do not capture semantic meaning. Additionally, the issue of stop words and phrases are particularly challenging as legal phrases of interest may incorporate these. Nevertheless, topic modelling has potential to provide broad overviews of contexts in legal corpora, which is valuable to legal informatics.

7 Conclusions

In this work, we utilised STM in order to cluster English and Welsh data protection cases in order to derive contexts where the rights are being invoked. STM is an unsupervised text clustering methodology able to associate covariates to the resultant topics, which in this case was the outcome of the case, the year of the case and the court level the case was decided. We find that there is modest but practical utility in the use of STM with legal texts. There are issues related to the retention of semantic meaning, as the method does not retain it. In future work, we hope to explore transformers such as BERT and ROBERTA [16, 20] which can retain semantic meaning. These models require little training and may facilitate more effective risk predictors and richer coherence in clusters.

Acknowledgements

The authors would like to thank the anonymous reviewers for their helpful and insightful comments. AC would like to acknowledge and thank the support of the Centre of Doctoral Training in Cyber Security, which is funded by EPSRC grant number (EP/P00881X/1).

References

1. Directive 95/46/EC of the European Parliament and of the Council of 24 October 1995 on the protection of individuals with regard to the processing of personal data and on the free movement of such data. <http://eur-lex.europa.eu>
2. Regulation on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). L119, 4/5/2016, p. 1–88 (2016)
3. Aletras, N., Tsarapatsanis, D., Preotiuc-Pietro, D., Lampos, V.: Predicting judicial decisions of the European Court of Human Rights: A natural language processing perspective. *PeerJ Computer Science* **2**, e93 (2016)
4. Alshammari, M., Simpson, A.: Towards an effective privacy impact and risk assessment methodology: risk analysis. In: Garcia-Alfaro, J., Herrera-Joancomartí, J., Livraga, G., Rios, R. (eds.) *Data Privacy Management, Cryptocurrencies and Blockchain Technology*. pp. 209–224. Springer International Publishing (2018)
5. Article 29 Working Party: Statement on the role of a risk-based approach in data protection legal frameworks. Accessed at: https://ec.europa.eu/justice/article-29/documentation/opinion-recommendation/files/2014/wp218_en.pdf (2014)
6. Article 29 Working Party: Revised Guidelines on Data Protection Impact Assessment (DPIA) and determining whether processing is “likely to result in a high risk” for the purposes of Regulation 2016/679. Accessed at: https://ec.europa.eu/newsroom/document.cfm?doc_id=44137 (2017)
7. Binns, R.: Data protection impact assessments: A meta-regulatory approach. *International Data Privacy Law* **7**(1), 22–35 (2017)
8. Blei, D., Lafferty, J.: Correlated topic models. *Advances in Neural Information Processing Systems* **18**, 147 (2006)
9. British and Irish Legal Information Institute: Home page. <https://www.bailii.org/> (2021), accessed Jan 2021
10. Cavoukian, A., Taylor, S., Abrams, M.E.: Privacy by design: essential for organizational accountability and strong business practices. *Identity in the Information Society* **3**(2), 405–413 (2010)
11. Ceross, A., Simpson, A.C.: The use of data protection regulatory actions as a data source for privacy economics. In: Tonetta, S., Schoitsch, E., Bitsch, F. (eds.) *Computer Safety, Reliability, and Security (SAFEComp)*. Lecture Notes in Computer Science (LNCS), vol. 10489, pp. 350–360. Springer (2017)
12. Ceross, A., Simpson, A.C.: Rethinking the proposition of privacy engineering. In: *Proceedings of the New Security Paradigms Workshop*. pp. 89–102. NSPW ’18, ACM, New York, NY, USA (2018)
13. Ceross, A., Zhu, T.: Prediction of monetary penalties for data protection cases in multiple languages. In: *Proceedings of the Eighteenth International Conference on Artificial Intelligence and Law*. p. 185–189. ICAIL ’21, Association for Computing Machinery, New York, NY, USA (2021)

14. Chang, J., Gerrish, S., Wang, C., Boyd-Graber, J.L., Blei, D.M.: Reading tea leaves: How humans interpret topic models. In: *Advances in Neural Information Processing Systems*. pp. 288–296 (2009)
15. Department for Digital, Culture, Media, & Sport: Government response to the consultation on the National Data Strategy. <https://www.gov.uk/government/consultations/uk-national-data-strategy-nds-consultation/outcome/government-response-to-the-consultation-on-the-national-data-strategy> (2021), accessed 08 September 2021
16. Devlin, J., Chang, M.W., Lee, K., Toutanova, K.: 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)*. pp. 4171–4186. Association for Computational Linguistics, Minneapolis, Minnesota (Jun 2019). <https://doi.org/10.18653/v1/N19-1423>, <https://www.aclweb.org/anthology/N19-1423>
17. Eisenstein, J., Ahmed, A., Xing, E.P.: Sparse additive generative models of text. In: *Proceedings of the 28th International Conference on Machine Learning*. pp. 1041–1048. Citeseer (2011)
18. Gellert, R.: Understanding the notion of risk in the General Data Protection Regulation. *Computer Law & Security Review* **34**(2), 279–288 (2018)
19. Katz, D.M., Bommarito II, M.J., Blackman, J.: A general approach for predicting the behavior of the Supreme Court of the United States. *PloS One* **12**(4), e0174698 (2017)
20. Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., Stoyanov, V.: Roberta: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692* (2019)
21. Martin, F., Johnson, M.: More efficient topic modelling through a noun only approach. In: *Proceedings of the Australasian Language Technology Association Workshop 2015*. pp. 111–115 (2015)
22. Meis, R., Heisel, M.: Supporting privacy impact assessments using problem-based privacy analysis. In: Lorenz, P., Cardoso, J., Maciaszek, L.A., van Sinderen, M. (eds.) *Software Technologies*. pp. 79–98. Springer International Publishing, Cham (2016)
23. Mimno, D., Wallach, H.M., Talley, E., Leenders, M., McCallum, A.: Optimizing semantic coherence in topic models. In: *Proceedings of the Conference on Empirical Methods in Natural Language Processing*. pp. 262–272. EMNLP ’11, Association for Computational Linguistics, Stroudsburg, PA, USA (2011)
24. Mimno, D.M., McCallum, A.: Topic models conditioned on arbitrary features with dirichlet-multinomial regression. In: *Proceedings of the 24th Conference on Uncertainty in Artificial Intelligence*. vol. 24, pp. 411–418. Citeseer (2008)
25. Posner, R.A.: *Frontiers of Legal Theory*. Harvard University Press (2001)
26. van Rest, J., Boonstra, D., Everts, M., van Rijn, M., van Paassen, R.: Designing privacy-by-design. In: Preneel, B., Ikonomou, D. (eds.) *Privacy Technologies and Policy: First Annual Privacy Forum (AFP)*. *Lecture Notes in Computer Science (LNCS)*, vol. 8319, pp. 55–72. Springer (2014)
27. Roberts, M.E., Stewart, B.M., Airoidi, E.M.: A model of text for experimentation in the social sciences. *Journal of the American Statistical Association* **111**(515), 988–1003 (2016)
28. Roberts, M.E., Stewart, B.M., Tingley, D., Lucas, C., Leder-Luis, J., Gadarian, S.K., Albertson, B., Rand, D.G.: Structural topic models for open-ended survey responses. *American Journal of Political Science* **58**(4), 1064–1082 (2014)

29. Šadl, U., Olsen, H.P.: Can quantitative methods complement doctrinal legal studies? using citation network and corpus linguistic analysis to understand international courts. *Leiden Journal of International Law* **30**(2), 327–349 (2017)
30. Schofield, A., Magnusson, M., Mimno, D.: Pulling out the stops: Rethinking stop-word removal for topic models. In: *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 2, Short Papers*. vol. 2, pp. 432–436 (2017)
31. Schofield, A., Mimno, D.: Comparing apples to apple: The effects of stemmers on topic models. *Transactions of the Association for Computational Linguistics* **4**, 287–300 (2016)
32. Spiekermann, S.: The challenges of privacy by design. *Communications of the ACM* **55**(7), 38–40 (2012)
33. The Information Commissioner’s Office: Data protection impact assessments. <https://ico.org.uk/media/for-organisations/guide-to-the-general-data-protection-regulation-gdpr/data-protection-impact-assessments-dpias-1-0.pdf> (March 2018)
34. Wallach, H.M., Murray, I., Salakhutdinov, R., Mimno, D.: Evaluation methods for topic models. In: *Proceedings of the 26th Annual International Conference on Machine Learning*. pp. 1105–1112. ICML ’09, ACM, New York, NY, USA (2009)
35. Wright, D.: The state of the art in privacy impact assessment. *Computer Law & Security Review* **28**(1), 54–61 (2012)