

PrivacyGLUE: A Benchmark Dataset for General Language Understanding in Privacy Policies

Anonymous ACL submission

Abstract

Benchmarks for general language understanding have been rapidly developing in recent years of NLP research, with well-known examples such as GLUE and SuperGLUE. While benchmarks have been proposed in the legal language domain, virtually no such benchmarks exist for privacy policies despite their increasing importance in modern digital life. This could be explained by privacy policies falling under the legal language domain, but we find evidence to the contrary that motivates a separate benchmark for privacy policies. Consequently, we propose PrivacyGLUE as the first comprehensive benchmark of relevant and high-quality privacy tasks for measuring general language understanding in the privacy language domain. Furthermore, we release performances from the BERT, RoBERTa, Legal-BERT, Legal-RoBERTa and PrivBERT transformer language models and perform model-pair agreement analysis to detect PrivacyGLUE task examples where models benefited from domain specialization. Our findings show PrivBERT outperforms other models by an average of 2 – 3% over all PrivacyGLUE tasks, shedding light on the importance of in-domain pretraining for privacy policies. We believe PrivacyGLUE can accelerate NLP research and improve general language understanding for humans and AI algorithms in the privacy language domain.

1 Introduction

Data privacy is evolving into a critical aspect of modern life with the United Nations (UN) describing it as a *human right in the digital age* (Gstrein and Beaulieu, 2022). Despite its importance, several studies have demonstrated high barriers to the understanding of privacy policies (Obar and Oeldorf-Hirsch, 2020) and estimate that an average person would require ~200 hours annually to read through all privacy policies encountered in their daily life (McDonald and Cranor, 2008). To

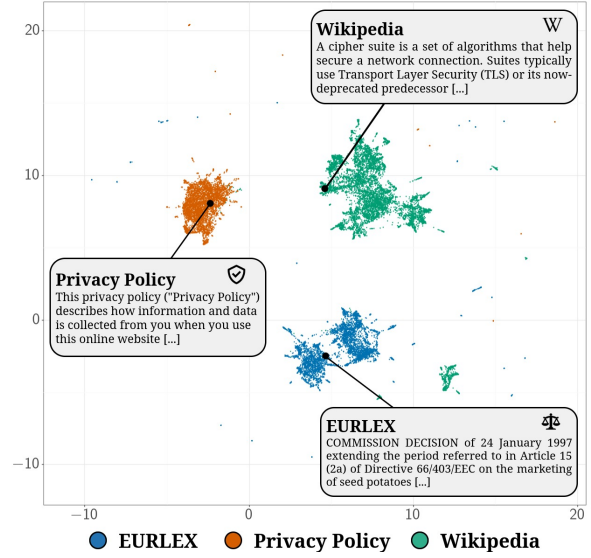


Figure 1: UMAP visualization of BERT embeddings from Wikipedia, European Legislation (EURLEX) and company privacy policy documents with a total of 2.5M tokens per corpus

address this, studies such as Wilson et al. (2016) recommend training Artificial Intelligence (AI) algorithms on appropriate benchmark datasets to assist humans in understanding privacy policies.

In recent years, benchmarks have been gaining popularity in Machine Learning and Natural Language Processing (NLP) communities because of their ability to holistically evaluate model performance over a variety of representative tasks. GLUE (Wang et al., 2018) and SuperGLUE (Wang et al., 2019) are examples of popular NLP benchmarks which measure the natural language understanding capabilities of SOTA models. NLP benchmarks are also developing rapidly in language domains, with LexGLUE (Chalkidis et al., 2022) being an example of a recent benchmark hosting several difficult tasks in the legal language domain. Interestingly, we do not find similar NLP benchmarks in the privacy language domain for privacy policies. While

Task	Source	Task Type	Train/Dev/Test Instances	# Classes
OPP-115	Wilson et al. (2016)	Multi-label sequence classification	2,185/550/697	12
PI-Extract	Bui et al. (2021)	Multi-task token classification	2,579/456/1,029	3/3/3 [†]
Policy-Detection	Amos et al. (2021)	Binary sequence classification	773/137/391	2
PolicyIE-A	Ahmad et al. (2021)	Multi-class sequence classification	4,109/100/1,041	5
PolicyIE-B	Ahmad et al. (2021)	Multi-task token classification	4,109/100/1,041	29/9 [†]
PolicyQA	Ahmad et al. (2020)	Reading comprehension	17,056/3,809/4,152	–
PrivacyQA	Ravichander et al. (2019)	Binary sequence classification	157,420/27,780/62,150	2

Table 1: Summary statistics of PrivacyGLUE benchmark tasks; [†] PI-Extract and PolicyIE-B consist of four and two subtasks respectively and the number of BIO token classes per subtask are separated by a forward slash character

this could be explained by privacy policies falling under the legal language domain due to their formal and jargon-heavy nature, we claim that privacy policies fall under a distinct language domain and cannot be subsumed under any other specialized NLP benchmark such as LexGLUE.

To investigate this claim, we gather documents from Wikipedia (Wikimedia Foundation, 2022), European Legislation (EURLEX; Chalkidis et al. 2019) and company privacy policies (Mazzola et al., 2022), with each corpus truncated to 2.5M tokens. Next, we feed these documents into BERT and gather contextualized embeddings, which are then projected to 2-dimensional space using UMAP (McInnes et al., 2018). In Figure 1, we observe that the three domain corpora cluster independently, providing evidence that privacy policies lie in a distinct language domain from both legal and wikipedia documents. With this motivation, we propose PrivacyGLUE as the first comprehensive benchmark for measuring general language understanding in the privacy language domain. Our main contributions are threefold:

1. Composition of seven high-quality and relevant PrivacyGLUE tasks, specifically OPP-115, PI-Extract, Policy-Detection, PolicyIE-A, PolicyIE-B, PolicyQA and PrivacyQA.
2. Benchmark performances of five transformer language models on all aforementioned tasks, specifically BERT, RoBERTa, Legal-BERT, Legal-RoBERTa and PrivBERT.
3. Model agreement analysis to detect PrivacyGLUE task examples where models benefited from domain specialization.

We release PrivacyGLUE as a fully configurable benchmark suite for straight-forward reproducibil-

ity and production of new results in our public GitHub repository¹. Our findings show that PrivBERT, the only model pretrained on privacy policies, outperforms other models by an average of 2 – 3% over all PrivacyGLUE tasks, shedding light on the importance of in-domain pretraining for privacy policies. Our model-pair agreement analysis explores specific examples where PrivBERT’s privacy-domain pretraining provided both competitive advantage and disadvantage. By benchmarking holistic model performances, we believe PrivacyGLUE can accelerate NLP research into the privacy language domain and ultimately improve general language understanding of privacy policies for both humans and AI algorithms.

2 Related work

NLP benchmarks have been gaining popularity in recent years because of their ability to holistically evaluate model performance over a variety of representative tasks. GLUE (Wang et al., 2018) and SuperGLUE (Wang et al., 2019) are examples of benchmarks that evaluate SOTA models on a range of natural language understanding tasks. The GEM benchmark (Gehrmann et al., 2021) looks beyond text classification and measures performance in Natural Language Generation tasks such as summarization and data-to-text conversion. The XTREME (Hu et al., 2020) and XTREME-R (Ruder et al., 2021) benchmarks specialize in measuring cross-lingual transfer learning on 40-50 typologically diverse languages and corresponding tasks. Popular NLP benchmarks often host public leaderboards with SOTA scores on supported tasks, thereby encouraging the community to apply new approaches for surpassing top scores.

¹Repository will be made public post-acceptance. Anonymous repository: <https://anonymous.4open.science/r/f4293357886f671347fa69fae3650543>

While the aforementioned benchmarks focus on problem types such as natural language understanding and generation, other benchmarks focus on language domains. The LexGLUE benchmark (Chalkidis et al., 2022) is an example of a benchmark that evaluates models on tasks from the legal language domain. LexGLUE consists of seven English-language tasks that are representative of the legal language domain and chosen based on size and legal specialization. Chalkidis et al. (2022) benchmarked several models such as BERT (Devlin et al., 2019) and Legal-BERT (Chalkidis et al., 2020), where Legal-BERT has a similar architecture to BERT but was pretrained on diverse legal corpora. A key finding of LexGLUE was that Legal-BERT outperformed other models which were not pretrained on legal corpora. In other words, they found that an in-domain pretrained model outperformed models that were pretrained out-of-domain.

In the privacy language domain, we tend to find isolated datasets from specialized studies. Zimmeck et al. (2019), Wilson et al. (2016), Bui et al. (2021) and Ahmad et al. (2021) are examples of studies that introduce annotated corpora for privacy-practice sequence and token classification tasks, while Ravichander et al. (2019) and Ahmad et al. (2020) release annotated corpora for privacy-practice question answering. Amos et al. (2021) is another recent study that released an annotated corpus of privacy policies. As of writing, no comprehensive NLP benchmark exists for general language understanding in privacy policies, making PrivacyGLUE the first consolidated NLP benchmark in the privacy language domain.

3 Datasets and Tasks

The PrivacyGLUE benchmark consists of seven natural language understanding tasks originating from six datasets in the privacy language domain. Summary statistics, detailed label information and representative examples are shown in Table 1, Table 5 (Appendix A) and Table 6 (Appendix B) respectively.

OPP-115 Wilson et al. (2016) was the first study to release a large annotated corpus of privacy policies. A total of 115 privacy policies were selected based on their corresponding company’s popularity on Google Trends. The selected privacy policies were annotated with 12 data privacy practices on a paragraph-segment level by experts in the pri-

vacy domain. As noted by Mousavi Nejad et al. (2020), one limitation of Wilson et al. (2016) was the lack of publicly released training and test data splits which are essential for machine learning and benchmarking. To address this, Mousavi Nejad et al. (2020) released their own training, validation and test data splits for researchers to easily reproduce OPP-115 results. PrivacyGLUE utilizes the "Majority" variant of data splits released by Mousavi Nejad et al. (2020) to compose the OPP-115 task. Given an input paragraph segment of a privacy policy, the goal of OPP-115 is to predict one or more data practice categories.

PI-Extract Bui et al. (2021) focuses on enhanced data practice extraction and presentation to help users better understand privacy policies. As part of their study, they released the PI-Extract dataset consisting of 4.1K sentences (97K tokens) and 2.6K expert-annotated data practices from 30 privacy policies in the OPP-115 dataset. Expert annotations were performed on a token-level for all sentences of selected privacy policies. PI-Extract is broken into four subtasks, where spans of tokens are independently tagged using the BIO scheme commonly used in Named Entity Recognition (NER). Subtasks I, II, III and IV require the classification of token spans for data-related entities that are collected, not collected, not shared and shared respectively. In the interest of diversifying tasks in PrivacyGLUE, we composed PI-Extract as a multi-task token classification problem where all four PI-Extract subtasks are to be jointly learned.

Policy-Detection Amos et al. (2021) developed a crawler for automated collection and curation of privacy policies. An important aspect of their system is the automated classification of documents into privacy policies and non-privacy-policy documents encountered during web crawling. To train such a privacy policy classifier, Amos et al. (2021) performed expert annotations of commonly encountered documents during web crawls and classified them into the aforementioned categories. The Policy-Detection dataset was released with a total of 1.3K annotated documents and is utilized in PrivacyGLUE as a binary sequence classification task.

PolicyIE Inspired by Wilson et al. (2016) and Bui et al. (2021), Ahmad et al. (2021) created PolicyIE, an English corpus composed by 5.3K sentence-level and 11.8K token-level data practice

Model	Source	# Params	Vocab. Size	Pretraining corpora [†]
BERT	Devlin et al. (2019)	110M	30K	Wikipedia, BC (16 GB)
RoBERTa	Liu et al. (2019)	125M	50K	Wikipedia, BC, CC-News, OWT (160 GB)
Legal-BERT	Chalkidis et al. (2020)	110M	30K	Legislation, Court Cases, Contracts (12 GB)
Legal-RoBERTa [‡]	Geng et al. (2021)	125M	50K	Patents, Court Cases (5 GB)
PrivBERT [‡]	Srinath et al. (2021)	125M	50K	Privacy policies (17 GB)

Table 2: Summary of models used in the PrivacyGLUE benchmark; all models used are base-sized variants of BERT/RoBERTa architectures; [†] BC = BookCorpus, CC-News = CommonCrawl-News, OWT = OpenWebText; [‡] models were initialized with the pretrained RoBERTa model

annotations over 31 privacy policies from websites and mobile applications. PolicyIE was designed to be used for machine learning in NLP, to ultimately make data privacy concepts easier for users to understand. We split the PolicyIE corpus into two tasks, namely *PolicyIE-A* and *PolicyIE-B*. Given an input sentence, PolicyIE-A entails multi-class data practice classification while PolicyIE-B entails multi-task token classification over distinct subtasks I and II, which require the classification of token spans for entities that participate in privacy practices and their conditions/purposes respectively. The motivation for composing PolicyIE-B as a multi-task problem is similar to that of PI-Extract.

PolicyQA Ahmad et al. (2020) argue in favour of short-span answers to user questions for long privacy policies. They release PolicyQA, a dataset of 25k reading comprehension examples curated from the OPP-115 corpus from Wilson et al. (2016). Furthermore, they provide 714 human-written questions optimized for a wide range of privacy policies. The final question-answer annotations follow the SQuAD-1.0 format (Rajpurkar et al., 2016), which improves the ease of adaptation into NLP pipelines. We utilize PolicyQA as PrivacyGLUE’s reading comprehension task.

PrivacyQA Similar to Ahmad et al. (2020), Ravichander et al. (2019) argue in favour of annotated question-answering data for training NLP models to answer user questions about privacy policies. They correspondingly released PrivacyQA, a corpus composed by 1.75K questions and more than 3.5K expert annotated answers. Unlike PolicyQA, PrivacyQA proposes a binary sequence classification task where a question-answer pair is classified as either relevant or irrelevant. Correspondingly, we treat PrivacyQA as a binary sequence classification task in PrivacyGLUE.

4 Experimental setup

The PrivacyGLUE benchmark was tested using the BERT, RoBERTa, Legal-BERT, Legal-RoBERTa and PrivBERT models which are summarized in Table 2. We describe the models used and task-specific approaches, and provide details on our benchmark configuration in Appendix E.

4.1 Models

BERT Proposed by Devlin et al. (2019), BERT is perhaps the most well-known transformer language model. BERT utilizes the WordPiece tokenizer (Wu et al., 2016) and is case-insensitive. It is pretrained with the Masked Language Model (MLM) and Next Sentence Prediction (NSP) tasks on the Wikipedia and BookCorpus corpora.

RoBERTa Liu et al. (2019) proposed RoBERTa as an improvement to BERT. RoBERTa uses dynamic token masking and eliminates the NSP task during pretraining. Furthermore, it uses a case sensitive byte-level Byte-Pair Encoding (Sennrich et al., 2016) tokenizer and is pretrained on larger corpora. Liu et al. (2019) reported improved results on various benchmarks using RoBERTa over BERT.

Legal-BERT Chalkidis et al. (2020) proposed Legal-BERT by pretraining BERT from scratch on legal corpora consisting of legislation, court cases and contracts. The sub-word vocabulary of Legal-BERT is learned from scratch using the SentencePiece (Kudo and Richardson, 2018) tokenizer to better support legal terminology. Legal-BERT was the best overall performing model in the LexGLUE benchmark as reported in Chalkidis et al. (2022).

Legal-RoBERTa Inspired by Legal-BERT, Geng et al. (2021) proposed Legal-RoBERTa by further pretraining RoBERTa on legal corpora, specifically patents and court cases. Legal-RoBERTa

is pretrained on less legal data than Legal-BERT while producing similar results on downstream fine-tuning legal domain tasks.

PrivBERT Due to the scarcity of large corpora in the privacy domain, [Srinath et al. \(2021\)](#) proposed PrivaSeer, a novel corpus of 1M English language website privacy policies crawled from the web. They subsequently proposed PrivBERT by further pretraining RoBERTa on the PrivaSeer corpus.

4.2 Task-specific approaches

Given the aforementioned models and tasks, we now describe our task-specific fine-tuning and evaluation approaches. Given an input sequence $s = \{w_1, w_2, \dots, w_N\}$ consisting of N sequential sub-word tokens, we feed s into a transformer encoder and obtain a contextual representation $\{h_0, h_1, \dots, h_N\}$ where $h_i \in \mathbb{R}^D$ and D is the output dimensionality of the transformer encoder. Here, h_0 refers to the contextual embedding for the starting token which is [CLS] for BERT-derived models and <s> for RoBERTa-derived models. For PolicyQA and PrivacyQA, the input sequence s is composed by concatenating the question and context/answer pairs respectively. The concatenated sequences are separated by a separator token, which is [SEP] for BERT-derived models and </s> for RoBERTa-derived models.

4.2.1 Sequence classification

The h_0 embedding is fed into a class-wise sigmoid classifier (1) and softmax classifier (2) for multi-label and binary/multi-class tasks respectively. The classifier has weights $W \in \mathbb{R}^{D \times C}$ and bias $b \in \mathbb{R}^C$ and is used to predict the probability vector $y \in \mathbb{R}^C$, where C refers to the number of output classes. We fine-tune models end-to-end by minimizing the binary cross-entropy loss and cross-entropy loss for multi-label and binary/multi-class tasks respectively.

$$y = \text{sigmoid}(W^\top h_0 + b) \quad (1)$$

$$y = \text{softmax}(W^\top h_0 + b) \quad (2)$$

We report the macro and micro-average F_1 scores for all sequence classification tasks since the former ignores class imbalance while the latter takes it into account.

4.2.2 Multi-task token classification

Each $h_i \in \{h_1, h_2, \dots, h_N\}$ token embedding is fed into J independent softmax classifiers with weights

$W_j \in \mathbb{R}^{D \times C_j}$ and bias $b_j \in \mathbb{R}^{C_j}$ to predict the token probability vector $y_{ij} \in \mathbb{R}^{C_j}$, where C_j refers to the number of output BIO classes per subtask $j \in \{1, 2, \dots, J\}$. We fine-tune models end-to-end by minimizing the cross-entropy loss across all tokens and subtasks.

$$y_{ij} = \text{softmax}(W_j^\top h_i + b_j) \quad (3)$$

We report the macro and micro-average F_1 scores for all multi-task token classification tasks by averaging the respective metrics for each subtask. Furthermore, we ignore cases where B or I prefixes are mismatched as long as the main token class is correct.

4.2.3 Reading comprehension

Each $h_i \in \{h_1, h_2, \dots, h_N\}$ token embedding is fed into two independent linear layers with weights $W_j \in \mathbb{R}^D$ and bias $b_j \in \mathbb{R}$ where $j \in \{1, 2\}$. These linear outputs are then concatenated per layer and a softmax function is applied to form a probability vector y_j across all tokens for answer-start and answer-end token probabilities respectively. We fine-tune models end-to-end by minimizing the cross-entropy loss on the gold answer-start and answer-end indices.

$$y_j = \text{softmax}\left(\begin{bmatrix} W_j \cdot h_1 + b_j & \dots & W_j \cdot h_N + b_j \end{bmatrix}\right) \quad (4)$$

Similar to SQuAD ([Rajpurkar et al., 2016](#)), we report the sample F_1 and exact match accuracy for our reading comprehension task. It is worth noting that [Rajpurkar et al. \(2016\)](#) refer to their reported F_1 score as a macro-average, whereas we refer to it as the sample-average as we believe this is a more accurate term.

5 Results

After running the PrivacyGLUE benchmark with 10 random seeds, we collect results on the test-sets of all tasks. Figure 2 shows the respective results in a graphical form while Table 7 in Appendix C shows the numerical results in a tabular form. In terms of absolute metrics, we observe that PrivBERT outperforms other models for all PrivacyGLUE tasks. We apply the Mann-Whitney U-test ([Mann and Whitney, 1947](#)) over random seed metric distributions and find that PrivBERT significantly outperforms other models on six out of seven PrivacyGLUE tasks with $p \leq 0.05$, where

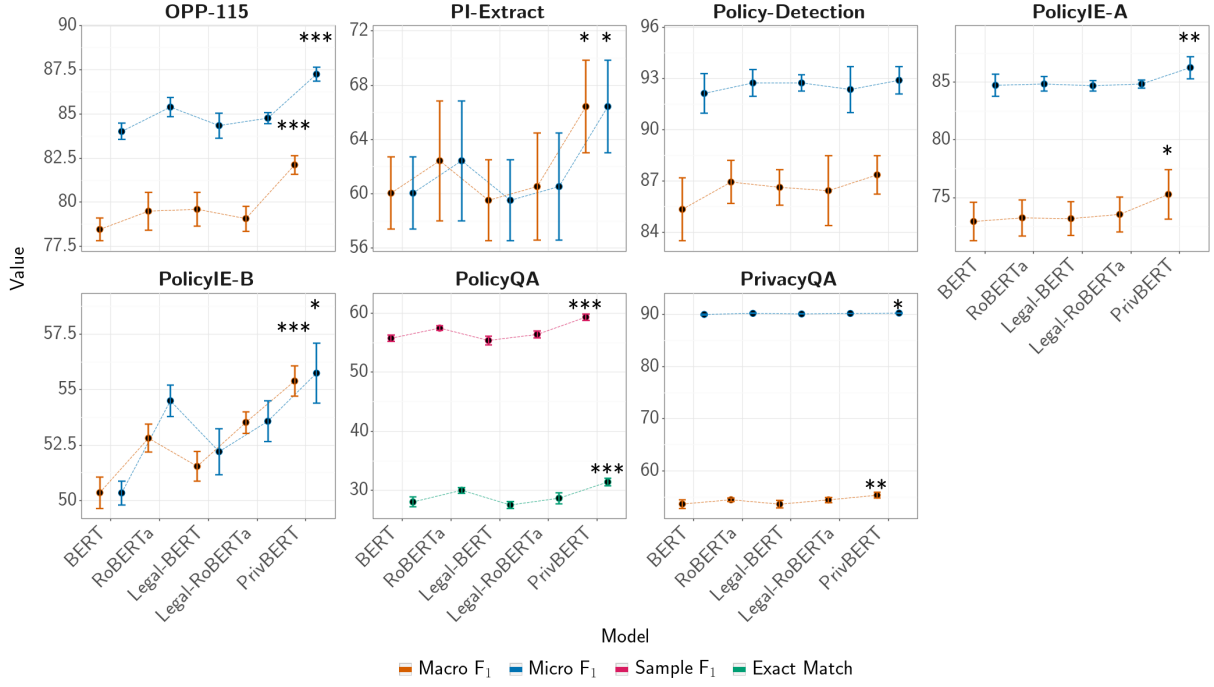


Figure 2: Test-set results of the PrivacyGLUE benchmark where points indicate mean performance and error bars indicate standard deviation over 10 random seeds; *** implies $p \leq 0.001$, ** implies $0.001 < p \leq 0.01$, * implies $0.01 < p \leq 0.05$ given an alternative hypothesis that PrivBERT has a greater performance metric than all other models in a task using the Mann-Whitney U-test

Policy-Detection was the task where the significance threshold was not met. We utilize the Mann-Whitney U-test because it does not require a normal distribution for test-set metrics, an assumption which has not been extensively validated for deep neural networks (Dror et al., 2019).

In Figure 2, we observe large differences between the two representative metrics for OPP-115, Policy-Detection, PolicyIE-A, PrivacyQA and PolicyQA. For the first four of the aforementioned tasks, this is because of data imbalance resulting in the micro-average F_1 being significantly higher since it can be skewed by the metric of the majority class. For PolicyQA, this occurs because the EM metric requires exact matches and is therefore much stricter than the sample F_1 metric. Furthermore, we observe an exceptionally large standard deviation on PI-Extract metrics compared to other tasks. This can be attributed to data imbalance between the four subtasks of PI-Extract, with the NOT_COLLECT and NOT_SHARE subtasks having less than 100 total examples each.

We apply the arithmetic, geometric and harmonic means to aggregated metric means and standard deviations as shown in Table 3. With this, we observe the following general ranking of models

Model	A-Mean		G-Mean		H-Mean	
	μ	σ	μ	σ	μ	σ
BERT	67.5	1.1	64.6	0.9	61.1	0.6
RoBERTa	69.0	1.2	66.4	0.7	63.2	0.3
Legal-BERT	67.9	1.1	64.9	0.8	61.2	0.4
Legal-RoBERTa	68.5	1.3	65.7	0.8	62.3	0.4
PrivBERT	70.8	1.2	68.3	0.8	65.2	0.5

Table 3: Macro-aggregation of means (μ) and standard deviations (σ) per model using the arithmetic mean (A-Mean), geometric mean (G-Mean) and harmonic mean (H-Mean)

from best to worst: PrivBERT, RoBERTa, Legal-RoBERTa, Legal-BERT and BERT. Interestingly, models derived from RoBERTa generally outperformed models derived from BERT. Using the arithmetic mean for simplicity, we observe that PrivBERT outperforms all other models by 2 – 3%.

6 Discussion

With the PrivacyGLUE benchmark results, we revisit our privacy vs. legal language domain claim from Section 1 and discuss our model-pair agreement analysis for detecting PrivacyGLUE task examples where models benefited from domain spe-

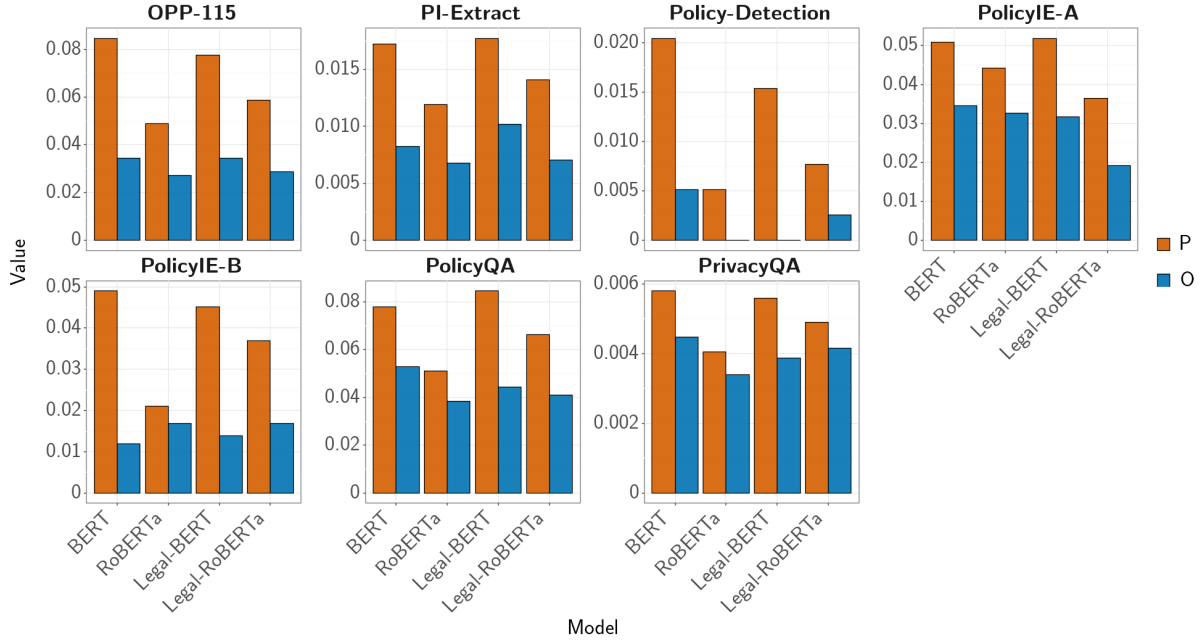


Figure 3: Model-pair agreement analysis of PrivBERT against other models over all PrivacyGLUE tasks; bars represent proportions of examples per model-pair and task which fell into categories P and O; all models on the x-axis are compared against PrivBERT

Category P	Category O
ID: 1978	ID: 33237
Question: Who can see my information?	Question: Could the wordscapes app contain malware?
Answer: We do not sell or rent your personal information to third parties for their marketing purposes without your explicit consent.	Answer: We encrypt the transmission of all information using secure socket layer technology (SSL).
Label: Relevant	Label: Relevant

Table 4: Test-set examples from PrivacyQA that fall under categories P and O for PrivBERT vs. BERT

cialization.

6.1 Privacy vs. legal language domain

We initially provided evidence from Figure 1 suggesting that the privacy language domain is distinct from the legal language domain. We believe that our PrivacyGLUE results further support this initial claim. If the privacy language domain was subsumed under the legal language domain, we could have observed Legal-RoBERTa and Legal-BERT performing competitively with PrivBERT. Instead, we observed that the legal models underperformed compared to both PrivBERT and RoBERTa, further indicating that the privacy language domain is distinct and requires its own NLP benchmark.

6.2 Model-pair agreement analysis

PrivBERT, the top performing model, differentiates itself from other models by its in-domain pretrain-

ing on the PrivaSeer corpus (Srinath et al., 2021). Therefore, we can infer that PrivBERT incorporated *knowledge* of privacy policies through its pretraining and became specialized for fine-tuning tasks in the privacy language domain. We investigate this specialization using model-pair agreement analysis to detect examples where PrivBERT had a competitive advantage over other models. Consequently, we detect examples where PrivBERT was disadvantaged due to its in-domain pretraining.

We compare $10 \times 10 = 100$ random seed combinations for all test-set pairs between PrivBERT and other models. Each prediction-pair can be classified into one of four mutually exclusive categories (B, P, O and N) shown below. Categories B and N represent examples that are either not challenging or too challenging for both PrivBERT and the other model respectively. Categories P and O are more interesting for us since they indicate examples where

PrivBERT had a competitive advantage and disadvantage over the other model respectively. Therefore, we focus on categories P and O in our analysis. We classify examples over all random seed combinations and take the majority occurrence for each category within its distribution.

Category B: Both PrivBERT and the other model were correct, i.e. (PrivBERT, Other Model)

Category P: PrivBERT was correct and the other model was wrong, i.e. (PrivBERT, $\neg$ Other Model)

Category O: Other model was correct and PrivBERT was wrong, i.e. ($\neg$ PrivBERT, Other Model)

Category N: Neither PrivBERT nor the other model was correct, i.e. ($\neg$ PrivBERT, $\neg$ Other Model)

Figure 3 shows a relative distribution of majority categories across model-pairs and PrivacyGLUE tasks. We observe that category P is always greater than category O, which correlates with PrivBERT outperforming all other models. We also observe that category P is often the greatest when compared against BERT, implying that PrivBERT has the most competitive advantage over BERT. Surprisingly, we also observe category O is often the greatest when compared against BERT, implying that BERT has the highest absolute advantage over PrivBERT. This is an insightful observation since we would have expected BERT to have the least competitive advantage given its lowest overall PrivacyGLUE performance.

To investigate PrivBERT’s competitive advantage and disadvantage against BERT, we extract several examples from categories P and O in the PrivacyQA task for brevity. Two interesting examples are listed in Table 4 and additional examples can be found in Table 8 in Appendix D. From Table 4, we speculate that PrivBERT specializes in example 1978 because it contains several privacy-specific terms such as "third parties" and "explicit consent". On the other hand, we speculate that BERT specializes in example 33237 since it contains more generic information regarding encryption and SSL, which also happens to be a topic in BERT’s Wikipedia pretraining corpus as seen in Figure 1 and Table 2.

Looking at further examples in Table 8, we can also observe that all sampled category P exam-

ples have the Relevant label while many sampled category O examples have the Irrelevant label. On further analysis of the PrivacyQA test-set, we find that 71% of category P examples have the Relevant label and 61% of category O samples have the Irrelevant label. We can infer that PrivBERT specializes in the minority Relevant label while BERT specializes in the majority Irrelevant label as the former label could require more privacy knowledge than the latter.

7 Conclusions and further work

In this paper, we describe the importance of data privacy in modern digital life and observe the lack of a NLP benchmark in the privacy language domain despite its distinctness. To address this, we propose PrivacyGLUE as the first comprehensive benchmark for measuring general language understanding in the privacy language domain. We release benchmark performances from the BERT, RoBERTa, Legal-BERT, Legal-RoBERTa and PrivBERT transformer language models. Our findings show that PrivBERT outperforms other models by an average of 2 – 3% over all PrivacyGLUE tasks, shedding light on the importance of in-domain pretraining for privacy policies. We apply model-pair agreement analysis to detect PrivacyGLUE examples where PrivBERT’s pretraining provides competitive advantage and disadvantage. By benchmarking holistic model performances, we believe PrivacyGLUE can accelerate NLP research into the privacy language domain and ultimately improve general language understanding of privacy policies for both humans and AI algorithms.

Looking forward, we envision several ways to further our study. Firstly, we intend to apply deep-learning explainability techniques such as Integrated Gradients (Sundararajan et al., 2017) on examples from Table 4, to explore PrivBERT’s and BERT’s token-level attention attributions for categories P and O. Additionally, we intend to benchmark large prompt-based transformer language models such as T5 (Raffel et al., 2020) and T0 (Sanh et al., 2022), as they incorporate large amounts of knowledge from the various sequence-to-sequence tasks that they were trained on. Finally, we plan to continue maintaining our PrivacyGLUE GitHub repository and host new model results from the community.

Limitations

To the best of our knowledge, our study has two main limitations. While we provide performances from transformer language models, our study does not provide human expert performances on PrivacyGLUE. This would have been a valuable contribution to judge how competitive language models are against human expertise. However, this limitation can be challenging to address due to the difficulty in finding experts and high costs for their services. Additionally, our study only focuses on English language privacy tasks and omits multilingual scenarios. Multilingual tasks would have been very interesting and relevant to explore, but also involve significant complexity since privacy experts for non-English languages may be harder to find.

Ethics Statement

Original work attribution

All datasets used to compose PrivacyGLUE are publicly available and originate from previous studies. We cite these studies in our paper and include references for them in our GitHub repository. Furthermore, we clearly illustrate how these datasets were used to form the PrivacyGLUE benchmark.

Social impact

PrivacyGLUE could be used to produce fine-tuned transformer language models, which could then be utilized in downstream applications to help users understand privacy policies and/or answer questions regarding them. We believe this could have a positive social impact as it would empower users to better understand lengthy and complex privacy policies. That being said, application developers should perform appropriate risk analyses when using fine-tuned transformer language models. Important points to consider include the varying performance ranges on PrivacyGLUE tasks and known examples of implicit bias, such as gender and racial bias, that transformer language models incorporate through their large-scale pretraining (Bender et al., 2021).

Software licensing

We release source code for PrivacyGLUE under version 3 of the GNU General Public License (GPL-3.0). We chose GPL-3.0 as it is a strong copyleft license that protects user freedoms such as the freedom to use, modify and distribute software.

References

- Wasi Ahmad, Jianfeng Chi, Tu Le, Thomas Norton, Yuan Tian, and Kai-Wei Chang. 2021. [Intent classification and slot filling for privacy policies](#). 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 4402–4417, Online. Association for Computational Linguistics.
- Wasi Ahmad, Jianfeng Chi, Yuan Tian, and Kai-Wei Chang. 2020. [PolicyQA: A reading comprehension dataset for privacy policies](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 743–749, Online. Association for Computational Linguistics.
- Ryan Amos, Gunes Acar, Elena Lucherini, Mihir Kshirsagar, Arvind Narayanan, and Jonathan Mayer. 2021. [Privacy Policies over Time: Curation and Analysis of a Million-Document Dataset](#). In *Proceedings of The Web Conference 2021*, WWW '21, page 22. Association for Computing Machinery.
- Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. [On the dangers of stochastic parrots: Can language models be too big?](#) In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, FAccT '21, page 610–623, New York, NY, USA. Association for Computing Machinery.
- Lukas Biewald. 2020. [Experiment tracking with weights and biases](#). Software available from wandb.com.
- Duc Bui, Kang G. Shin, Jong-Min Choi, and Junbum Shin. 2021. [Automated extraction and presentation of data practices in privacy policies](#). *Proceedings on Privacy Enhancing Technologies*, 2021(2):88–110.
- Ilias Chalkidis, Manos Fergadiotis, Prodromos Malakasiotis, Nikolaos Aletras, and Ion Androutsopoulos. 2020. [LEGAL-BERT: The muppets straight out of law school](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 2898–2904, Online. Association for Computational Linguistics.
- Ilias Chalkidis, Manos Fergadiotis, Prodromos Malakasiotis, and Ion Androutsopoulos. 2019. [Large-scale multi-label text classification on EU legislation](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 6314–6322, Florence, Italy. Association for Computational Linguistics.
- Ilias Chalkidis, Abhik Jana, Dirk Hartung, Michael Bommarito, Ion Androutsopoulos, Daniel Katz, and Nikolaos Aletras. 2022. [LexGLUE: A benchmark dataset for legal language understanding in English](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 4310–4330, Dublin, Ireland. Association for Computational Linguistics.

675	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding . In <i>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)</i> , pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.	735
676		736
677		737
678		738
679		
680		739
681		740
682		741
683		742
684	Rotem Dror, Segev Shlomov, and Roi Reichart. 2019. Deep dominance - how to properly compare deep neural models . In <i>Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics</i> , pages 2773–2785, Florence, Italy. Association for Computational Linguistics.	743
685		744
686		745
687		746
688		
689		
690	Sebastian Gehrmann, Tosin Adewumi, Karmanya Aggarwal, Pawan Sasanka Ammanamanchi, Anuoluwapo Aremu, Antoine Bosselut, Khyathi Raghavi Chandu, Miruna-Adriana Clinciu, Dipanjan Das, Kaustubh Dhole, Wanyu Du, Esin Durmus, Ondřej Dušek, Chris Chinenye Emezue, Varun Gangal, Cristina Garbacea, Tatsunori Hashimoto, Yufang Hou, Yacine Jernite, Harsh Jhamtani, Yangfeng Ji, Shailza Jolly, Mihir Kale, Dhruv Kumar, Faisal Ladhak, Aman Madaan, Mounica Maddela, Khyati Mahajan, Saad Mahamood, Bodhisattwa Prasad Majumder, Pedro Henrique Martins, Angelina McMillan-Major, Simon Mille, Emiel van Miltenburg, Moin Nadeem, Shashi Narayan, Vitaly Nikolaev, Andre Niyongabo Rubungo, Salomey Osei, Ankur Parikh, Laura Perez-Beltrachini, Niranjan Ramesh Rao, Vikas Raunak, Juan Diego Rodriguez, Sashank Santhanam, João Sedoc, Thibault Sellam, Samira Shaikh, Anastasia Shmorina, Marco Antonio Sobrevilla Cabeza, Hendrik Strobelt, Nishant Subramani, Wei Xu, Diyi Yang, Akhila Yerukola, and Jiawei Zhou. 2021. The GEM benchmark: Natural language generation, its evaluation and metrics . In <i>Proceedings of the 1st Workshop on Natural Language Generation, Evaluation, and Metrics (GEM 2021)</i> , pages 96–120, Online. Association for Computational Linguistics.	751
691		752
692		753
693		754
694		755
695		756
696		757
697		
698		
699		
700		
701		
702		
703		
704		
705		
706		
707		
708		
709		
710		
711		
712		
713		
714		
715		
716		
717	Saibo Geng, Rémi Lebre, and Karl Aberer. 2021. Legal transformer models may not always help . <i>CoRR</i> , abs/2109.06862.	758
718		759
719		
720	Oskar J Gstrein and Anne Beaulieu. 2022. How to protect privacy in a datafied society? a presentation of multiple legal and conceptual approaches. <i>Philosophy & Technology</i> , 35(1):1–38.	760
721		761
722		762
723		
724	Junjie Hu, Sebastian Ruder, Aditya Siddhant, Graham Neubig, Orhan Firat, and Melvin Johnson. 2020. XTREME: A massively multilingual multi-task benchmark for evaluating cross-lingual generalisation . In <i>Proceedings of the 37th International Conference on Machine Learning</i> , volume 119 of <i>Proceedings of Machine Learning Research</i> , pages 4411–4421. PMLR.	763
725		764
726		765
727		766
728		767
729		768
730		769
731		770
732	Taku Kudo and John Richardson. 2018. SentencePiece: A simple and language independent subword tokenizer and detokenizer for neural text processing . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing: System Demonstrations</i> , pages 66–71, Brussels, Belgium. Association for Computational Linguistics.	771
733		772
734		773
	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. <i>arXiv preprint arXiv:1907.11692</i> .	774
		775
		776
		777
		778
		779
		780
		781
		782
		783
	Ilya Loshchilov and Frank Hutter. 2017. Fixing weight decay regularization in adam . <i>CoRR</i> , abs/1711.05101.	784
		785
		786
		787
		788
		789
	Henry B Mann and Donald R Whitney. 1947. On a test of whether one of two random variables is stochastically larger than the other. <i>The annals of mathematical statistics</i> , pages 50–60.	
	Luca Mazzola, Andreas Waldis, Atreya Shankar, Diamantis Argyris, Alexander Denzler, and Michiel Van Roey. 2022. Privacy and customer’s education: Nlp for information resources suggestions and expert finder systems. In <i>HCI for Cybersecurity, Privacy and Trust</i> , pages 62–77, Cham. Springer International Publishing.	
	Aleecia M McDonald and Lorrie Faith Cranor. 2008. The cost of reading privacy policies. <i>Isjlp</i> , 4:543.	
	L. McInnes, J. Healy, and J. Melville. 2018. UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction . <i>ArXiv e-prints</i> .	
	Najmeh Mousavi Nejad, Pablo Jabat, Rostislav Nedelchev, Simon Scerri, and Damien Graux. 2020. Establishing a strong baseline for privacy policy classification. In <i>IFIP International Conference on ICT Systems Security and Privacy Protection</i> , pages 370–383. Springer.	
	Jonathan A Obar and Anne Oeldorf-Hirsch. 2020. The biggest lie on the internet: Ignoring the privacy policies and terms of service policies of social networking services. <i>Information, Communication & Society</i> , 23(1):128–147.	
	Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Kopf, Edward Yang, Zachary DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. 2019. Pytorch: An imperative style, high-performance deep learning library . In <i>Advances in Neural Information Processing Systems</i> , volume 32. Curran Associates, Inc.	
	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 . <i>Journal of Machine Learning Research</i> , 21(140):1–67.	

790	Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and	Alex Wang, Yada Pruksachatkun, Nikita Nangia, Aman-	848
791	Percy Liang. 2016. SQuAD: 100,000+ questions for	preet Singh, Julian Michael, Felix Hill, Omer Levy,	849
792	machine comprehension of text . In <i>Proceedings of</i>	and Samuel Bowman. 2019. Superglue: A stickier	850
793	<i>the 2016 Conference on Empirical Methods in Natural</i>	benchmark for general-purpose language understand-	851
794	<i>Language Processing</i> , pages 2383–2392, Austin,	ing systems . In <i>Advances in Neural Information</i>	852
795	Texas. Association for Computational Linguistics.	<i>Processing Systems</i> , volume 32. Curran Associates,	853
		Inc.	854
796	Abhilasha Ravichander, Alan W Black, Shomir Wilson,	Alex Wang, Amanpreet Singh, Julian Michael, Felix	855
797	Thomas Norton, and Norman Sadeh. 2019. Question	Hill, Omer Levy, and Samuel Bowman. 2018. GLUE:	856
798	answering for privacy policies: Combining compu-	A multi-task benchmark and analysis platform for nat-	857
799	tational and legal perspectives . In <i>Proceedings of</i>	ural language understanding . In <i>Proceedings of the</i>	858
800	<i>the 2019 Conference on Empirical Methods in Natural</i>	<i>2018 EMNLP Workshop BlackboxNLP: Analyzing</i>	859
801	<i>Language Processing and the 9th International</i>	<i>and Interpreting Neural Networks for NLP</i> , pages	860
802	<i>Joint Conference on Natural Language Processing</i>	353–355, Brussels, Belgium. Association for Com-	861
803	<i>(EMNLP-IJCNLP)</i> , pages 4949–4959, Hong Kong,	putational Linguistics.	862
804	China. Association for Computational Linguistics.		
805	Sebastian Ruder, Noah Constant, Jan Botha, Aditya Sid-	Wikimedia Foundation. 2022. Wikimedia downloads .	863
806	dhand, Orhan Firat, Jinlan Fu, Pengfei Liu, Junjie		
807	Hu, Dan Garrette, Graham Neubig, and Melvin John-	Shomir Wilson, Florian Schaub, Aswarth Abhilash	864
808	son. 2021. XTREME-R: Towards more challenging	Dara, Frederick Liu, Sushain Cherivirala, Pedro	865
809	and nuanced multilingual evaluation . In <i>Proceedings</i>	Giovanni Leon, Mads Schaarup Andersen, Sebas-	866
810	<i>of the 2021 Conference on Empirical Methods in</i>	tian Zimmeck, Kanthashree Mysore Sathyendra,	867
811	<i>Natural Language Processing</i> , pages 10215–10245,	N. Cameron Russell, Thomas B. Norton, Eduard	868
812	Online and Punta Cana, Dominican Republic. Asso-	Hovy, Joel Reidenberg, and Norman Sadeh. 2016.	869
813	ciation for Computational Linguistics.	The creation and analysis of a website privacy policy	870
		corpus . In <i>Proceedings of the 54th Annual Meet-</i>	871
814	Victor Sanh, Albert Webson, Colin Raffel, Stephen	<i>ing of the Association for Computational Linguistics</i>	872
815	Bach, Lintang Sutawika, Zaid Alyafeai, Antoine	<i>(Volume 1: Long Papers)</i> , pages 1330–1340, Berlin,	873
816	Chaffin, Arnaud Stiegler, Arun Raja, Manan Dey,	Germany. Association for Computational Linguistics.	874
817	M Saiful Bari, Canwen Xu, Urmish Thakker,		
818	Shanya Sharma Sharma, Eliza Szczechla, Taewoon	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien	875
819	Kim, Gunjan Chhablani, Nihal Nayak, Debajyoti	Chaumond, Clement Delangue, Anthony Moi, Pier-	876
820	Datta, Jonathan Chang, Mike Tian-Jian Jiang, Han	ric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz,	877
821	Wang, Matteo Manica, Sheng Shen, Zheng Xin Yong,	Joe Davison, Sam Shleifer, Patrick von Platen, Clara	878
822	Harshit Pandey, Rachel Bawden, Thomas Wang, Tr-	Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le	879
823	ishala Neeraj, Jos Rozen, Abheesht Sharma, An-	Scao, Sylvain Gugger, Mariama Drame, Quentin	880
824	ndrea Santilli, Thibault Fevry, Jason Alan Fries, Ryan	Lhoest, and Alexander M. Rush. 2020. Transform-	881
825	Teehan, Teven Le Scao, Stella Biderman, Leo Gao,	ers: State-of-the-art natural language processing . In	882
826	Thomas Wolf, and Alexander M Rush. 2022. Multi-	<i>Proceedings of the 2020 Conference on Empirical</i>	883
827	task prompted training enables zero-shot task gener-	<i>Methods in Natural Language Processing: System</i>	884
828	alization . In <i>International Conference on Learning</i>	<i>Demonstrations</i> , pages 38–45, Online. Association	885
829	<i>Representations</i> .	for Computational Linguistics.	886
830	Rico Sennrich, Barry Haddow, and Alexandra Birch.	Yonghui Wu, Mike Schuster, Zhifeng Chen, Quoc V Le,	887
831	2016. Neural machine translation of rare words with	Mohammad Norouzi, Wolfgang Macherey, Maxim	888
832	subword units . In <i>Proceedings of the 54th Annual</i>	Krikun, Yuan Cao, Qin Gao, Klaus Macherey, et al.	889
833	<i>Meeting of the Association for Computational Lin-</i>	2016. Google’s neural machine translation system:	890
834	<i>guistics (Volume 1: Long Papers)</i> , pages 1715–1725,	Bridging the gap between human and machine trans-	891
835	Berlin, Germany. Association for Computational Lin-	lation. <i>arXiv preprint arXiv:1609.08144</i> .	892
836	guistics.		
837	Mukund Srinath, Shomir Wilson, and C Lee Giles. 2021.	Sebastian Zimmeck, Peter Story, Daniel Smullen, Ab-	893
838	Privacy at scale: Introducing the PrivaSeer corpus	hilasha Ravichander, Ziqi Wang, Joel R Reidenberg,	894
839	of web privacy policies . In <i>Proceedings of the 59th</i>	N Cameron Russell, and Norman Sadeh. 2019. Maps:	895
840	<i>Annual Meeting of the Association for Computational</i>	Scaling privacy compliance analysis to a million apps.	896
841	<i>Linguistics and the 11th International Joint Confer-</i>	<i>Proc. Priv. Enhancing Tech.</i> , 2019:66.	897
842	<i>ence on Natural Language Processing (Volume 1:</i>		
843	<i>Long Papers)</i> , pages 6829–6839, Online. Association		
844	for Computational Linguistics.		
845	Mukund Sundararajan, Ankur Taly, and Qiqi Yan. 2017.		
846	Axiomatic attribution for deep networks . <i>CoRR</i> ,		
847	abs/1703.01365.		

A Detailed label information

Task	Labels
OPP-115	Data Retention, Data Security, Do Not Track, First Party Collection/Use, International and Specific Audiences Introductory/Generic, Policy Change, Practice not covered, Privacy contact information, Third Party Sharing/Collection, User Access, Edit and Deletion, User Choice/Control
PI-Extract	Subtask-I: {B,I}-COLLECT, 0 Subtask-II: {B,I}-NOT_COLLECT, 0 Subtask-III: {B,I}-NOT_SHARE, 0 Subtask-IV: {B,I}-SHARE, 0
Policy-Detection	Not Policy, Policy
PolicyIE-A	Other, data-collection-usage, data-security-protection, data-sharing-disclosure, data-storage-retention-deletion
PolicyIE-B	Subtask-I: {B,I}-data-protector, {B,I}-data-protected, {B,I}-data-collector, {B,I}-data-collected, {B,I}-data-receiver, {B,I}-data-retained, {B,I}-data-holder, {B,I}-data-provider, {B,I}-data-sharer, {B,I}-data-shared, storage-place, {B,I}-retention-period, {B,I}-protect-against, {B,I}-action, 0 Subtask-II: {B,I}-purpose-argument, {B,I}-polarity, {B,I}-method, {B,I}-condition-argument, 0
PrivacyQA	Irrelevant, Relevant

Table 5: Breakdown of labels for each PrivacyGLUE task; PolicyQA is omitted from this table since it is a reading comprehension task and does not have explicit labels like other tasks

B PrivacyGLUE task examples

Task	Input	Target
OPP-115	Revision Date: March 24th 2015	Introductory/Generic, Policy Change
PI-Extract	We may collect and share your IP address but not your email address with our business partners .	Subtask-I: 0 0 0 0 0 B-COLLECT I-COLLECT I-COLLECT 0 0 0 0 0 0 0 0 0 0 Subtask-II: 0 0 0 0 0 0 0 0 0 0 B-NOT_COLLECT I-NOT_COLLECT I-NOT_COLLECT 0 0 0 0 0 Subtask-III: 0 0 0 0 0 0 0 0 0 0 B-NOT_SHARE I-NOT_SHARE I-NOT_SHARE 0 0 0 0 0 Subtask-IV: 0 0 0 0 0 B-SHARE I-SHARE I-SHARE 0 0 0 0 0 0 0 0 0 0
Policy-Detection	Log in through another service: * Facebook * Google	Not Policy
PolicyIE-A	To backup and restore your Pocket AC camera log	data-collection-usage
PolicyIE-B	Access to your personal information is restricted .	Subtask-I: 0 0 B-data-provider B-data-protected I-data-protected 0 B-action 0 Subtask-II: B-method 0 0 0 0 0 0 0
PolicyQA	Question: How do they secure my data? Context: Users can visit our site anonymously	Answer: Users can visit our site anonymously
PrivacyQA	Question: What information will you collect about my usage? Answer: Location information	Relevant

Table 6: Representative examples of each PrivacyGLUE benchmark task

Task	Metric [†]	BERT	RoBERTa	Legal-BERT	Legal-RoBERTa	PrivBERT
OPP-115	m-F ₁	78.4 $\pm$ 0.6	79.5 $\pm$ 1.1	79.6 $\pm$ 1.0	79.1 $\pm$ 0.7	82.1$\pm$0.5
	μ -F ₁	84.0 $\pm$ 0.5	85.4 $\pm$ 0.5	84.3 $\pm$ 0.7	84.7 $\pm$ 0.3	87.2$\pm$0.4
PI-Extract	m-F ₁	60.0 $\pm$ 2.7	62.4 $\pm$ 4.4	59.5 $\pm$ 3.0	60.5 $\pm$ 3.9	66.4$\pm$3.4
	μ -F ₁	60.0 $\pm$ 2.7	62.4 $\pm$ 4.4	59.5 $\pm$ 3.0	60.5 $\pm$ 3.9	66.4$\pm$3.4
Policy-Detection	m-F ₁	85.3 $\pm$ 1.8	86.9 $\pm$ 1.3	86.6 $\pm$ 1.0	86.4 $\pm$ 2.0	87.3$\pm$1.1
	μ -F ₁	92.1 $\pm$ 1.2	92.7 $\pm$ 0.8	92.7 $\pm$ 0.5	92.4 $\pm$ 1.3	92.9$\pm$0.8
PolicyIE-A	m-F ₁	72.9 $\pm$ 1.7	73.2 $\pm$ 1.6	73.2 $\pm$ 1.5	73.5 $\pm$ 1.5	75.3$\pm$2.2
	μ -F ₁	84.7 $\pm$ 1.0	84.8 $\pm$ 0.6	84.7 $\pm$ 0.5	84.8 $\pm$ 0.3	86.2$\pm$1.0
PolicyIE-B	m-F ₁	50.3 $\pm$ 0.7	52.8 $\pm$ 0.6	51.5 $\pm$ 0.7	53.5 $\pm$ 0.5	55.4$\pm$0.7
	μ -F ₁	50.3 $\pm$ 0.5	54.5 $\pm$ 0.7	52.2 $\pm$ 1.0	53.6 $\pm$ 0.9	55.7$\pm$1.3
PolicyQA	s-F ₁	55.7 $\pm$ 0.5	57.4 $\pm$ 0.4	55.3 $\pm$ 0.7	56.3 $\pm$ 0.6	59.3$\pm$0.5
	EM	28.0 $\pm$ 0.9	30.0 $\pm$ 0.5	27.5 $\pm$ 0.6	28.6 $\pm$ 0.9	31.4$\pm$0.6
PrivacyQA	m-F ₁	53.6 $\pm$ 0.8	54.4 $\pm$ 0.3	53.6 $\pm$ 0.8	54.4 $\pm$ 0.5	55.3$\pm$0.6
	μ -F ₁	90.0 $\pm$ 0.1	90.2 $\pm$ 0.0	90.0 $\pm$ 0.1	90.2 $\pm$ 0.1	90.2$\pm$0.1

Table 7: Test-set results of the PrivacyGLUE benchmark; [†] m-F₁ refers to macro-average F₁, μ -F₁ refers to the micro-average F₁, s refers to sample-average F₁, EM refers to the exact match accuracy, metrics are reported as percentages with the following format: mean $\pm$ standard deviation

D Additional PrivacyQA examples from categories P and O

Category P	Category O
ID: 9227 Question: Will the app use my data for marketing purposes? Answer: We will never share with or sell the information gained through the use of Apple HealthKit, such as age, weight and heart rate data, to advertisers or other agencies without your authorization. Label: Relevant	ID: 8749 Question: Will my fitness coach share my information with others? Answer: Develop new services. Label: Irrelevant
ID: 10858 Question: What information will this app have access to of mine? Answer: Information you make available to us when you open a Keep account, as set out above; Label: Relevant	ID: 47271 Question: Who will have access to my medical information? Answer: 23andMe may share summary statistics, which do not identify any particular individual or contain individual-level information, with our qualified research collaborators. Label: Irrelevant
ID: 18704 Question: Does it share my personal information with others? Answer: We may also disclose Non-Identifiable Information: Label: Relevant	ID: 54904 Question: What data do you keep and for how long? Answer: We may keep activity data on a non-identifiable basis to improve our services. Label: Irrelevant
ID: 45935 Question: Will my test results be shared with any third party entities? Answer: 23andMe may share summary statistics, which do not identify any particular individual or contain individual-level information, with our qualified research collaborators. Label: Relevant	ID: 57239 Question: Do you sell any of our data? Answer: (c) Advertising partners: to enable the limited advertisements on our service, we may share a unique advertising identifier that is not attributable to you, with our third party advertising partners, and advertising service providers, along with certain technical data about you (your language preference, country, city, and device data), based on our legitimate interest. Label: Relevant
ID: 50467 Question: Can I delete my personally identifying information? Answer: (Account Deletion), we allow our customers to delete their accounts at any time. Label: Relevant	ID: 59334 Question: Does the app protect my account details from being accessed by other people? Answer: Note that chats with bots and Public Accounts, and communities are not end-to-end encrypted, but we do encrypt such messages when sent to the Viber servers and when sent from the Viber servers to the third party (the Public Account owner and/or additional third party tool (eg CRM solution) integrated by such owner). Label: Irrelevant

Table 8: Additional test-set examples from PrivacyQA that fall under categories P and O for PrivBERT vs. BERT; note that these examples are not paired and can therefore be compared in any order between categories

E Benchmark configuration

We run PrivacyGLUE benchmark tasks with the following configuration:

- We train all models for 20 epochs with a batch size of 16. We utilize a linear learning rate scheduler with a warmup ratio of 0.1 and peak learning rate of 3×10^{-5} . We utilize AdamW (Loshchilov and Hutter, 2017) as our optimizer. Finally, we monitor respective metrics on the validation datasets and utilize early stopping if the validation metric does not improve for 5 epochs.
- We use Python v3.8.13, CUDA v11.7, PyTorch v1.12.1 (Paszke et al., 2019) and Transformers v4.19.4 (Wolf et al., 2020) as our core software dependencies.
- We use the following HuggingFace model tags: bert-base-uncased, roberta-base, nlpueb/legal-bert-base-uncased, saibo/legal-roberta-base, mukund/privbert for BERT, RoBERTa, Legal-BERT, Legal-RoBERTa and PrivBERT respectively.
- We use 10 random seeds for each benchmark run, i.e. $\{0, 1, 2, 3, 4, 5, 6, 7, 8, 9\}$. This provides a distribution of results that can be used for statistical significance testing.
- We run the PrivacyGLUE benchmark on a Lambda workstation with $4 \times$ NVIDIA RTX A4000 (16 GB VRAM) GPUs for ~ 180 hours.
- We use Weights and Biases v0.13.3 (Biewald, 2020) to monitor model metrics during training and for intermediate report generation.