

EROS: Entity-Driven Controlled Policy Document Summarization

Joykirat Singh*, Sehban Fazili*, Rohan Jain, Md Shad Akhtar

Indraprastha Institute of Information Technology Delhi (IIIT Delhi), India.

{joykirat19166, sehban21143, rohan19095, shad.akhtar}@iiitd.ac.in

Abstract

Privacy policy documents have a crucial role in educating individuals about the collection, usage, and protection of users' personal data by organizations. However, they are notorious for their lengthy, complex, and convoluted language especially involving privacy-related entities. Hence, they pose a significant challenge to users who attempt to comprehend organization's data usage policy. In this paper, we propose to enhance the interpretability and readability of policy documents by using controlled abstractive summarization – we enforce the generated summaries to include critical privacy-related entities (e.g., *data* and *medium*) and organization's rationale (e.g., *target* and *reason*) in collecting those entities. To achieve this, we develop PD-Sum, a policy-document summarization dataset with marked privacy-related entity labels. Our proposed model, EROS, identifies critical entities through a span-based entity extraction model and employs them to control the information content of the summaries using proximal policy optimization (PPO). Comparison shows encouraging improvement over various baselines. Furthermore, we furnish qualitative and human evaluations to establish the efficacy of EROS.

Keywords: Controlled Summarization, Entity Extraction, Reinforcement Learning, PPO

1. Introduction

In today's Internet era, access to information has never been so convenient. Every day, an overwhelming amount of users are exploring the Internet horizon for entertainment or business purposes. Realizing an excellent opportunity to increase their customer base, many organizations offer their products or services in a convenient online setting. In majority of the cases, customers are required to sign up to acquire the services on offer and while doing so, they have to agree to the term-and-conditions (T&C) or policies of the service providers.

A privacy policy is a crucial component of any organization that allows it to legally collect, process, store, and/or distribute personal information. It outlines how an organization will handle personal data and how it will comply with applicable data protection laws and regulations. Little that they know, on many occasions, customers are, advertently or inadvertently, granting full access to their sensitive and private data (e.g., name, contact information, location, etc.) to the service providers without reading or understanding the privacy policy document. Moreover, some companies collect data with distributional rights as well and make a fortune by selling user's data to third parties without their realization but with their inadvertent consent^{1,2}. The primary reason for such ignorance on the user's part is their busy and packed schedule as well as lengthy and technical/legal language, which are usually difficult to comprehend by a common user.

Motivation and Problem Definition: Privacy policies are essential for both businesses and individuals. For businesses, having a privacy policy can protect them from legal issues related to data privacy and usage. On the other hand, it provides transparency to individuals about how their personal information will be managed and protected by the organizations; thus enabling them to make informed decisions prior to registering for the service. Despite its importance, very few users read these lengthy and non-trivial documents and fall prey to their inadvertent consent.

Summarizing these documents is a straightforward remedy of the lengthy document but it needs to ensure that every aspect of the data usage/management must also be present in the summary to make it useful. However, given the complicated nature of the policy document, it's non-trivial to obtain every critical privacy-related information in a summary. For instance, some policy documents define different data items (*viz.* name, age, contact details, etc.) at the beginning of the document but refrain from reporting the management of data items until the end of the document or in different paragraph or context; thus making it challenging for any summarization system to deal with such cases. In such cases, controlled abstractive summarization techniques (He et al., 2020; Liu and Chen, 2021; Zhang et al., 2023) can potentially enhance accessibility and transparency by generating concise and coherent summaries of policy documents. Acknowledging the severity of the problem, in this paper, we propose an **Entity-driven Controlled policy document Summarization** system *aka.* EROS. EROS operates in two stages: 1) it extracts various entities or data items and their ra-

*First two authors contributed equally

¹Brave under fire for alleged sale of copyrighted data

²Twitter fined 150m in US for selling user's data

(2021) explored fine-tuning strategies for language models like BERT and GPT, revealing substantial performance gains through limited labelled data utilization. Dong et al. (2021) proposed a hierarchical transformer model for summarizing lengthy documents, achieving leading results on multiple benchmarks. Zhang et al. (2020) devised PE-GASUS, leveraging gap sentence extraction and transformer-based gap filling pre-training, attaining state-of-the-art performance on various benchmarks. These studies depict the impactful role of pre-trained encoders in advancing summarization techniques.

Entity extraction is a fundamental task in information extraction. The problem has been modelled in multiple ways such as sequence labelling Lin et al. (2020); Bui et al. (2021), span level prediction Eberts and Ulges (2020); Zhong and Chen (2020); Zhu and Li (2022), question answering Li et al. (2020) as well as dependency parsing task Yu et al. (2020).

Recently, a paradigm shift has been observed from sequence labelling tasks to span-based prediction of entities. In span-based task, such as Eberts and Ulges (2020), all possible spans are selected and further classified whether that span represents an entity or not followed by relation classification, if required. To tackle the problem of exact spans being treated correctly and partial spans being treated incorrectly, Zhu and Li (2022), proposes a way to regularize span-based prediction tasks. The annotated spans are assigned a full probability and the nearby tokens are also assigned some probability of being correct. Zhong and Chen (2020) extracts entities along with relation instead of the traditional approach of extracting entities and then using the extracted entities for relation classification.

Reinforcement learning (RL) has gained traction in summarization (Wang and et al., 2018; Wan and et al., 2018). Rondeau and et al. (2018) introduced RL-driven translation with simulated human feedback. Liu and et al. (2020) addressed RL’s reward scarcity using human feedback. Gunasekara and et al. (2021) presented a versatile framework using RL for abstractive summarization through question-answering rewards. These efforts highlight RL’s effectiveness in improving summarization. In another work, He et al. (2020) proposed controlled summarization to let the user interact with the summary either in the form of keywords or direct prompts. Saito et al. (2020) develops a combination model consisting of a saliency model that extracts a token sequence from a source text and a seq-to-seq model that takes the sequence as an additional input text.

In comparison to existing works, our approach integrates pre-trained encoders, reinforcement learn-

ing, and modified loss functions. Our model incorporates a penalty mechanism in addition to a refined BART architecture for controlled summarization. In order to provide more exact and controlled summaries, this combined technique builds on the advantages of each component resulting in more precise and controlled summary generation. Additionally, our method uniquely incorporates insights from an Entity Extraction task, enhancing the model’s ability to capture information from policy documents.

3. Dataset Construction

To the best of our knowledge, the domain of privacy-driven policy document summarization has not been studied so far; hence, we recognize the need for a dataset to facilitate our research and thereby, develop $PD-Sum$, tailor-made dataset for the policy document summarization.

Data collection and Filtering: We collect policy documents of different websites curated by Amos et al. (2021). These documents outline how websites manage (i.e., collect, use, or disclose) personal information. After collection, we observed several issues and hence, applied a filter to discard policy documents as follows:

- Many websites have identical policy documents, we discard all but one.
- In case there are URLs linking to other websites, disregard them.
- Skip documents that lack meaningful/significant information or are incomplete.
- Refrain from including any policy content that is not relevant to the topic at hand.

Subsequent to the filtering process, 1921 policy document remains in $PD-Sum$.

Data Annotation: To facilitate the entity-driven controlled summarization, we need two sets of annotations: a) identification of privacy-related entities; and b) a summary of the document. Considering the users’ concerns, we identified five fundamental entities regarding data privacy and security and proposed a schema (c.f. Figure 1b) to capture their relationships:

- **Data:** It defines the type of information that an organization usually collects – *name, email, contact number, address, location, photos, system details, browsing history, search queries/patterns, keystrokes*, etc. Further, we observe that some of these data are compulsory as part of the service agreement, while others are optional and users can deny access without any interruption in service. We identify data items according to these two categories, i.e., ‘*data compulsory*’ and ‘*data optional*’. Further, we identify

data items as '*data others*' which are not associated with a target requesting the data explicitly may not highlight its usage, e.g., the data item '*personal information*' in the sentence '*We may share personal information with our clients*'.

- **Source:** It signifies the provider of the information. While a majority of the time, the user (e.g., '*you*') is the direct source, in some cases, the source can be indirect, e.g., "*inviting your friends to join the website by sharing contact information*" (*friends* will be source indirect), "*requiring someone else information to ship products to their address*" (*someone else* is source indirect). These entities are very low in number and are majorly associated with sharing someone else information.
- **Medium:** It defines the way data is collected such as '*while visiting the website*', '*responding to a survey*', '*filling a form*', etc.
- **Target:** It specifies who will consume the data. Similar to the source entity, a target can be direct (*the organization itself*) or indirect (*any third-party vendor outside the organization*). Though the direct target is somewhat benign as the users know their data are being used for some specific purposes by the service provider, the indirect target can be extremely detrimental as there is no transparency about the usage of data in an unknown capacity.
- **Reason:** It clarifies the purpose of data collection by the parent organization such as '*improving customer experience*'. We observe that with indirect targets, reasons are usually hidden or extremely vague.

Following the above schema, two annotators³ with good English proficiency annotated the whole dataset using LabelBox⁴ as the annotation tool. At first, we tokenize the sentences using NLTK tokenizer Bird et al. (2009), and subsequently, for each identified entity, we record their start and end indices as span. Further, to ensure the consistency of annotations between them, the annotators independently annotated a small set of documents in the pilot phase and discussed their common understanding. Next, they annotated 10 documents separately and achieved a Cohen Kappa inter-annotator agreement score of 0.74. Subsequently, we manually annotate 150 documents having 8000 sentences with 9094 distinct entity labels. In the next step, we learn an entity-recognition model trained on the annotated 150 documents to obtain pseudo entity labels for the remaining 1771 documents. An illustrative example of annotation for a paragraph of the document is presented in Figure 1c. Table 1 lists the distribution of the entities in PD-Sum.

³Annotators were undergraduate student volunteers and in the age group of 20-30.

⁴<https://labelbox.com/>

	Train	Test
Documents	#Doc	1536 385
	#Sent in Doc	116251 27487
	Avg. token/Doc	1455.6 1375.82
	Total entities (Doc)	130856 31146
	– Data other	26836 6458
	– Data Compulsory	7715 1866
	– Data Optional	274 68
	– Reason	10720 2451
	– Medium	11924 2922
	– Target Direct	33927 8159
	– Target Indirect	12421 3129
	– Source Direct	27002 6080
	– Source Indirect	37 13
Summary	#Sent in Summ	18691 4500
	Avg. token/Summ	198.9 192.46
	Total entities (Summ)	28197 6952
	– Data other	6244 1580
	– Data Compulsory	1784 497
	– Data Optional	0 0
	– Reason	1449 338
	– Medium	2469 586
	– Target Direct	7991 1938
	– Target Indirect	2453 648
	– Source Direct	5807 1365
	– Source Indirect	0 0

Table 1: Dataset statistics of PD-Sum.

In the second stage, we focus on manually annotating 1921 documents with a brief yet informative summary. We provide clear guidelines to our annotators, encouraging them to extract key data-related phrases and concepts from the document while preserving the essence of the content in a concise and clear manner. The goal is to keep the summaries as crisp and to the point as possible. We emphasize using simple and everyday language, steering away from complex or overly technical terminology. Our aim is to make the summaries accessible and easily understandable to a broad audience. This approach ensures that the information conveyed in the summaries remains approachable and digestible. In terms of document length, we've found that policy documents typically contain an average of 1500 tokens. On the other hand, the annotated summaries, following our guidelines, average about 200 tokens. This significant reduction in token count maintains a balance between providing necessary information and keeping the summaries succinct and user-friendly.

4. Proposed Methodology - EROS

Our proposed controlled summarization model, EROS, is depicted in Figure 2. It works in two stages. At first, we train an entity extraction model that aims to predicts all spans of entities in a given document. This model employ BERT-based span prediction framework with the contrastive loss. Additionally, we supplement the prediction via an entity classification model in a joint learning setup. In the second stage, we employ a BART-based summarizer model to generate a summary. To assess the quality of generated summary and to ensure

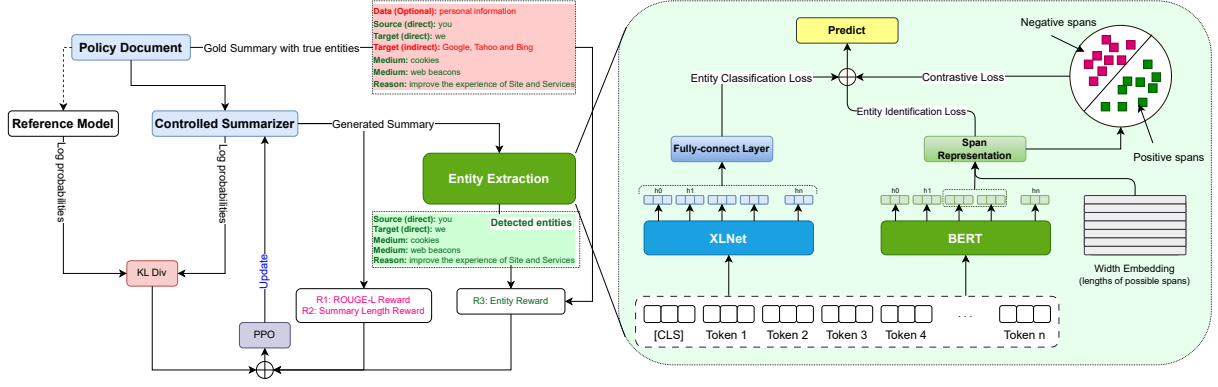


Figure 2: **Left:** Proposed Model for EROS. The reference model is a frozen pre-trained BART-based model with modified loss. We initialize the controlled summarization model in a similar way, which is subsequently updated through a PPO framework on a combination of rewards and KL-divergence loss. **Right:** Entity extraction model jointly learns a entity classification and entity identification module with the assistance of contrastive loss. Further to minimize the effect of false positives in identification, we supplement it with a entity classification module in a joint framework.

induction of critical entity information in the summary, we introduce our pretrained entity extraction module in the pipeline. The extracted entities are then compared with the true (gold) entities of the reference summary and an entity reward is computed. Moreover, we compute two other rewards (ROUGE-L and summary length) to ensure a relevant and concise summary. The accumulated reward is then added with the KL-divergence score between the log-probabilities of the controlled summarizer and a reference model (a BART-based summarizer trained on a huge out-of-domain corpus). Subsequently, we employ PPO to update the controlled summarizer. This process repeats for a few step till (near-)convergence. In the following subsections, we elaborate on these steps.

Entity extraction module

Recent years have seen a paradigm shift in the task of entity recognition from token-level tagging—which conceptualizes it as a sequence labelling task—to span-level prediction [Fu et al. \(2021\)](#).

A span in a sentence is represented by the start and end token of a sentence. Given a sentence $X = \{x_1, \dots, x_n\}$ with n tokens, we define a span of an entity as $s_i^y = \{x_{b_i}, x_{b_i+1}, \dots, x_{e_i}\}$, where b_i and e_i denote the start and end index of the span s_i with a corresponding tag $y \in \{\text{source-direct}, \text{source-indirect}, \text{data-optional}, \text{data-compulsory}, \text{medium}, \text{target-direct}, \text{target-indirect}, \text{reason}\}$.

Spans come in different lengths. To avoid overfitting for a particular length, we adopt an enumeration strategy, where all the possible m spans with a maximum length l are being considered as valid spans for predicting entities. For example, in the sentence, “we will collect name” with maximum span length as 4, possible spans are: $s_1^y = \{x_1, x_1\}$, $s_2^y = \{x_1, x_2\}$, $s_3^y = \{x_1, x_3\}$,

$s_4^y = \{x_1, x_4\}$, $s_5^y = \{x_2, x_2\}$, $s_6^y = \{x_2, x_3\}$, $s_7^y = \{x_2, x_4\}$, $s_8^y = \{x_3, x_3\}$, $s_9^y = \{x_3, x_4\}$, and $s_{10}^y = \{x_4, x_4\}$. From gold labels, we know that out of these 10 spans, only s_1 and s_4 are valid spans; therefore, their y labels will be source-target and data-compulsory, respectively. For all other spans, the y labels will be “invalid (0)”.

We obtain token representation from BERT and subsequently, compute span embeddings as $z_i^b = [\mathbf{h}_{b_i}; \mathbf{h}_{e_i}]$. Additionally, to provide information regarding the width of each span, we induce a learnable width encoding vector z_i^w according to their width. Thus, the final representation becomes $s_i = [z_i^b; z_i^w]$. Our initial experiments showed encouraging results; however, we also observe a significant number of false-positive, especially, for a sentence with no entity at all, in our predictions. To mitigate such issue, we incorporate a binary classification task to identify if the sentence contains an entity in a joint learning framework.

Contrastive Loss: We compute a similarity score between pairs of input spans, and then minimize the distance between similar pairs while maximizing the distance between dissimilar pairs using the contrastive loss:

$$P(y | s_i) = \frac{\text{score}(s_i, y)}{\sum_{y' \in Y} \text{score}(s_i, y')} \quad (1)$$

where $\text{score}(s_i, y_k) = \exp(s_i^T y_k)$, is a function that measures the compatibility between a learnable label representation of the class k y_k and span s_i .

Span Prediction: Finally, the span representations s_i are fed into a softmax function to get the probability considering the label y . For optimization, we combine the three losses – entity identification cross-entropy loss (ℓ_1), binary classification loss

(ℓ_2), and the contrastive loss (ℓ_3) through the following weighting mechanism, where α_1 , α_2 , and α_3 are hyperparameters.

$$\mathcal{L}^e = \alpha_1 \ell_1 + \alpha_2 \ell_2 + \alpha_3 \ell_3, \quad \text{where } \sum \alpha_i = 1 \quad (2)$$

Entity-Driven Controlled Summarization

As the foundational model, we employ BART in our experiment. Further, we induced a modified loss function specially designed for the entity-driven controlled summary generation. To elaborate, we obtain gold entities (e_i) from the PD-Sum dataset and integrate them into the loss function of the BART model. This entails augmenting the traditional cross-entropy loss with a penalty component derived from the extracted entities. It enables the BART model to comprehend the presence and importance of entities in the summaries, thereby refining its summary generation capabilities while maintaining control over the process. Mathematically it can be seen as follows:

$$\begin{aligned} \text{CE} &= - \sum (y \cdot \log(x)) \\ \text{TP} &= \sum_{e_i} (1.0 - \text{step}(e_i \in S_G)) \\ \mathcal{L}^s &= \lambda \cdot \text{CE} + (1 - \lambda) \cdot \text{TP} \end{aligned} \quad (3)$$

where CE, TP, λ , and S_G are the cross-entropy, token penalty score, weight of the loss, and the generated summary, respectively. We compute TP by penalizing the model for each missing entity e_i in S_G . The step function will return 1 only if the entity is part of the summary, otherwise, a value of 0 will be returned.

Further to supplement the controlled summary generation process, we adopted a feedback mechanism, in the form of reinforcement learning, to reward/penalize the model for inducing/not-inducing the privacy-related entities in the summaries. We use proximal policy optimization (PPO) to enforce the model to improve the generation quality. First introduced by Schulman et al. (2017), PPO refines policy adjustments by combining ratio-based enhancement with a clipped surrogate objective; thus, ensuring controlled updates. Incorporating an auxiliary value function, PPO enhances policy updates by estimating advantages and rewards more accurately, particularly in complex scenarios.

The proposed model shown in Figure 2, contains a policy model (i.e., the controlled summarizer model that is being trained), a reference model, a reward model, and a value function. The value function is used to describe the reward at timestep t . On the other hand, the reference model is used to calculate the KL divergence between the original model and the policy model. The main idea is to ensure that the active model does not deviate a lot from its original distribution.

Training <i>EEPD</i>		Training <i>EROS</i>		Generating Text (<i>EROS</i>)	
Hyper-parameter	Value	Hyper-parameter	Value	Hyper-parameter	Value
<i>n_class</i>	10	<i>warmup</i>	0.1	<i>max_seq_length</i>	1024
<i>bert_dropout</i>	0.2	<i>learning_rate</i>	5.41e-6	<i>min_new_tokens</i>	200
<i>xlnet_dropout</i>	0.2	<i>adaptive_kl_coef</i>	True	<i>top_p</i>	0.9
<i>LR</i>	1e-5	<i>gamma</i>	0.99	<i>do_sample</i>	True
<i>maxlen</i>	512	<i>batch_size</i>	8	<i>top_k</i>	10
<i>maxnorm</i>	1.0	<i>mini_batch_size</i>	2	<i>use_cache</i>	True
<i>batchSize</i>	4			<i>num_beams</i>	1
<i>max_spanLen</i>	10				
<i>spanLen_emb_dim</i>	300				
α_1	0.5				
α_2	0.25				
α_3	0.25				
<i>l</i>	10				

Table 2: Hyperparameters for training *EEPD*, *EROS*, and generating summary from *EROS*

Reward calculation: We compute three rewards to maintain the coverage, conciseness, and relevance in the summary. The coverage reward ensures the readability of the generated summary –we compute the ROUGE-L score, which is based on the longest common subsequence (LCS) between two sequences, as the first reward ($R1 = \text{ROUGE-L}(S_G, S_R)$). A longer LCS indicates that the generated summary conveys similar meaning and concepts as the reference summary. The conciseness reward ($R2$) limits the model to generate an adequate length summary and avoid generating lengthy jargon.

$$R_2 = \frac{1 - |(\text{len}(S_G) - \text{len}(S_R))|}{\max(\text{len}(S_G), \text{len}(S_R))} \quad (4)$$

Finally, we compute the entity reward ($R3$) as follows: Let E_{total} be the total number of entities predicted from the generated summary, and $E_{correct}$ and $E_{incorrect}$ be the number of entities present and not present in the gold summary respectively.

$$R_3 = \frac{E_{correct} - \beta * E_{incorrect}}{E_{total}} \quad (5)$$

where β is a negative factor for penalizing incorrect entities. We set $\beta = 0.3$ for our experiments.

5. Experiments and Results

We train *EROS* on 1536 documents, while we use 385 documents for evaluating the performance. A summary of the hyperparameters used for training is listed in Table 2. We evaluate both components of our model separately and their results are furnished in Tables 3 and 4, respectively – we compute precision, recall, and F1-score for the entity extraction module, we measure the quality of summarization through BLEU, ROUGE, METEOR, and BERTScore. We also perform extensive comparative analysis against the following baselines for both models. In all cases, we re-train/fine-tune these models on PD-Sum.

Baselines:

Model	Rouge			BLEU				METEOR	BS
	R1	R2	RL	B1	B2	B3	B4		
Extractive Oracle	0.43	0.30	0.42	0.14	0.12	0.11	0.10	0.25	0.881
Bert2Bert	0.25	0.04	0.22	0.22	0.08	0.10	0.16	0.25	0.818
T5-Summarizer	0.44	0.24	0.42	0.32	0.24	0.19	0.17	0.34	0.877
PEGASUS	0.35	0.17	0.32	0.18	0.12	0.09	0.08	0.23	0.859
BART	0.46	0.29	0.44	0.31	0.25	0.22	0.20	0.33	0.882
CTRL-SUM [bigpatent]	0.200	0.037	0.180	0.174	0.067	0.022	0.008	0.142	0.798
CTRL-SUM [cnndm]	0.340	0.148	0.318	0.266	0.164	0.121	0.099	0.246	0.837
CTRL-SUM [PD-Sum]	0.355	0.124	0.327	0.244	0.134	0.077	0.044	0.216	0.848
T5-Loss	0.45	0.26	0.43	0.32	0.24	0.19	0.17	0.34	0.877
PEGASUS-Loss	0.38	0.21	0.36	0.23	0.17	0.14	0.12	0.26	0.865
BART-Loss	0.48	0.31	0.46	0.31	0.25	0.22	0.20	0.34	0.889
EROS	0.519	0.341	0.500	0.424	0.337	0.292	0.262	0.438	0.891
EROS w/o R1	0.434	0.232	0.412	0.351	0.241	0.185	0.148	0.33	0.871
EROS w/o R2	0.414	0.211	0.395	0.328	0.217	0.16	0.124	0.295	0.8648
EROS w/o R3	0.510	0.324	0.491	0.412	0.322	0.275	0.244	0.415	0.888

Table 3: Rouge, Bleu, Meteor and BertScore scores for baselines and our proposed **EROS** model. The term “Loss” signifies application of a customized loss function (c.f. equation 3). CTRL-SUM [D] signifies that the model is trained on dataset D .

- **Entity Extraction:** We evaluate **EEPD** against six entity extraction models covering both sequence-labelling and span-based frameworks: **Sequence-labelling:** Tagging each token as **Begin**, **Intermedite**, or **Other** – **BERT** Devlin et al. (2018), **SpanBERT** Joshi et al. (2020), and **PrivBERT** Srinath et al. (2021). **Span-based:** Start/end index-based extraction models – **Boundary Smoothing** Zhu and Li (2022), **PURE** Zhong and Chen (2020), **SPERT** Eberts and Ulges (2020).
- **Entity-Driven Controlled Summarization:** For the controlled summarization model, we compare **EROS**’s effectiveness against the following baseline approaches: **Extractive Oracle** (Hirao et al., 2017): It employs extractive methods to gather essential information from the source text for summarization. This model offers efficiency by avoiding new sentence generation and at the expense of missing overall context and flow. **PEGASUS** (Zhang et al., 2020): Pre-trained transformer-based sequence-to-sequence architecture designed by Google AI. The model uses a novel pre-training objective known as “gap-sentences generation”. **CTRL-SUM** (He et al., 2020): It uses control tokens at inference time to enable user to control the summary. In the original work, it has been trained on two huge datasets – BIGPATENT patent documents (Sharma et al., 2019) and CNN/Dailymail news articles (Hermann et al., 2015). Further, in this work, we also train CTRL-SUM on PD-Sum. In addition, we also compare with **Bert2Bert** (Chen et al., 2022), **T5** (Raffel et al., 2020), and **BART** (Lewis et al., 2019) summarizers.

Result Analysis: Table 4 contains the comparative result of **EEPD** and various baselines. PrivBert

		Precision	Recall	F1-score
BIO	BERT	0.31	0.38	0.34
	SpanBERT	0.31	0.39	0.35
	PrivBERT	0.37	0.44	0.40
Span Based	SPERT	0.10	0.68	0.17
	Boundary Smoothing	0.40	0.48	0.44
	PURE	0.35	0.12	0.17
	SpanNER	0.49	0.56	0.52
	SpanNER + Identification	0.47	0.66	0.55
	EEPD	0.54	0.62	0.58
	– Identification	0.48	0.57	0.52

Table 4: Comparative results of entity extraction module.

reports the best F1-score of 0.40 in the sequence-labelling framework (i.e., in a BIO setup), whereas, BoundarySmoothing yields +4% better F1-score at 0.44 in the span-based setting. Further, SpanNER, with the identification module, records the best F1-score of 0.58 among all baselines. In comparison, **EEPD** reports the state-of-the-art performance at 0.58 F1-score – an increment of +3% over the best baseline. We also observe the effect of the entity identification module on the overall performance – a decrement of –6% is observed on removing the identification component from **EEPD**.

For the controlled summarization task, we furnish the results in Table 3. We compute the traditional ROUGE, METEOR, BLEU and BertScore scores to evaluate the generated summaries. Pre-trained Ctrl-sum He et al. (2020) failed to perform well when compared to fine-tuned models. Among all baselines (except with the modified loss, TP (c.f. Section 4)), BART reports the best performance across the three metrics – ROUGE-L score (0.44), BLEU-4 (0.20), METEOR (0.33) and BertScore(0.882). Further, we observe that the incorporation of modified TP loss obtains compa-

Ref Summary: When you register with the Site or use any of our Services, you may be asked to provide us with Personal Information. It is entirely optional to provide this information. If you do not provide the requested information, you may be unable to use some or all of the Site's features. To improve the Site and Services, we use cookies (a small text file placed on your computer to identify your computer and browser) and Web beacons (electronic files placed on a Web site that monitor usage). If you delete or disable cookies, some features of the Site or Services may not function properly. You have the option of making some of your Personal Information available to others. Users of the Site and commercial search engines such as Google, Yahoo!, and Bing may access the information you choose to make available. Some third-party services may provide us with information from your accounts, allowing us to improve and personalise your experience on the Site. When you visit our Site, we may allow third-party companies to serve ads and/or collect certain anonymous information....
Entities in Ref summary: Data compulsory: [e-mail, personally identifiable information]; Data others: [personal information]; Source Direct: [you, info@day-finder.com]; Target Direct: [we, us]; Target Indirect: [service providers, third parties]; Medium: [register with the Site or use any of our Services, cookies, web beacons, visit our site]; Reason: [improve the site and services, improve and personalise your experience on the Site, Usage Data regarding the Site and Services].
Entities in BART summary: Data compulsory: [None]; Data others: [None]; Source Direct: [you, info@day-finder.com]; Target Direct: [we, us]; Target Indirect: [None]; Medium: [using the site or services, cookies, web beacons, visit our site]; Reason: [improve the experience of the site and services].
Entities in BART-Loss summary: Data compulsory: [e-mail, personally identifiable information]; Data others: [None]; Source Direct: [you]; Target Direct: [we, us]; Target Indirect: [service providers]; Medium: [cookies, web beacons]; Reason: [None].
Entities in EROS w/o R3 summary: Data compulsory: [e-mail, personally identifiable information]; Data others: [personal information]; Source Direct: [you]; Target Direct: [we, us]; Target Indirect: [service providers]; Medium: [register with the Site or use any of our Services, cookies, web beacons, visit our site]; Reason: [None].
Entities in EROS summary: Data compulsory: [e-mail, personally identifiable information]; Data others: [personal information]; Source Direct: [you, info@day-finder.com]; Target Direct: [we, us]; Target Indirect: [service providers, third parties]; Medium: [register with the Site or use any of our Services, cookies, web beacons, visit our site]; Reason: [Usage Data regarding the Site and Services].

Table 5: Qualitative analysis of generated summaries. Due to space constraints, we could not provide the generated summaries of models. source direct, target direct, medium - blue, reason, data, data compulsory, target indirect, data optional. Best viewed in color.

able results in the majority of the cases and better several setups – ROUGE-L: T5, BART, and PEGASUS improved; BLEU-4: PEGASUS improved; METEOR: BART and PEGASUS improved and BertScore: BART and PEGASUS improved. On the other hand, EROS yields the best scores across all metrics – improvement of +0.04 in ROUGE-L (0.50), +0.062 in BLEU-4 (0.262), +0.098 in METEOR (0.438) and +0.0017 in BertScore (0.8907).

Model	En / Sum	Len dev
Gold Summary	18.057	-
BART	10.329	-0.2308
BART-Loss	10.355	-0.266
EROS	14.436	+0.163

Table 6: Rate of entities and conciseness of the generated against the gold summary.

Table 6 shows a comparative analysis on the conciseness and the rate of captured entities in the generated summaries. We observe +4% increase in rate of entities in EROS against the two baselines.

Human Evaluation: We performed a human evaluation on a subset of randomly chosen samples from the PD-Sum’s test set. We compare the summaries of EROS and two baselines i.e., BART and BART-Loss. We ask our evaluators to assess the generated summaries against the reference summaries on four parameters – the informativeness of the summary (INF), its conciseness (CON), its fluency and grammatical correctness (FL), and the inclusion of relevant entities (EC). The first two metrics assess the quantity of the information content in the generated summary, whereas the third metric ensures the linguistics quality of the summary. Additionally, the last metric explicitly

Model	INFO	CON	FLU	EC
BART	3.0	3.35	3.75	2.90
BART-Loss	3.75	3.70	4.0	3.40
EROS	4.20	3.15	4.05	4.15

Table 7: Human evaluation on *informativeness*, *conciseness*, *fluency*, and *entity coverage*.

evaluates the presence of privacy-related entities in the summary. For each parameter, all evaluators assign a rating on a scale of 1 (worst) to 5 (best) based on the quality of the summaries. Subsequently, we aggregate the scores through averaging and report the observations in Table 7. We observe that EROS outperforms the other two baseline models in three out of four metrics. It records a comparatively inferior score for conciseness, suggesting that EROS’s summaries are relatively lengthier than others. However, it is better in informativeness, grammatical correctness, and inclusion of relevant entities.

6. Conclusion

In this work, we presented a novel approach for abstractive summarization of privacy policy documents. Our approach aimed to address the challenge of generating controlled and informative summaries that capture the essence of complex privacy policies. To achieve this, we introduced a customized loss function and incorporated a reinforcement learning framework, enabling us to optimize the relevance of the generated summaries. To facilitate the evaluation and advancement of research in this domain, we also introduced a new dataset for controlled summarization generation. The experimental results obtained from our comprehen-

sive evaluations highlight the effectiveness of the proposed approach. Our model achieved state-of-the-art performance on the custom dataset. The controlled generation of summaries allows for improved accessibility and transparency for users, enabling them to quickly grasp the key points of privacy policies without getting overwhelmed by excessive information. The findings of our work demonstrate the potential of our approach to make a significant impact in the field of privacy policy summarization. By addressing the critical need for concise and user-friendly representations of privacy policies, we contribute to enhancing user understanding. The implications of our work extend to various domains where privacy policies play a crucial role, including data protection, online services, and legal compliance.

Limitation

Our research initiative necessitates the utilization of proprietary and sensitive data, encompassing company privacy policies, to train our specialized model. This dataset is comprehensive, containing detailed policies of various companies. This inclusion is vital for our model to effectively understand and learn the intricacies of privacy policies and their implementations. However, the training process for our model is computationally intensive, requiring substantial computing resources and time. The utilization of a reinforcement learning approach, while effective, significantly extends the duration needed for the training phase. This computational demand poses a notable challenge in terms of time and resource allocation. We understand that the relevant entities, such as data items, target usage, etc., are extremely critical in comprehending the policy document and hence, any model must include these components in the generated summaries. However, in some cases, the generated summaries may not be exhaustive in terms of these non-trivial entities and care must be taken prior to agreeing to the terms and conditions of the policy document.

Ethical Consideration

In this section, we address the ethical concerns and safeguards associated with our research on privacy policy text summarization. Ensuring the ethical conduct of our research is of paramount importance, given the sensitive nature of the data and the potential implications for privacy.

- **Data Collection and Usage:** Data was obtained from publicly available sources or with explicit permission, and sensitive information was excluded from our dataset. We respect any limitations set by data sources concerning data usage.

- **Transparency and Explainability:** We recognize the importance of providing transparent and understandable summaries. Efforts were made to ensure that our summarization process can be explained to users, and address any algorithmic biases that may arise, striving for fairness.
- **Fairness and Bias Mitigation:** We are mindful of potential biases and employ measures to mitigate them, particularly those related to gender, race, and other sensitive attributes. Fairness in our summaries is a core consideration.
- **Compliance with Regulations:** Our research adhered to all relevant data protection laws and regulations, such as GDPR, HIPAA, or regional data privacy laws, to ensure the ethical use of data.
- **Potential Risks and Mitigation Strategies:** We acknowledge the potential risks, such as unintended disclosure of sensitive information. To mitigate these risks, we implement strict controls and perform thorough risk assessments.

By taking these ethical considerations into account, we aimed to conduct our research on privacy policy text summarization with the utmost respect for privacy, transparency, and ethical integrity.

Acknowledgement

Authors acknowledge the partial support of Infosys foundation through Center for AI (CAI)-IIIT Delhi.

Bibliographical References

- Ryan Amos, Gunes Acar, Eli 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*, pages 2165–2176.
- Steven Bird, Ewan Klein, and Edward Loper. 2009. *Natural language processing with Python: analyzing text with the natural language toolkit*. "O'Reilly Media, Inc."
- Duc Bui, Kang G Shin, Jong-Min Choi, and Junbum Shin. 2021. Automated extraction and presentation of data practices in privacy policies. *Proc. Priv. Enhancing Technol.*, 2021(2):88–110.
- Cheng Chen, Yichun Yin, Lifeng Shang, Xin Jiang, Yujia Qin, Fengyu Wang, Zhi Wang, Xiao Chen, Zhiyuan Liu, and Qun Liu. 2022. [bert2BERT: Towards reusable pretrained language models](#). In

- Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2134–2148, Dublin, Ireland. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.
- Li Dong, Furu Wei, Ming Zhou, and Ke Xu. 2021. Hierarchical transformers for long document summarization. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics*.
- Markus Eberts and Adrian Ulges. 2020. [Span-based joint entity and relation extraction with transformer pre-training](#).
- Jinlan Fu, Xuanjing Huang, and Pengfei Liu. 2021. [SpanNER: Named entity re-/recognition as span prediction](#). 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 7183–7195, Online. Association for Computational Linguistics.
- Chamara Gunasekara and et al. 2021. Using question answering rewards to improve abstractive summarization. In *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 518–526.
- Junxian He, Wojciech Kryscinski, Bryan McCann, Nazneen Rajani, and Caiming Xiong. 2020. [Ctrlsum: Towards generic controllable text summarization](#). *ArXiv*, abs/2012.04281.
- Karl Moritz Hermann, Tomas Kocisky, Edward Grefenstette, Lasse Espeholt, Will Kay, Mustafa Suleyman, and Phil Blunsom. 2015. Teaching machines to read and comprehend. *Advances in neural information processing systems*, 28.
- Tsutomu Hirao, Masaaki Nishino, Jun Suzuki, and Masaaki Nagata. 2017. [Enumeration of extractive oracle summaries](#). In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 1, Long Papers*, pages 386–396, Valencia, Spain. Association for Computational Linguistics.
- Mandar Joshi, Danqi Chen, Yinhan Liu, Daniel S Weld, Luke Zettlemoyer, and Omer Levy. 2020. Spanbert: Improving pre-training by representing and predicting spans. *Transactions of the Association for Computational Linguistics*, 8:64–77.
- Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Ves Stoyanov, and Luke Zettlemoyer. 2019. [Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension](#).
- Xiaoya Li, Jingrong Feng, Yuxian Meng, Qinghong Han, Fei Wu, and Jiwei Li. 2020. [A unified MRC framework for named entity recognition](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 5849–5859, Online. Association for Computational Linguistics.
- Bill Yuchen Lin, Dong-Ho Lee, Ming Shen, Ryan Moreno, Xiao Huang, Prashant Shiralkar, and Xiang Ren. 2020. [Triggerer: Learning with entity triggers as explanations for named entity recognition](#). *CoRR*, abs/2004.07493.
- Fei Liu and et al. 2020. Learning to summarize from human feedback. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*.
- Yang Liu, Ming Li, Xin Liu, Yao Chen, and Maosong Sun. 2020. Presumm: A bert-based unsupervised text summarization model. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*.
- Zhengyuan Liu and Nancy F. Chen. 2021. [Controllable neural dialogue summarization with personal named entity planning](#). In *Conference on Empirical Methods in Natural Language Processing*.
- 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](#).
- Mathieu-Auguste Rondeau and et al. 2018. Reinforcement learning for bandit neural machine translation with simulated human feedback. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*.
- Itsumi Saito, Kyosuke Nishida, Kosuke Nishida, and Junji Tomita. 2020. Abstractive summarization with combination of pre-trained sequence-to-sequence and saliency models. *arXiv preprint arXiv:2003.13028*.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. [Proximal policy optimization algorithms](#). *CoRR*, abs/1707.06347.

- Eva Sharma, Chen Li, and Lu Wang. 2019. Bigpatent: A large-scale dataset for abstractive and coherent summarization. *arXiv preprint arXiv:1906.03741*.
- Mukund Srinath, Shomir Wilson, and C Lee Giles. 2021. [Privacy at scale: Introducing the PriVaSeer corpus of web 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 6829–6839, Online. Association for Computational Linguistics.
- Xiaojun Wan and et al. 2018. Improving abstractive document summarization with salient information modeling. In *Proceedings of the 27th International Conference on Computational Linguistics*.
- Chenguang Wang and et al. 2018. A hierarchical reinforced sequence operation method for abstractive summarization. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*.
- Liyuan Wang, Yue Zhang, Fuli Xu, and Xu Chen. 2021. Fine-tuning pretrained language models: Weight initializations, data orders, and early stopping. *arXiv preprint arXiv:2103.05447*.
- Juntao Yu, Bernd Bohnet, and Massimo Poesio. 2020. [Named entity recognition as dependency parsing](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 6470–6476, Online. Association for Computational Linguistics.
- Jingqing Zhang, Yao Zhao, Mohammad Saleh, and Peter J. Liu. 2020. [Pegasus: Pre-training with extracted gap-sentences for abstractive summarization](#).
- Yusen Zhang, Yang Liu, Ziyi Yang, Yuwei Fang, Yulong Chen, Dragomir Radev, Chenguang Zhu, Michael Zeng, and Rui Zhang. 2023. [Macsum: Controllable summarization with mixed attributes](#).
- Zexuan Zhong and Danqi Chen. 2020. A frustratingly easy approach for entity and relation extraction. *arXiv preprint arXiv:2010.12812*.
- Enwei Zhu and Jinpeng Li. 2022. [Boundary smoothing for named entity recognition](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7096–7108, Dublin, Ireland. Association for Computational Linguistics.