

Like a Good Nearest Neighbor: Practical Content Moderation with Sentence Transformers

Anonymous ACL submission

Abstract

Modern text classification systems have impressive capabilities but are infeasible to deploy and use reliably due to their dependence on prompting and billion-parameter language models. SetFit (Tunstall et al., 2022) is a recent, practical approach that fine-tunes a Sentence Transformer under a contrastive learning paradigm and achieves similar results to more unwieldy systems. Text classification is important for addressing the problem of domain drift in detecting harmful content, which plagues all social media platforms. Here, we propose Like a Good Nearest Neighbor (LAGONN), an inexpensive modification to SetFit that requires no additional parameters or hyperparameters but modifies input with information about its nearest neighbor, for example, the label and text, in the training data, making novel data appear similar to an instance on which the model was optimized. LAGONN is effective at the task of detecting harmful content and generally improves SetFit’s performance. To demonstrate LAGONN’s value, we conduct a thorough study of text classification systems in the context of content moderation under four label distributions.¹

1 Introduction

Text classification is the most important tool for NLP practitioners, and there has been substantial progress in advancing the state-of-the-art, especially with the advent of large, pretrained language models (PLM) (Devlin et al., 2019). Modern research focuses on in-context learning (Brown et al., 2020), pattern exploiting training (Schick and Schütze, 2021a,b, 2022), adapter-based fine-tuning with learned label embeddings (Karimi Mahabadi et al., 2022), and parameter efficient fine-tuning (Liu et al., 2022a). These methods have achieved impressive results on the SuperGLUE (Wang et al., 2019) and RAFT (Alex et al., 2021)

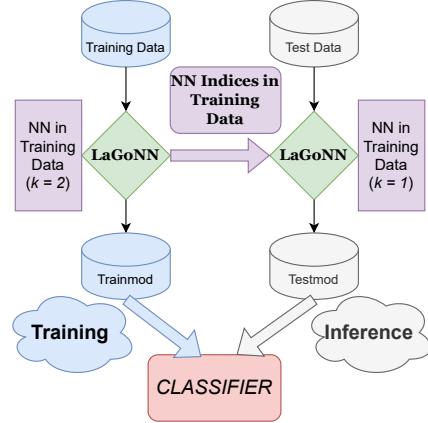


Figure 1: We embed training data, retrieve the text, gold label, and distance for each instance from its second nearest neighbor ($k=2$) and modify the original text with this information. Then we embed the modified training data and train a classifier. During inference, the NN from the training data is selected ($k=1$), the original text is modified with the text, gold label, and distance from the NN, and the classifier is called.

few-shot benchmarks, but most are difficult to use because of their reliance on billion-parameter PLMs and prompting. Constructing prompts is not trivial and may require domain expertise.

One exception to these cumbersome systems is SetFit. SetFit does not rely on prompting or billion-parameter PLMs, and instead fine-tunes a pretrained Sentence Transformer (ST) (Reimers and Gurevych, 2019) under a contrastive learning paradigm. SetFit has comparable performance to more unwieldy systems while being one to two orders of magnitude faster to train and run inference.

An important application of text classification is aiding or automating content moderation, which is the task of determining the appropriateness of user-generated content on the Internet (Roberts, 2017). From fake news to toxic comments to hate speech, it is difficult to browse social media without being exposed to potentially dangerous posts that may have an effect on our ability to reason (Ecker

¹Code and data: [https://github.com/\[REDACTED\]](https://github.com/[REDACTED])

et al., 2022). Misinformation spreads at alarming rates (Vosoughi et al., 2018), and an ML system should be able to quickly aid human moderators. While there is work in NLP with this goal (Markov et al., 2022; Shido et al., 2022; Ye et al., 2023), a general, practical and open-sourced method that is effective across multiple domains remains an open challenge. Novel fake news topics or racial slurs emerge and change constantly. Retraining of ML-based systems is required to adapt this concept drift, but this is expensive, not only in terms of computation, but also in terms of the human effort needed to collect and label data.

SetFit’s performance, speed, and low cost would make it ideal for effective content moderation, however, this type of text classification poses a challenge for even state-of-the-art approaches. For example, detecting hate speech on Twitter (Basile et al., 2019), a subtask on the RAFT few-shot benchmark, appears to be the most difficult dataset; at time of writing, it is the only task where the human baseline has not been surpassed, yet SetFit is among the top ten most performant systems.²

Here, we propose a modification to SetFit, called Like a Good Nearest Neighbor (LAGONN). LAGONN introduces no parameters or hyperparameters and instead modifies input text by retrieving information about the nearest neighbor (NN) seen during optimization (see Figure 1). Specifically, we append the label, distance, and text of the NN in the training data to a new instance and encode this modified version with an ST. By making input data appear more similar to instances seen during training, we inexpensively exploit the ST’s pretrained or fine-tuned knowledge when considering a novel example. Our method can also be applied to the linear probing of an ST, requiring no expensive fine-tuning of the large embedding model. Finally, we propose a simple alteration to the SetFit training procedure, where we fine-tune the ST on a subset of the training data. This results in a more efficient and performant text classifier that can be used with LAGONN. We summarize our contributions as follows:

1. We propose LAGONN, an inexpensive modification to SetFit- or ST-based text classification.
2. We suggest an alternative training procedure

²<https://huggingface.co/spaces/ought/raft-leaderboard> (see "Tweet Eval Hate").

to the standard fine-tuning of SetFit, that can be used with or without LAGONN, and results in a cheaper system with similar performance to the more expensive SetFit.

3. We perform an extensive study of LAGONN, SetFit, and standard transformer fine-tuning in the context of content moderation under different label distributions.

2 Related Work

There is not much work on using sentence embeddings as features for classification despite the pioneering work being roughly five years old (Perone et al., 2018). STs are pretrained with the objective of maximizing the distance between semantically distinct text and minimizing the distance between text that is semantically similar in feature space. They are composed of a Siamese and triplet architecture that encodes text into dense vectors which can be used as features for ML. STs were first used to encode text for classification by Piao (2021), however, the authors relied on pretrained representations.

SetFit uses a contrastive learning paradigm (Koch et al., 2015) to optimize the ST embedding model. The ST is fine-tuned with a distance-based loss function, like cosine similarity, such that examples with different labels are separated in feature space. Input text is then encoded with the fine-tuned ST and a classifier, such as logistic regression, is trained. This approach creates a strong, few-shot text classification system, transforming the ST from a sentence encoder to a topic encoder.

Most related to LAGONN is work done by Xu et al. (2021), who showed that retrieving and concatenating text from training data and external sources, such as ConceptNet (Speer et al., 2017) and the Wiktionary³ definition, can be viewed as a type of external attention that does not modify the architecture of the Transformer in question answering. Liu et al. (2022b) used PLMs, including STs, and k -NN lookup to prepend examples that are similar to a GPT-3 query sample to aid in prompt engineering for in-context learning. Wang et al. (2022) demonstrated that prepending and appending training data can benefit PLMs in the tasks of summarization, language modelling, machine translation, and question answering, using BM25 as their retrieval model for speed (Manning et al., 2008; Robertson and Zaragoza, 2009).

³<https://www.wiktionary.org/>

Training Data		Test Data	
	"I love this." [positive 0.0] (0)		"So good!" [?] (?)
	"This is great!" [positive 0.5] (0)		"Just terrible!" [?] (?)
	"I hate this." [negative 0.7] (1)		"Never again." [?] (?)
	"This is awful!" [negative 1.2] (1)		"This rocks!" [?] (?)
LAGoNN Configuration		Train Modified	
LABEL		"I love this. [SEP] [positive 0.5]" (0)	
TEXT		"I love this. [SEP] [positive 0.5] This is great!" (0)	
BOTH		"I love this. [SEP] [positive 0.5] This is great! [SEP] [negative 0.7] I hate this." (0)	
		Test Modified	
LABEL		"So good! [SEP] [positive 1.5]" (?)	
TEXT		"So good! [SEP] [positive 1.5] I love this." (?)	
BOTH		"So good! [SEP] [positive 1.5] I love this. [SEP] [negative 2.7] This is awful!" (?)	

Table 1: Toy training and test data and different LAGoNN configurations considering the first training example. Train and Test Modified are altered instances that are input into the final embedding model for training and inference, respectively. The input format is "original text [SEP] [NN gold label distance] NN instance text". Input text is in quotation marks, the NN’s gold label and distance from the training data are in square brackets, and the integer label is in parenthesis (see Appendix A.4 for examples of LAGoNN BOTH modified text).

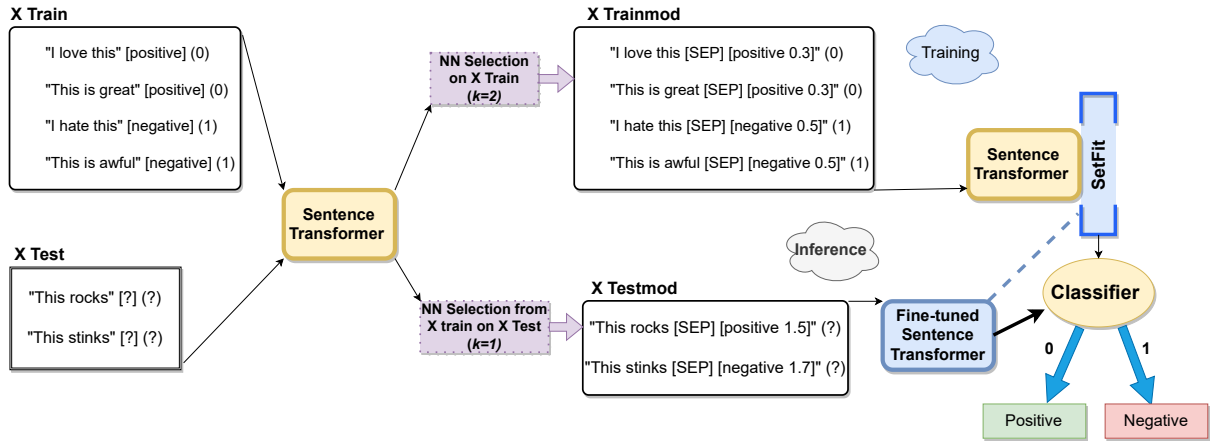


Figure 2: LAGoNN LABEL uses an ST to encode training data, performs NN lookup, appends the second NN’s ($k=2$) gold label and distance, and optionally SetFit to fine-tune the embedding model. We then embed this new instance and train a classifier. During inference, we use the embedding model to modify the test data with its NN’s gold label and distance from the training data ($k=1$), compute the final representation, and call the classifier. Input text is in quotation marks, the NN’s gold label and distance are in brackets, and the integer label is in parenthesis.

We alter the SetFit training procedure by using fewer examples to adapt the embedding model for many-shot learning. LAGoNN decorates input text with its nearest neighbor’s gold label, Euclidean distance, and text from the training data to exploit the ST’s optimized representations. Compared to retrieval-based methods, LAGoNN uses the same model for both retrieval and encoding, which can be fine-tuned via SetFit. We only retrieve information from the training data for text classification.

3 Like a Good Nearest Neighbor

Xu et al. (2021) formulate a type of external attention, where textual information is retrieved from multiple sources and added to text input to give the model stronger reasoning ability without altering the internal architecture. Inspired by this approach, LAGoNN exploits pretrained and fine-tuned knowledge through external attention, but the information we retrieve comes only from data used during optimization. We consider an embedding function, f , that is called on both training and test

data, $f(X_{train})$ and $f(X_{test})$. Considering its success and speed on realistic, few-shot data and our goal of practical content moderation, we choose an ST that can be fine-tuned with SetFit as our embedding function.

Encoding training data and nearest neighbors LAGONN first uses a pretrained Sentence Transformer to embed training text in feature space, $f(X_{train})$. We perform NN lookup with scikit-learn (Buitinck et al., 2013) on the resulting embeddings and query the second closest NN ($k=2$). We do not use the NN because it is the example itself.

Nearest neighbor information We extract text from the second nearest neighbor and use it to decorate the original example. We experimented with different text that LAGONN could use. The first configuration we consider is the gold label and Euclidean distance of the NN, which we call LABEL. We then considered the gold label, distance, and the text of the NN, which we refer to as TEXT. Finally, we tried the same format as TEXT but for all possible labels, which we call BOTH (see Table 1 and Figure 2).⁴ Information from the second NN is appended to the text following a separator token to indicate this instance is composed of multiple sequences. While the BOTH and TEXT configurations are arguably the most interesting, we find LABEL to result in the most performant version of LAGONN, and this is the version about which we report results.

Training LAGONN encodes the modified training data and optionally fine-tunes the embedding model via SetFit, $f(X_{trainmod})$. After fine-tuning, we train a classifier $CLF(f(X_{trainmod}))$, like logistic regression.

Inference LAGONN uses information from the nearest neighbor in the training data to modify input text. We compute the embeddings on the test data, $f(X_{test})$, and query the NN lookup, selecting the NN ($k=1$) in the training data and extracting information from the training text. LAGONN then decorates the input instance with information from the NN in the training data. Finally, we encode the modified data with the embedding model and call the classifier, $CLF(f(X_{testmod}))$.

Intuition As f is the same function, we hypothesize that LAGONN’s modifications will make

⁴LAGONN requires a mapping from the label to the text the label represents, for example, 0 – positive and 1 – negative.

a novel instance more semantically similar to its NNs in the training data. The resulting representation should be more akin to an instance on which the embedding model and classifier were optimized. Our method also leverages both distance-based (NN lookup) and probabilistic algorithms (logistic regression) for its final prediction.

4 Experiments

4.1 Data and label distributions

In our experiments, we study LAGONN’s performance on four binary and one ternary classification dataset related to the task of content moderation. Each dataset is composed of a training, validation, and test split.

Here, we provide a summary of the five datasets we studied. LIAR was created from Politifact⁵ for fake news detection and is composed of the data fields *context*, *speaker*, and *statement*, which are labeled with varying levels of truthfulness (Wang, 2017). We used a collapsed version of this dataset where a statement can only be true or false. We did not use *speaker*, but did use *context* and *statement*, separated by a separator token. Quora Insincere Questions⁶ is composed of neutral and toxic questions, where the author is not asking in good faith. Hate Speech Offensive⁷ has three labels and is composed of tweets that can contain either neutral text, offensive language, or hate speech (Davidson et al., 2017). Amazon Counterfactual⁸ contains sentences from product reviews, and the labels can be "factual" or "counterfactual" (O’Neill et al., 2021). "Counterfactual" indicates that the customer said something that cannot be true. Finally, Toxic Conversations⁹ is a dataset of comments where the author wrote a comment with unintended bias¹⁰ (see Table 2).

We study our system by simulating growing training data over ten discrete steps sampled under four different label distributions: extreme, imbalanced, moderate, and balanced (see Table 3). On

⁵<https://www.politifact.com/>

⁶<https://www.kaggle.com/c/quora-insincere-questions-classification>

⁷https://huggingface.co/datasets/hate_speech_offensive

⁸https://huggingface.co/datasets/SetFit/amazon_counterfactual_en

⁹https://huggingface.co/datasets/SetFit/toxic_conversations

¹⁰<https://www.kaggle.com/c/jigsaw-unintended-bias-in-toxicity-classification/overview>

Dataset (and Detection Task)	Number of Labels
LIAR (Fake News)	2
Insincere Questions (Toxicity)	2
Hate Speech Offensive	3
Amazon Counterfactual (English)	2
Toxic Conversations	2

Table 2: Summary of datasets and number of labels. We provide the type of task in parenthesis in unclear cases.

each step we add 100 examples (100 on the first, 200 on the second, etc.) from the training split sampled under one of the four ratios.¹¹ On each step, we train our method with the sampled data and evaluate on the test split. Considering growing training data has two benefits: 1) We can simulate a streaming data scenario, where new data is labeled and added for training and 2) We can investigate each method’s sensitivity to the number of training examples. We sampled over five seeds, reporting the mean and standard deviation.

Regime	Binary	Ternary
Extreme	0: 98% 1: 2%	0: 95%, 1: 2%, 2: 3%
Imbalanced	0: 90% 1: 10%	0: 80%, 1: 5%, 2: 15%
Moderate	0: 75% 1: 25%	0: 65%, 1: 10%, 2: 25%
Balanced	0: 50% 1: 50%	0: 33%, 1: 33%, 2: 33%

Table 3: Label distributions for sampling training data. 0 represents neutral while 1 and 2 represent different types of undesirable text.

4.2 Baselines

We compare LAGONN against a number of strong baselines, detailed below. We used default hyperparameters in all cases unless stated otherwise.

RoBERTa RoBERTa-base is a pretrained language model (Liu et al., 2019) that we fine-tuned with the transformers library (Wolf et al., 2020). We select two versions of RoBERTa-base: an expensive version, where we perform standard fine-tuning on each step (RoBERTa_{full}) and a cheaper version, where we freeze the model body after step one and update the classification head on subsequent steps (RoBERTa_{freeze}). We set the learning rate to $1e^{-5}$, train for a maximum of 70 epochs, and use early stopping, selecting the best model after training. We consider RoBERTa_{full} an upper bound as it has the most trainable parameters and requires the most time to train of all our methods.

¹¹For Hate Speech Offensive, 0 and 2 denote undesirable text and 1 denotes neither.

Linear probe We perform linear probing of a pretrained Sentence Transformer by fitting logistic regression with default hyperparameters on the training embeddings on each step. We choose this baseline because LAGONN can be applied as a modification in this scenario. We select MPNET (Song et al., 2020) as the ST, for SetFit, and for LAGONN.¹² We refer to this method as Probe.

Logistic regression Here, we perform standard fine-tuning with SetFit on the first step, and then on subsequent steps, freeze the embedding model and retrain only the classification head. We choose this baseline as LAGONN also uses logistic regression as its final classifier and refer to this method as Log Reg.

k -nearest neighbors Similar to the above baseline, we fine-tune the embedding model via SetFit, but swap out the classification head for a k NN classifier, where $k = 3$. We select this baseline as LAGONN also relies on an NN lookup. $k = 3$ was chosen during our development stage as it yielded the strongest performance. We refer to this method as k NN.

SetFit For this baseline we perform standard fine-tuning with SetFit on each step. On the first step, this method is equivalent to Log Reg.

LAGONN cheap This method modifies data via LAGONN before fitting a logistic regression classifier. Even without adapting the embedding model, as the training data grow, modifications made to the test data may change. We refit the classification head on each step and refer to this method as LAGONN_{cheap}, which is comparable to Probe.

LAGONN On the first step, we use LAGONN to modify our training data and then perform standard fine-tuning with SetFit. On subsequent steps, we freeze the embedding model and use it to modify our data. We fit logistic regression on each step and refer to this method as LAGONN. It is comparable to Log Reg.

LAGONN expensive This version is identical to LAGONN, except we fine-tune the embedding model on each step. We refer to this method as LAGONN_{exp} and it is comparable to SetFit. On the first step, this method is equivalent to LAGONN.

¹²<https://huggingface.co/sentence-transformers/paraphrase-mpnet-base-v2>

Method	InsincereQs				AmazonCF			
	1 st	5 th	10 th	Average	1 st	5 th	10 th	Average
RoBERTa _{full}	19.9 _{8.4}	30.9 _{7.9}	42.0 _{7.4}	33.5 _{6.7}	21.8 _{6.6}	63.9 _{10.2}	72.3 _{3.0}	59.6 _{16.8}
SetFit	24.1 _{6.3}	29.2 _{6.7}	36.7 _{7.3}	31.7 _{3.4}	22.3 _{8.8}	64.2 _{3.3}	68.6 _{4.6}	56.8 _{14.9}
LAGONN _{exp}	30.7 _{8.9}	37.6 _{6.1}	39.0 _{6.1}	36.1 _{2.3}	26.1 _{17.5}	68.4 _{4.4}	74.9 _{2.9}	63.2 _{16.7}
RoBERTa _{freeze}	19.9 _{8.4}	34.1 _{5.4}	37.9 _{5.9}	32.5 _{5.5}	21.8 _{6.6}	41.0 _{12.7}	51.3 _{10.7}	40.6 _{8.9}
kNN	6.8 _{0.42}	15.9 _{3.4}	16.9 _{4.3}	14.4 _{3.0}	10.3 _{0.2}	15.3 _{4.2}	18.4 _{3.7}	15.6 _{2.4}
Log Reg	24.1 _{6.3}	31.7 _{4.9}	36.1 _{5.4}	31.8 _{3.6}	22.3 _{8.8}	32.4 _{11.5}	42.3 _{8.8}	34.5 _{5.9}
LAGONN	30.7 _{8.9}	39.3 _{4.9}	41.2 _{4.7}	38.4 _{3.0}	26.1 _{17.5}	31.1 _{19.4}	33.0 _{19.1}	30.9 _{2.3}
Probe	24.3 _{8.4}	39.8 _{5.6}	44.8 _{4.2}	38.3 _{6.2}	24.2 _{9.0}	46.3 _{4.4}	54.6 _{2.0}	45.1 _{10.3}
LAGONN _{cheap}	23.6 _{7.8}	40.7 _{5.9}	45.3 _{4.4}	38.6 _{6.6}	20.1 _{6.9}	38.3 _{4.9}	47.8 _{3.4}	38.2 _{9.5}
<i>Balanced</i>								
RoBERTa _{full}	47.1 _{4.2}	52.1 _{3.6}	55.7 _{2.6}	52.5 _{2.9}	73.6 _{2.1}	78.6 _{3.9}	82.4 _{1.1}	78.9 _{2.2}
SetFit	43.5 _{4.2}	47.1 _{4.6}	48.5 _{3.9}	48.0 _{1.7}	73.8 _{4.4}	69.8 _{4.0}	64.1 _{4.6}	69.6 _{3.6}
LAGONN _{exp}	42.8 _{5.3}	47.6 _{2.9}	47.0 _{1.7}	46.2 _{2.0}	76.0 _{3.0}	73.4 _{2.6}	72.3 _{2.9}	72.5 _{3.4}
RoBERTa _{freeze}	47.1 _{4.2}	52.1 _{0.4}	53.3 _{1.7}	51.5 _{2.1}	73.6 _{2.1}	76.8 _{1.6}	77.9 _{1.0}	76.5 _{1.3}
kNN	22.3 _{2.3}	30.2 _{2.3}	30.9 _{1.8}	29.5 _{2.5}	41.7 _{3.4}	57.9 _{3.3}	58.3 _{3.3}	56.8 _{5.1}
Log Reg	43.5 _{4.2}	53.8 _{2.2}	55.5 _{1.6}	52.8 _{3.5}	73.8 _{4.4}	79.2 _{1.9}	80.1 _{1.0}	78.6 _{1.8}
LAGONN	42.8 _{5.3}	54.1 _{2.9}	56.3 _{1.3}	53.4 _{3.7}	76.0 _{3.0}	80.1 _{2.0}	81.4 _{1.1}	79.8 _{1.4}
Probe	47.5 _{1.6}	52.4 _{1.7}	55.3 _{1.1}	52.2 _{2.5}	52.4 _{3.4}	64.7 _{2.5}	67.5 _{0.4}	63.4 _{4.4}
LAGONN _{cheap}	49.3 _{2.6}	54.4 _{1.4}	57.6 _{0.7}	54.2 _{2.7}	48.1 _{3.4}	62.0 _{2.0}	65.3 _{0.8}	60.5 _{5.0}

Table 4: Average performance (average precision $\times 100$) on Insincere Questions and Amazon Counterfactual. The first, fifth, and tenth step are followed by the average over all ten steps. The average gives insight into the overall strongest performer by aggregating all steps. We group methods with a comparable number of trainable parameters together. The extreme label distribution results are followed by balanced (see Appendix A.2 for additional results).

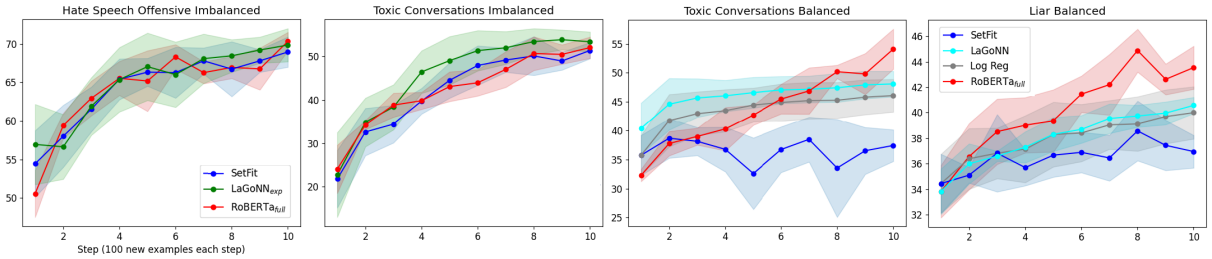


Figure 3: Average performance in the imbalanced and balanced regimes relative to comparable methods. We include RoBERTa_{full} results for reference. The metric is macro-F1 for Hate Speech Offensive, average precision elsewhere.

5 Results

Table 4 and Figure 3 show our results. In the cases of the extreme and imbalanced regimes, SetFit’s performance steadily increases with the number of training examples. As the label distribution shifts to the balanced regime, however, SetFit’s performance quickly saturates or even degrades as the number of training examples grows. LAGONN, RoBERTa_{full}, and Log Reg, other fine-tuned PLM classifiers, do not exhibit this behavior. LAGONN_{exp}, being based on SetFit, exhibits a similar trend, but the performance degradation is mitigated; on the 10th step of Amazon Counterfac-

tual in Table 4 SetFit’s performance decreased by 9.7, while LAGONN_{exp} only fell by 3.7.

LAGONN and LAGONN_{exp} generally outperform Log Reg and SetFit, respectively, often resulting in a more stable model, as reflected in the standard deviation. We find that LAGONN and LAGONN_{exp} exhibit stronger predictive power with fewer examples than RoBERTa_{full} despite having fewer trainable parameters. For example, on the first step of Insincere Questions under the extreme setting, LAGONN’s performance is more than 10 points higher.

LAGONN_{cheap} outperforms all other methods

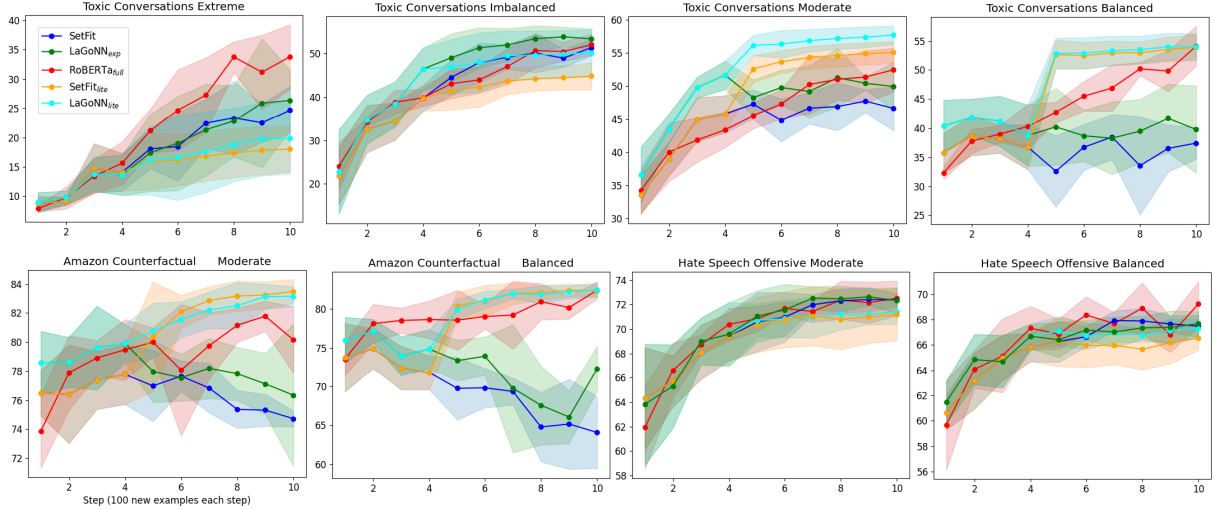


Figure 4: Average performance for all sampling regimes on Toxic Conversations and the moderate and balanced regimes for Amazon Counterfactual and Hate Speech Offensive. More expensive models, such as LAGONN_{exp}, SetFit, and RoBERTa_{full} perform best when the label distribution is imbalanced. As the distribution becomes more balanced, inexpensive models, such as LAGONN_{lite}, show similar or improved performance. The metric is macro-F1 for Hate Speech Offensive, average precision elsewhere (see Appendix A.3 for additional results).

on the Insincere Questions dataset for all balance regimes, despite being the third fastest (see Table 5) and having the second fewest trainable parameters. We attribute this result to the fact that this dataset is composed of questions from Quora¹³ and our ST backbone was pretrained on similar data. This intuition is supported by Probe, the cheapest method, which despite having the fewest trainable parameters, shows comparable performance.

5.1 SetFit for efficient many-shot learning

Respectively comparing SetFit to Log Reg and LAGONN_{exp} to LAGONN suggests that fine-tuning the ST embedding model on moderate or balanced data hurts model performance as the number of training samples grows. We therefore hypothesize that randomly sampling a subset of training data to fine-tune the encoder, freezing, embedding the remaining data, and training the classifier will result in a stronger model.

To test our hypothesis, we add two models to our experimental setup: SetFit_{lite} and LAGONN_{lite}. SetFit_{lite} and LAGONN_{lite} are respectively equivalent to SetFit and LAGONN_{exp}, except after the fourth step (400 samples), we freeze the encoder and only retrain the classifier on subsequent steps, similar to Log Reg and LAGONN.

Figure 4 shows our results with these two new models. As expected, in the cases of extreme and imbalanced distributions, LAGONN_{exp}, SetFit,

and RoBERTa_{exp}, are the strongest performers on Toxic Conversations. We note very different results for both LAGONN_{lite} and SetFit_{lite} compared to LAGONN_{exp} and SetFit on Toxic Conversations and Amazon Counterfactual under the moderate and balanced label distributions. As their expensive counterparts start to plateau or degrade on the fourth step, the predictive power of these two new models dramatically increases, showing improved or comparable performance to RoBERTa_{full}, despite being optimized on less data; for example, LAGONN_{lite} reaches an average precision of approximately 55 after being optimized on only 500 examples. RoBERTa_{full} does not exhibit similar performance until the tenth step. Finally, we point out that LAGONN-based methods generally provide a performance boost for SetFit-based classification.

5.2 LAGONN’s computational expense

LAGONN is more computationally expensive than Sentence Transformer- or SetFit-based text classification. LAGONN introduces additional inference with the encoder, NN-lookup, and string modification. As the computational complexity of transformers increases with sequence length (Vaswani et al., 2017), additional expense is created when LAGONN appends textual information before inference with the ST. In Table 5, we provide a speed comparison between Probe, Log Reg, SetFit, and LAGONN classification computed on the same

¹³<https://www.quora.com/>

Method	Time in seconds
Probe	22.9
LAGONN _{cheap}	44.2
Log Reg	42.9
LAGONN	63.4
SetFit	207.3
LAGONN _{exp}	238.0
RoBERTa _{full}	446.9

Table 5: Speed comparison between LAGONN and comparable methods. Time includes training each method on 1,000 examples and performing inference on 51,000 examples.

hardware.¹⁴ On average, LAGONN introduced 24.2 additional seconds of computation compared to its relative counterpart.

6 Discussion

Modern research has achieved impressive results on a variety of text classification tasks and with limited training data. SetFit is one such example and can be used practically, but based on our results, the task of text classification for content moderation presents a challenge even for state-of-the-art approaches. It is imperative that we develop reliable methods that can be feasibly and quickly applied. These methods should be as inexpensive as possible such that we can re-tune them for novel forms of hate speech, toxicity, and fake news.

Our results suggest that LAGONN_{exp} or SetFit, relatively expensive techniques, can detect harmful content when dealing with imbalanced label distributions, as is common with realistic datasets. This finding is intuitive from the perspective that less common instances are more difficult to learn and require more effort. The exception to this would be our examination of Insincere Questions, where LAGONN_{cheap} excelled. This highlights the fact that we can inexpensively extract pretrained knowledge if PLMs are chosen with care for related tasks.

Standard fine-tuning with SetFit does not help performance on more balanced datasets that are not few-shot. SetFit was developed for few-shot learning, but we have observed that it should not be applied "out of the box" to balanced, non-few-shot data. This can be detrimental to performance and has a direct effect on our approach. However, we have observed that LAGONN can stabilize Set-

Fit's predictions and reduce its performance drop. Figures 3 and 4 show that when the label distribution is moderate or balanced (see Table 3), SetFit plateaus, yet less expensive systems, such as LAGONN, continue to learn. We believe this is due to SetFit's fine-tuning objective, which optimizes a Sentence Transformer using cosine similarity loss to separate examples belonging to different labels in feature space by assuming independence between labels. This may be too strong an assumption as we optimize with more examples, which is counter-intuitive for data-hungry transformers. RoBERTa_{full}, optimized with cross-entropy loss, generally showed improved performance as we added training data.

When dealing with balanced data, it is sufficient to fine-tune the Sentence Transformer via SetFit with 50 to 100 examples per label, while 150 to 200 instances appear to be sufficient when the training data are moderately balanced. The encoder can then be frozen and all available data embedded to train a classifier. This improves performance and is more efficient than full-model fine-tuning. LAGONN is directly applicable to this case, boosting the performance of SetFit_{lite} without introducing trainable parameters. In this setup, all models fine-tuned on Hate Speech Offensive exhibited similar, upward-trending learning curves, but we note the speed of LAGONN relative to RoBERTa_{full} or SetFit (see Figure 4 and Table 5).

7 Conclusion

We have proposed LAGONN, a simple and inexpensive modification to Sentence Transformer- or SetFit-based text classification. LAGONN does not introduce any trainable parameters or new hyperparameters, but typically improves SetFit's performance. To demonstrate the merit of LAGONN, we examined text classification systems in the context of content moderation under four label distributions on five datasets and with growing training data. To our knowledge, this is the first work to examine SetFit in this way. When the training labels are imbalanced, expensive systems, such as LAGONN_{exp} are performant. However, when the distribution is balanced, standard fine-tuning with SetFit can actually hurt model performance. We have therefore proposed an alternative fine-tuning procedure to which LAGONN can be easily utilized, resulting in a powerful, but inexpensive system capable of detecting harmful content.

¹⁴We used a 40 GB NVIDIA A100 Tensor Core GPU.

8 Limitations

In the current work, we have only considered text data, but social media content can of course consist of text, images, and videos. As LAGONN depends only on an embedding model, an obvious extension to our approach would be examining the modifications we suggest, but on multimodal data. This is an interesting direction that we leave for future research. We have also considered English data, but harmful content can appear in any language. The authors demonstrated that SetFit is performant on multilingual data, the only necessary modification being the underlying pretrained ST. We therefore suspect that LAGONN would behave similarly on non-English data, but this is not something we have tested ourselves. In order to examine our system’s performance under different label-balance distributions, we restricted ourselves to binary and ternary text classification tasks, and LAGONN therefore remains untested when there are more than three labels. We did not study our method when there are fewer than 100 examples, and investigating LAGONN in a few-shot learning setting is fascinating topic for future study. Finally, we note that our system could be misused to detect undesirable content that is not necessarily harmful. For example, a social media website could detect and silence users who complain about the platform. This is not our intended use case, but could result from any classifier, and potential misuse is an unfortunate drawback of all technology.

9 Ethics Statement

It is our sincere goal that our work contributes to the social good in multiple ways. We first hope to have furthered research on text classification that can be feasibly applied to combat undesirable content, such as misinformation, on the Internet, which could potentially cause someone harm. To this end, we have tried to describe our approach as accurately as possible and released our code and data, such that our work is transparent and can be easily reproduced and expanded upon. We hope that we have also created a useful but efficient system which reduces the need to expend energy in the form expensive computation. For example, LAGONN does not rely on billion-parameter language models that demand thousand-dollar GPUs to use. LAGONN makes use of GPUs no more than SetFit, despite being more computationally expensive. We have additionally proposed a simple method to make

SetFit, an already relatively inexpensive method, even more efficient.

References

- Neel Alex, Eli Lifland, Lewis Tunstall, Abhishek Thakur, Pegah Maham, C. Jess Riedel, Emmie Hine, Carolyn Ashurst, Paul Sedille, Alexis Carlier, Michael Noetel, and Andreas Stuhlmüller. 2021. [RAFT: A real-world few-shot text classification benchmark](#). In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*.
- Valerio Basile, Cristina Bosco, Elisabetta Fersini, Debora Nozza, Viviana Patti, Francisco Manuel Rangel Pardo, Paolo Rosso, and Manuela Sanguinetti. 2019. [SemEval-2019 task 5: Multilingual detection of hate speech against immigrants and women in Twitter](#). In *Proceedings of the 13th International Workshop on Semantic Evaluation*, pages 54–63, Minneapolis, Minnesota, USA. Association for Computational Linguistics.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Matheus Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901. Curran Associates, Inc.
- Lars Buitinck, Gilles Louppe, Mathieu Blondel, Fabian Pedregosa, Andreas Mueller, Olivier Grisel, Vlad Niculae, Peter Prettenhofer, Alexandre Gramfort, Jacques Grobler, Robert Layton, Jake VanderPlas, Arnaud Joly, Brian Holt, and Gaël Varoquaux. 2013. [API design for machine learning software: experiences from the scikit-learn project](#). In *ECML PKDD Workshop: Languages for Data Mining and Machine Learning*, pages 108–122.
- Thomas Davidson, Dana Warmusley, Michael Macy, and Ingmar Weber. 2017. [Automated hate speech detection and the problem of offensive language](#). In *Proceedings of the 11th International AAAI Conference on Web and Social Media, ICWSM ’17*, pages 512–515.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages

616	4171–4186, Minneapolis, Minnesota. Association for	Christian S. Perone, Roberto Pereira Silveira, and	671
617	Computational Linguistics.	Thomas S. Paula. 2018. Evaluation of sentence em-	672
		beddings in downstream and linguistic probing tasks.	673
618	Ullrich K. H. Ecker, Stephan Lewandowsky, John	arXiv preprint arXiv:1806.06259.	674
619	Cook, Philipp Schmid, Lisa K. Fazio, Nadia Brashier,		
620	Panayiota Kendeou, Emily K. Vraga, and Michelle A.	Guangyuan Piao. 2021. Scholarly text classification	675
621	Amazeen. 2022. The psychological drivers of mis-	with sentence bert and entity embeddings. In <i>Trends</i>	676
622	information belief and its resistance to correction.	<i>and Applications in Knowledge Discovery and Data</i>	677
623	<i>Nature Reviews Psychology</i> , 1(1):13–29.	<i>Mining</i> , pages 79–87, Cham. Springer International	678
		Publishing.	679
624	Rabeeh Karimi Mahabadi, Luke Zettlemoyer, James	Nils Reimers and Iryna Gurevych. 2019. Sentence-	680
625	Henderson, Lambert Mathias, Marzieh Saeidi,	BERT: Sentence embeddings using Siamese BERT-	681
626	Veselin Stoyanov, and Majid Yazdani. 2022. Prompt-	networks. In <i>Proceedings of the 2019 Conference on</i>	682
627	free and efficient few-shot learning with language	<i>Empirical Methods in Natural Language Processing</i>	683
628	models. In <i>Proceedings of the 60th Annual Meet-</i>	<i>and the 9th International Joint Conference on Natu-</i>	684
629	<i>ing of the Association for Computational Linguistics</i>	<i>ral Language Processing (EMNLP-IJCNLP)</i> , pages	685
630	<i>(Volume 1: Long Papers)</i> , pages 3638–3652, Dublin,	3982–3992, Hong Kong, China. Association for Com-	686
631	Ireland. Association for Computational Linguistics.	putational Linguistics.	687
632	Gregory Koch, Richard Zemel, Ruslan Salakhutdinov,	Sarah T. Roberts. 2017. Content Moderation , pages 1–4.	688
633	et al. 2015. Siamese neural networks for one-shot im-	Springer International Publishing, Cham.	689
634	age recognition. In <i>ICML Deep Learning Workshop</i> ,		
635	volume 2, page 0. Lille.	Stephen Robertson and Hugo Zaragoza. 2009. The	690
		probabilistic relevance framework: Bm25 and be-	691
636	Haokun Liu, Derek Tam, Mohammed Muqeeth, Jay Mo-	yond. <i>Found. Trends Inf. Retr.</i> , 3(4):333–389.	692
637	hta, Tenghao Huang, Mohit Bansal, and Colin Raffel.		
638	2022a. Few-shot parameter-efficient fine-tuning is	Timo Schick and Hinrich Schütze. 2021a. Exploiting	693
639	better and cheaper than in-context learning. <i>arXiv</i>	cloze-questions for few-shot text classification and	694
640	<i>preprint arXiv:2205.05638.</i>	natural language inference. In <i>Proceedings of the</i>	695
		<i>16th Conference of the European Chapter of the Asso-</i>	696
641	Jiachang Liu, Dinghan Shen, Yizhe Zhang, Bill Dolan,	<i>ciation for Computational Linguistics: Main Volume</i> ,	697
642	Lawrence Carin, and Weizhu Chen. 2022b. What	pages 255–269, Online. Association for Computa-	698
643	makes good in-context examples for GPT-3? In	tional Linguistics.	699
644	<i>Proceedings of Deep Learning Inside Out (DeeLIO</i>		
645	<i>2022): The 3rd Workshop on Knowledge Extrac-</i>	Timo Schick and Hinrich Schütze. 2021b. It’s not just	700
646	<i>tion and Integration for Deep Learning Architectures</i> ,	size that matters: Small language models are also few-	701
647	pages 100–114, Dublin, Ireland and Online. Associa-	shot learners. In <i>Proceedings of the 2021 Conference</i>	702
648	tion for Computational Linguistics.	<i>of the North American Chapter of the Association</i>	703
		<i>for Computational Linguistics: Human Language</i>	704
649	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	<i>Technologies</i> , pages 2339–2352, Online. Association	705
650	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	for Computational Linguistics.	706
651	Luke Zettlemoyer, and Veselin Stoyanov. 2019.		
652	Roberta: A robustly optimized bert pretraining ap-	Timo Schick and Hinrich Schütze. 2022. True few-shot	707
653	proach. <i>arXiv preprint arXiv:1907.11692.</i>	learning with Prompts—A real-world perspective.	708
		<i>Transactions of the Association for Computational</i>	709
654	Christopher D. Manning, Prabhakar Raghavan, and Hin-	<i>Linguistics</i> , 10:716–731.	710
655	rich Schütze. 2008. <i>Introduction to Information Re-</i>		
656	<i>trieval.</i> Cambridge University Press, USA.	Yusuke Shido, Hsien-Chi Liu, and Keisuke Umezawa.	711
		2022. Textual content moderation in C2C market-	712
657	Todor Markov, Chong Zhang, Sandhini Agarwal, Tyna	place. In <i>Proceedings of the Fifth Workshop on</i>	713
658	Eloundou, Teddy Lee, Steven Adler, Angela Jiang,	<i>e-Commerce and NLP (ECNLP 5)</i> , pages 58–62,	714
659	and Lilian Weng. 2022. A holistic approach	Dublin, Ireland. Association for Computational Lin-	715
660	to undesired content detection. <i>arXiv preprint</i>	guistics.	716
661	<i>arXiv:2208.03274.</i>		
		Kaitao Song, Xu Tan, Tao Qin, Jianfeng Lu, and Tie-	717
662	James O’Neill, Polina Rozenshtein, Ryuichi Kiryo, Mo-	Yan Liu. 2020. Mpnnet: Masked and permuted pre-	718
663	toko Kubota, and Danushka Bollegala. 2021. I wish	training for language understanding. In <i>Advances in</i>	719
664	I would have loved this one, but I didn’t – a multilin-	<i>Neural Information Processing Systems</i> , volume 33,	720
665	gual dataset for counterfactual detection in product	pages 16857–16867. Curran Associates, Inc.	721
666	review. In <i>Proceedings of the 2021 Conference on</i>		
667	<i>Empirical Methods in Natural Language Processing</i> ,	Robyn Speer, Joshua Chin, and Catherine Havasi. 2017.	722
668	pages 7092–7108, Online and Punta Cana, Domini-	Conceptnet 5.5: An open multilingual graph of gen-	723
669	can Republic. Association for Computational Lin-	eral knowledge. <i>Proceedings of the AAAI Conference</i>	724
670	guistics.	<i>on Artificial Intelligence</i> , 31(1).	725

Lewis Tunstall, Nils Reimers, Unso Eun Seo Jo, Luke Bates, Daniel Korat, Moshe Wasserblat, and Oren Pereg. 2022. [Efficient few-shot learning without prompts](#). *arXiv preprint arXiv:2209.11055*.

Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc.

Soroush Vosoughi, Deb Roy, and Sinan Aral. 2018. [The spread of true and false news online](#). *Science*, 359(6380):1146–1151.

Alex Wang, Yada Pruksachatkun, Nikita Nangia, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel Bowman. 2019. [Superglue: A stickier benchmark for general-purpose language understanding systems](#). In *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc.

Shuohang Wang, Yichong Xu, Yuwei Fang, Yang Liu, Siqi Sun, Ruochen Xu, Chenguang Zhu, and Michael Zeng. 2022. [Training data is more valuable than you think: A simple and effective method by retrieving from training data](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3170–3179, Dublin, Ireland. Association for Computational Linguistics.

William Yang Wang. 2017. [“Liar, liar pants on fire”: A new benchmark dataset for fake news detection](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 422–426, Vancouver, Canada. Association for Computational Linguistics.

Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander Rush. 2020. [Transformers: State-of-the-art natural language processing](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 38–45, Online. Association for Computational Linguistics.

Yichong Xu, Chenguang Zhu, Shuohang Wang, Siqi Sun, Hao Cheng, Xiaodong Liu, Jianfeng Gao, Pengcheng He, Michael Zeng, and Xuedong Huang. 2021. [Human parity on commonsenseqa: Augmenting self-attention with external attention](#). *arXiv preprint arXiv:2112.03254*, abs/2112.03254.

Meng Ye, Karan Sikka, Katherine Atwell, Sabit Hassan, Ajay Divakaran, and Malihe Alikhani. 2023. [Multi-lingual content moderation: A case study on reddit](#). *arXiv preprint arXiv:2302.09618*.

A Appendix

A.1 Observations about LAGONN

Our original goal was to construct a system that did not need to be updated after step one and could simply perform inference on subsequent steps, an active learning setup. While the performance of this version of LAGONN did not degrade, it also did not appear to learn anything and we found it necessary to update parameters on each step. We additionally tried fine-tuning the embedding model via SetFit first before modifying data, however, this hurt performance in all cases. We include this information for transparency and because we find it interesting.

A.2 Additional results for initial experiments

Here we provide additional results from our initial experimental setup that, due to space limitations, could not be included in the main text. We note that a version of LAGONN outperforms or has the same performance of all methods, including our upper bound RoBERTa_{full}, on 54% of all displayed results, and is the best performer relative to Sentence Transformer-based methods on 72%. This excludes LAGONN_{cheap}. This method showed strong performance on the Insincere Questions dataset, but hurts performance in other cases. In cases, when SetFit-based methods do outperform our system, the performances are comparable, yet they can be quite dramatic when LAGONN-based methods are the strongest. Below, we report the mean average precision $\times 100$ for all methods over five seeds with the standard deviation, except in the case of Hate Speech Offensive, where the evaluation metric is the macro-F1. Each table shows the results for a given dataset and a given label-balance distribution on the first, fifth, and tenth step followed by the average for all ten steps. The Liar dataset seems to be the most difficult for all methods. This is expected because it likely does not include enough context to determine the truth of a statement.

Method	Insincere-Questions			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	39.8 _{5.5}	53.1 _{4.6}	55.7 _{1.2}	50.6 _{4.4}
SetFit	43.7 _{2.7}	52.2 _{1.9}	53.8 _{0.9}	51.4 _{2.9}
LAGoNN _{exp}	44.5 _{4.5}	52.7 _{2.4}	55.4 _{2.0}	51.8 _{3.0}
RoBERTa _{freeze}	39.8 _{5.5}	44.1 _{3.6}	46.3 _{2.4}	44.0 _{2.0}
kNN	23.9 _{2.2}	30.3 _{3.0}	31.6 _{2.4}	30.0 _{2.1}
Log Reg	43.7 _{2.7}	47.6 _{1.6}	50.1 _{2.1}	47.6 _{1.8}
LAGoNN	44.5 _{4.5}	48.1 _{2.2}	50.3 _{1.7}	48.1 _{1.9}
Probe	40.4 _{4.2}	49.4 _{2.3}	52.3 _{1.7}	49.0 _{3.3}
LAGoNN _{cheap}	40.8 _{4.3}	51.1 _{2.4}	54.5 _{1.4}	50.4 _{4.0}

Table 6

Method	Insincere Questions			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	48.1 _{2.3}	54.7 _{1.9}	57.5 _{1.5}	53.9 _{2.9}
SetFit	48.9 _{1.7}	53.9 _{0.7}	54.2 _{1.5}	52.3 _{1.6}
LAGoNN _{exp}	49.8 _{1.6}	52.2 _{1.9}	53.2 _{3.3}	52.0 _{1.4}
RoBERTa _{freeze}	48.1 _{2.3}	50.2 _{2.2}	52.0 _{1.4}	50.2 _{1.4}
kNN	28.0 _{2.4}	33.9 _{2.8}	33.6 _{2.0}	33.5 _{1.9}
Log Reg	48.9 _{1.7}	53.6 _{1.9}	55.8 _{1.7}	53.3 _{2.2}
LAGoNN	49.8 _{1.6}	54.4 _{1.3}	56.9 _{0.5}	54.2 _{2.2}
Probe	45.7 _{2.1}	52.3 _{1.8}	54.4 _{1.1}	51.4 _{2.5}
LAGoNN _{cheap}	45.7 _{2.2}	54.4 _{1.6}	56.4 _{0.6}	53.2 _{3.2}

Table 7

Method	Amazon Counterfactual			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	68.2 _{4.5}	81.0 _{1.7}	82.2 _{1.0}	79.2 _{3.9}
SetFit	72.0 _{2.1}	78.4 _{2.8}	78.8 _{1.2}	78.0 _{2.1}
LAGoNN _{exp}	74.3 _{3.8}	80.1 _{1.4}	79.0 _{1.6}	79.5 _{1.9}
RoBERTa _{freeze}	68.2 _{4.5}	75.0 _{2.2}	77.0 _{2.4}	74.2 _{2.6}
kNN	51.0 _{4.1}	60.0 _{3.1}	61.3 _{2.1}	59.7 _{3.0}
Log Reg	72.0 _{2.1}	74.4 _{2.3}	76.7 _{1.8}	74.8 _{1.4}
LAGoNN	74.3 _{3.8}	76.1 _{3.6}	77.3 _{3.2}	76.1 _{1.0}
Probe	46.6 _{2.8}	60.3 _{1.4}	64.2 _{1.2}	59.2 _{5.2}
LAGoNN _{cheap}	38.2 _{3.2}	55.3 _{1.8}	61.0 _{1.2}	54.4 _{6.7}

Table 8

Method	Amazon Counterfactual			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	73.9 _{2.5}	80.0 _{1.0}	80.1 _{2.3}	79.1 _{2.1}
SetFit	76.5 _{1.6}	77.0 _{2.4}	74.7 _{0.5}	76.5 _{1.0}
LAGoNN _{exp}	78.6 _{2.2}	78.0 _{2.1}	76.3 _{4.9}	78.2 _{1.0}
RoBERTa _{freeze}	73.9 _{2.5}	76.6 _{1.4}	78.5 _{0.7}	76.4 _{1.7}
kNN	54.5 _{3.1}	64.2 _{1.9}	66.6 _{1.3}	64.7 _{3.5}
Log Reg	76.5 _{1.6}	80.6 _{0.5}	81.2 _{0.3}	80.0 _{1.4}
LAGoNN	78.6 _{2.2}	81.2 _{1.4}	81.6 _{1.1}	80.8 _{0.9}
Probe	52.3 _{2.0}	64.1 _{1.8}	67.2 _{1.4}	63.1 _{4.3}
LAGoNN _{cheap}	47.3 _{3.4}	60.7 _{1.5}	65.2 _{1.4}	59.5 _{5.2}

Table 9

Method	Toxic Conversations			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	7.9 _{0.5}	21.2 _{3.7}	33.8 _{5.5}	21.9 _{9.3}
SetFit	8.8 _{1.2}	18.1 _{3.4}	24.7 _{4.1}	17.6 _{5.5}
LAGoNN _{exp}	8.9 _{1.7}	17.4 _{6.6}	26.4 _{5.2}	17.9 _{6.0}
RoBERTa _{freeze}	7.9 _{0.5}	12.8 _{2.4}	19.1 _{3.2}	13.5 _{3.5}
kNN	7.9 _{0.0}	8.7 _{0.4}	8.7 _{0.2}	8.5 _{0.3}
Log Reg	8.8 _{1.2}	13.1 _{2.5}	16.3 _{3.0}	13.0 _{2.6}
LAGoNN	8.9 _{1.7}	13.8 _{3.9}	17.1 _{4.8}	13.4 _{2.6}
Probe	13.1 _{2.8}	24.6 _{2.6}	30.1 _{2.1}	23.9 _{5.6}
LAGoNN _{cheap}	11.3 _{2.2}	21.7 _{2.7}	27.4 _{2.3}	21.3 _{5.3}

Table 10

Method	Toxic Conversations			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	24.1 _{5.6}	43.1 _{3.4}	52.1 _{2.5}	42.4 _{8.2}
SetFit	21.8 _{6.6}	44.5 _{4.1}	51.4 _{1.9}	42.1 _{9.3}
LAGoNN _{exp}	22.7 _{9.8}	49.1 _{5.6}	53.4 _{2.3}	45.6 _{9.8}
RoBERTa _{freeze}	24.1 _{5.6}	31.2 _{4.4}	34.0 _{4.0}	30.5 _{3.1}
kNN	11.5 _{2.5}	14.7 _{4.0}	15.3 _{3.2}	14.6 _{1.1}
Log Reg	21.8 _{6.6}	26.7 _{5.3}	30.2 _{4.0}	26.6 _{2.7}
LAGoNN	22.7 _{9.8}	27.6 _{8.9}	30.3 _{3.7}	27.4 _{2.4}
Probe	23.3 _{2.7}	33.0 _{2.8}	37.1 _{1.8}	32.5 _{4.2}
LAGoNN _{cheap}	20.5 _{3.2}	31.1 _{3.2}	35.6 _{1.8}	30.5 _{4.6}

Table 11

Method	Toxic Conversations			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	34.2 _{3.4}	45.5 _{1.9}	52.4 _{3.3}	45.7 _{5.6}
SetFit	33.6 _{2.9}	47.2 _{2.2}	46.6 _{3.3}	44.3 _{4.3}
LAGoNN _{exp}	36.6 _{4.2}	48.2 _{2.7}	49.9 _{3.7}	48.0 _{4.4}
RoBERTa _{freeze}	34.2 _{3.4}	38.4 _{2.1}	39.5 _{1.8}	38.0 _{1.5}
kNN	19.4 _{1.9}	21.5 _{3.4}	22.4 _{2.9}	21.6 _{0.8}
Log Reg	33.6 _{2.9}	39.2 _{2.9}	41.6 _{2.7}	38.6 _{2.4}
LAGoNN	36.6 _{4.2}	42.7 _{3.7}	45.0 _{3.5}	42.0 _{2.5}
Probe	29.0 _{2.7}	36.1 _{1.2}	39.1 _{1.5}	35.5 _{3.3}
LAGoNN _{cheap}	26.1 _{2.7}	34.3 _{1.3}	37.5 _{1.8}	33.6 _{3.6}

Table 12

Method	Toxic Conversations			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	32.3 _{1.1}	42.7 _{1.8}	54.1 _{3.4}	43.8 _{6.3}
SetFit	35.7 _{3.4}	32.6 _{6.2}	37.4 _{2.7}	36.5 _{1.9}
LAGoNN _{exp}	40.4 _{4.4}	40.2 _{6.6}	39.8 _{7.5}	40.0 _{1.2}
RoBERTa _{freeze}	32.3 _{1.1}	39.2 _{1.5}	41.0 _{0.6}	38.5 _{2.4}
kNN	17.4 _{0.8}	23.7 _{2.6}	24.3 _{2.7}	23.1 _{2.0}
Log Reg	35.7 _{3.4}	44.5 _{2.9}	46.1 _{2.8}	43.6 _{2.9}
LAGoNN	40.4 _{4.4}	46.6 _{2.7}	48.1 _{2.2}	46.1 _{2.2}
Probe	29.5 _{2.4}	35.9 _{0.9}	40.2 _{0.9}	36.1 _{3.5}
LAGoNN _{cheap}	26.8 _{2.7}	34.5 _{1.3}	38.5 _{0.8}	34.4 _{3.7}

Table 13

Method	Hate Speech Offensive			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	30.2 _{1.4}	43.5 _{2.5}	51.2 _{2.2}	44.3 _{7.4}
SetFit	30.3 _{0.8}	44.0 _{1.3}	51.1 _{2.0}	43.8 _{6.5}
LAGONN _{exp}	30.3 _{0.7}	40.7 _{2.9}	49.1 _{4.4}	42.2 _{6.2}
RoBERTa _{freeze}	30.2 _{1.4}	33.5 _{3.1}	34.4 _{3.4}	33.1 _{1.4}
kNN	31.5 _{1.2}	35.9 _{2.7}	37.4 _{2.0}	35.8 _{1.7}
Log Reg	30.3 _{0.8}	38.4 _{2.5}	41.1 _{1.5}	37.8 _{3.3}
LAGONN	30.3 _{0.7}	35.7 _{2.6}	39.1 _{2.4}	35.6 _{2.7}
Probe	29.0 _{0.2}	34.7 _{1.5}	40.1 _{2.1}	35.1 _{3.8}
LAGONN _{cheap}	29.0 _{0.1}	36.9 _{1.8}	40.5 _{2.1}	36.2 _{3.7}

Table 14

Method	Hate Speech Offensive			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	50.6 _{3.0}	65.2 _{3.9}	70.3 _{1.2}	64.2 _{5.3}
SetFit	54.4 _{4.3}	66.3 _{1.8}	68.9 _{2.0}	64.3 _{4.5}
LAGONN _{exp}	57.0 _{5.2}	67.0 _{4.4}	69.8 _{2.1}	64.9 _{4.6}
RoBERTa _{freeze}	50.6 _{3.0}	54.1 _{1.6}	55.3 _{2.3}	54.1 _{1.3}
kNN	55.6 _{4.8}	57.3 _{2.3}	58.8 _{3.6}	57.4 _{1.1}
Log Reg	54.4 _{4.3}	57.0 _{3.9}	58.2 _{3.8}	57.2 _{1.1}
LAGONN	57.0 _{5.2}	58.2 _{4.1}	58.3 _{3.4}	58.3 _{6.6}
Probe	46.5 _{2.2}	57.8 _{1.7}	60.3 _{1.2}	56.5 _{4.5}
LAGONN _{cheap}	47.1 _{1.3}	56.5 _{2.2}	59.5 _{2.5}	55.6 _{3.8}

Table 15

Method	Hate Speech Offensive			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	61.9 _{3.4}	70.8 _{1.0}	72.5 _{1.4}	69.9 _{3.2}
SetFit	64.3 _{4.2}	70.6 _{2.4}	72.4 _{0.5}	69.8 _{2.8}
LAGONN _{exp}	63.8 _{4.9}	71.0 _{2.1}	72.3 _{1.0}	70.0 _{3.0}
RoBERTa _{freeze}	61.9 _{3.4}	63.2 _{4.1}	64.1 _{4.5}	63.2 _{0.6}
kNN	64.3 _{4.0}	63.3 _{2.9}	63.9 _{2.5}	63.7 _{0.4}
Log Reg	64.3 _{4.2}	67.3 _{3.2}	67.6 _{2.3}	66.9 _{1.1}
LAGONN	63.8 _{4.9}	65.0 _{5.3}	66.7 _{5.9}	65.3 _{0.9}
Probe	55.6 _{1.7}	63.8 _{0.8}	66.1 _{0.3}	63.2 _{3.0}
LAGONN _{cheap}	56.0 _{3.6}	62.2 _{1.4}	66.0 _{0.9}	62.3 _{2.9}

Table 16

Method	Hate Speech Offensive			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	59.7 _{3.5}	66.9 _{1.2}	69.2 _{1.8}	66.4 _{2.7}
SetFit	60.7 _{1.3}	66.3 _{1.6}	67.5 _{0.9}	65.9 _{2.2}
LAGONN _{exp}	61.5 _{1.7}	66.4 _{1.4}	67.7 _{0.9}	66.1 _{1.8}
RoBERTa _{freeze}	59.7 _{3.5}	60.4 _{2.7}	63.1 _{2.3}	61.0 _{1.3}
kNN	60.7 _{1.3}	59.6 _{2.8}	59.5 _{2.5}	59.5 _{0.5}
Log Reg	60.7 _{1.3}	62.5 _{0.7}	63.4 _{1.0}	62.3 _{1.0}
LAGONN	61.5 _{1.7}	62.8 _{1.5}	64.2 _{1.0}	63.0 _{0.9}
Probe	54.9 _{1.4}	58.5 _{0.9}	60.9 _{0.4}	58.7 _{1.7}
LAGONN _{cheap}	54.2 _{2.3}	58.6 _{0.6}	60.6 _{0.5}	58.5 _{1.8}

Table 17

Method	Liar			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	32.0 _{2.7}	34.7 _{2.9}	35.1 _{4.3}	33.7 _{1.0}
SetFit	31.2 _{3.8}	30.4 _{3.1}	31.8 _{2.9}	31.5 _{0.7}
LAGONN _{exp}	30.6 _{4.7}	30.3 _{2.0}	31.3 _{2.0}	31.1 _{0.6}
RoBERTa _{freeze}	32.0 _{2.7}	32.8 _{4.5}	34.2 _{5.0}	33.2 _{0.7}
kNN	27.0 _{0.5}	27.3 _{0.8}	27.9 _{0.8}	27.4 _{0.3}
Log Reg	31.2 _{3.8}	33.7 _{5.1}	35.7 _{5.1}	34.3 _{1.6}
LAGONN	30.6 _{4.7}	32.0 _{4.6}	33.7 _{5.4}	32.6 _{0.9}
Probe	30.7 _{2.0}	30.6 _{3.9}	31.7 _{2.9}	31.1 _{0.4}
LAGONN _{cheap}	30.7 _{2.0}	30.5 _{3.8}	31.4 _{2.6}	31.0 _{0.4}

Table 18

Method	Liar			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	31.4 _{3.2}	35.8 _{2.6}	40.0 _{4.3}	36.2 _{2.4}
SetFit	32.3 _{4.5}	35.9 _{3.1}	36.4 _{2.2}	35.2 _{1.1}
LAGONN _{exp}	32.3 _{4.6}	35.7 _{3.4}	36.5 _{2.3}	35.7 _{1.4}
RoBERTa _{freeze}	31.4 _{3.2}	34.1 _{2.6}	35.6 _{3.2}	34.0 _{1.4}
kNN	27.0 _{0.2}	28.5 _{1.0}	29.0 _{1.0}	28.7 _{0.7}
Log Reg	32.3 _{4.5}	36.5 _{3.1}	38.5 _{3.4}	36.3 _{2.0}
LAGONN	32.3 _{4.6}	34.9 _{2.2}	36.9 _{2.5}	35.3 _{1.4}
Probe	30.7 _{3.0}	32.8 _{1.8}	35.0 _{1.6}	33.5 _{1.5}
LAGONN _{cheap}	30.4 _{3.0}	32.9 _{1.8}	35.4 _{1.7}	33.5 _{1.7}

Table 19

Method	Liar			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	33.9 _{3.1}	38.4 _{2.7}	43.9 _{2.2}	39.5 _{3.0}
SetFit	33.0 _{2.6}	37.2 _{1.8}	38.7 _{1.5}	37.4 _{1.6}
LAGONN _{exp}	34.1 _{3.4}	38.7 _{2.3}	39.0 _{1.8}	37.8 _{1.5}
RoBERTa _{freeze}	33.9 _{3.1}	35.3 _{2.6}	36.8 _{2.2}	35.4 _{1.0}
kNN	29.2 _{0.8}	29.7 _{1.5}	30.0 _{0.6}	29.8 _{0.3}
Log Reg	33.0 _{2.6}	37.2 _{3.9}	39.4 _{3.5}	37.0 _{1.8}
LAGONN	34.1 _{3.4}	37.0 _{3.1}	38.6 _{3.0}	36.8 _{1.3}
Probe	31.6 _{1.1}	34.7 _{2.5}	37.0 _{2.5}	34.9 _{1.7}
LAGONN _{cheap}	31.4 _{0.9}	35.3 _{2.3}	37.6 _{2.0}	35.3 _{1.9}

Table 20

Method	Liar			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	33.8 _{2.1}	39.4 _{2.4}	43.5 _{1.7}	40.2 _{3.2}
SetFit	34.4 _{2.3}	36.7 _{1.7}	37.0 _{1.3}	36.5 _{1.1}
LAGONN _{exp}	33.8 _{1.8}	34.2 _{2.7}	37.2 _{1.9}	36.2 _{1.4}
RoBERTa _{freeze}	33.8 _{2.1}	36.6 _{1.6}	38.6 _{1.5}	36.7 _{1.5}
kNN	30.1 _{0.4}	31.3 _{2.1}	30.6 _{1.1}	30.9 _{0.4}
Log Reg	34.4 _{2.3}	38.3 _{2.5}	40.0 _{2.0}	37.9 _{1.6}
LAGONN	33.8 _{1.8}	38.3 _{1.3}	40.6 _{0.6}	38.1 _{2.0}
Probe	32.1 _{1.9}	35.2 _{1.4}	37.2 _{2.5}	35.2 _{1.7}
LAGONN _{cheap}	31.9 _{1.9}	36.0 _{1.0}	37.5 _{2.5}	35.7 _{1.8}

Table 21

Method	Insincere Questions			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	19.9 _{8.4}	30.9 _{7.9}	42.0 _{7.4}	33.5 _{6.7}
SetFit	24.1 _{6.3}	29.2 _{6.7}	36.7 _{7.3}	31.7 _{3.4}
LAGONN _{exp}	30.7 _{8.9}	37.6 _{6.1}	39.0 _{6.1}	36.1 _{2.3}
SetFit _{lite}	24.1 _{6.3}	38.1 _{6.3}	41.1 _{6.5}	35.6 _{5.5}
LAGONN _{lite}	30.7 _{8.9}	41.8 _{8.3}	43.4 _{8.5}	39.3 _{4.4}
RoBERTa _{freeze}	19.9 _{8.4}	34.1 _{5.4}	37.9 _{5.2}	32.5 _{5.4}
kNN	6.8 _{0.4}	15.9 _{3.4}	16.9 _{4.3}	14.4 _{3.0}
Log Reg	24.1 _{6.3}	31.7 _{4.9}	36.1 _{5.4}	31.8 _{3.6}
LAGONN	30.7 _{8.9}	39.3 _{4.9}	41.2 _{4.7}	38.4 _{3.0}
Probe	24.3 _{8.4}	39.8 _{5.6}	44.8 _{4.2}	38.3 _{6.2}
LAGONN _{cheap}	23.6 _{7.8}	40.7 _{5.9}	45.3 _{4.4}	38.6 _{6.6}

Table 22

A.3 Additional results for second experiment

Here we provide additional results from our second set of experiments that, due to space limitations, could not be included in the main text. We note that a version of LAGONN outperforms or has the same performance of all methods, including our upper bound RoBERTa_{full}, on 60% of all displayed results, and is the best performer relative to Sentence Transformer-based methods on 65%. This excludes LAGONN_{cheap}. This method showed strong performance on the Insincere Questions dataset, but hurts performance in other cases. In cases when SetFit-based methods do outperform our system, the performances are comparable, yet they can be quite different when LAGONN-based methods are the strongest. Below, we report the mean average precision $\times 100$ for all methods over five seeds with the standard deviation, except in the case of Hate Speech Offensive, where the evaluation metric is the macro-F1. Each table shows the results for given dataset and a given label-balance distribution on the first, fifth, and tenth step followed by the average for all ten steps. Liar appears to be the most difficult dataset for all methods. This is expected because it likely does not include enough context to determine the truth of a statement.

Method	Insincere Questions			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	39.8 _{5.5}	53.1 _{4.6}	55.7 _{1.2}	50.6 _{4.4}
SetFit	43.7 _{2.7}	52.2 _{1.9}	53.8 _{0.9}	51.4 _{2.9}
LAGONN _{exp}	44.5 _{4.5}	52.7 _{2.4}	55.4 _{2.0}	51.8 _{3.0}
SetFit _{lite}	43.7 _{2.7}	52.9 _{2.6}	55.8 _{1.8}	52.2 _{3.4}
LAGONN _{lite}	44.5 _{4.5}	53.5 _{2.7}	55.9 _{2.4}	52.6 _{3.5}
RoBERTa _{freeze}	39.8 _{5.5}	44.1 _{3.6}	46.3 _{2.4}	44.0 _{2.0}
kNN	23.9 _{2.2}	30.3 _{3.0}	31.6 _{2.4}	30.0 _{2.1}
Log Reg	43.7 _{2.7}	47.6 _{1.6}	50.1 _{2.1}	47.6 _{1.8}
LAGONN	44.5 _{4.5}	48.1 _{2.2}	50.3 _{1.7}	48.1 _{1.9}
Probe	40.4 _{4.2}	49.4 _{2.3}	52.3 _{1.7}	49.0 _{3.3}
LAGONN _{cheap}	40.8 _{4.3}	51.1 _{2.4}	54.5 _{1.4}	50.4 _{4.0}

Table 23

Method	Insincere Questions			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	48.1 _{2.3}	54.7 _{1.9}	57.5 _{1.5}	53.9 _{2.9}
SetFit	48.9 _{1.7}	53.9 _{0.7}	54.2 _{1.5}	52.3 _{1.6}
LAGONN _{exp}	49.8 _{1.6}	52.2 _{1.9}	53.2 _{3.3}	52.0 _{1.4}
SetFit _{lite}	48.9 _{1.7}	56.5 _{1.4}	58.7 _{0.6}	55.0 _{3.5}
LAGONN _{lite}	49.8 _{1.6}	56.1 _{2.8}	58.3 _{1.5}	54.6 _{3.5}
RoBERTa _{freeze}	48.1 _{2.3}	50.2 _{2.2}	52.0 _{1.4}	50.2 _{1.4}
kNN	28.0 _{2.4}	33.9 _{2.8}	33.6 _{2.0}	33.5 _{1.9}
Log Reg	48.9 _{1.7}	53.6 _{1.9}	55.8 _{1.7}	53.3 _{2.2}
LAGONN	49.8 _{1.6}	54.4 _{1.3}	56.9 _{0.5}	54.2 _{2.2}
Probe	45.7 _{2.1}	52.3 _{1.8}	54.4 _{1.1}	51.4 _{2.5}
LAGONN _{cheap}	45.7 _{2.2}	54.4 _{1.6}	56.4 _{0.6}	53.2 _{3.2}

Table 24

Method	Insincere Questions			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	47.1 _{4.2}	52.1 _{3.6}	55.7 _{2.6}	52.5 _{2.9}
SetFit	43.5 _{4.2}	47.1 _{4.6}	48.5 _{3.9}	48.0 _{1.7}
LAGoNN _{exp}	42.8 _{5.3}	47.6 _{2.9}	47.0 _{1.7}	46.2 _{2.0}
SetFit _{lite}	43.5 _{4.2}	54.6 _{2.4}	59.6 _{0.9}	53.6 _{5.8}
LAGoNN _{lite}	42.8 _{5.3}	53.5 _{3.7}	58.6 _{2.5}	52.2 _{6.4}
RoBERTa _{freeze}	47.1 _{4.2}	52.1 _{0.4}	53.3 _{1.1}	51.5 _{2.1}
kNN	22.3 _{2.3}	30.2 _{2.3}	30.9 _{1.8}	29.5 _{2.5}
Log Reg	43.5 _{4.2}	53.8 _{2.2}	55.5 _{1.6}	52.8 _{3.5}
LAGoNN	42.8 _{5.3}	54.1 _{2.9}	56.3 _{1.3}	53.4 _{3.7}
Probe	47.5 _{1.6}	52.4 _{1.7}	55.3 _{1.1}	52.2 _{2.5}
LAGoNN _{cheap}	49.3 _{2.6}	54.4 _{1.4}	57.6 _{0.7}	54.2 _{2.7}

Table 25

Method	Amazon Counterfactual			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	21.8 _{6.6}	63.9 _{10.2}	72.3 _{3.0}	59.6 _{16.8}
SetFit	22.3 _{8.8}	64.2 _{3.3}	68.6 _{4.6}	56.8 _{14.9}
LAGoNN _{exp}	26.1 _{17.5}	68.4 _{4.4}	74.9 _{2.9}	63.2 _{16.7}
SetFit _{lite}	22.3 _{8.8}	62.4 _{5.1}	67.5 _{5.2}	56.5 _{14.7}
LAGoNN _{lite}	26.1 _{17.5}	68.3 _{4.3}	68.9 _{4.3}	60.6 _{15.1}
RoBERTa _{freeze}	21.8 _{6.6}	41.0 _{12.7}	51.3 _{10.7}	40.6 _{8.9}
kNN	10.3 _{0.2}	15.3 _{4.2}	18.4 _{3.7}	15.6 _{2.4}
Log Reg	22.3 _{8.8}	32.4 _{11.5}	42.3 _{8.8}	34.5 _{5.9}
LAGoNN	26.1 _{17.5}	31.1 _{19.4}	33.0 _{19.1}	30.9 _{2.3}
Probe	24.2 _{9.0}	46.3 _{4.4}	54.6 _{2.0}	45.1 _{10.3}
LAGoNN _{cheap}	20.1 _{6.9}	38.3 _{4.9}	47.8 _{3.4}	38.2 _{9.5}

Table 26

Method	Amazon Counterfactual			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	68.2 _{4.5}	81.0 _{1.7}	82.2 _{1.0}	79.2 _{3.9}
SetFit	72.0 _{2.1}	78.4 _{2.8}	78.8 _{1.2}	78.0 _{2.1}
LAGoNN _{exp}	74.3 _{3.8}	80.1 _{1.4}	79.0 _{1.6}	79.5 _{1.9}
SetFit _{lite}	72.0 _{2.1}	79.1 _{1.4}	81.6 _{1.3}	79.1 _{2.7}
LAGoNN _{lite}	74.3 _{3.8}	79.2 _{1.7}	81.9 _{1.1}	80.2 _{2.2}
RoBERTa _{freeze}	68.2 _{4.5}	75.0 _{2.2}	77.0 _{2.4}	74.2 _{2.6}
kNN	51.0 _{4.1}	60.0 _{3.1}	61.3 _{2.1}	59.7 _{3.0}
Log Reg	72.0 _{2.1}	74.4 _{2.3}	76.7 _{1.8}	74.8 _{1.4}
LAGoNN	74.3 _{3.8}	76.1 _{3.6}	77.3 _{3.2}	76.1 _{1.0}
Probe	46.6 _{2.8}	60.3 _{1.4}	64.2 _{1.2}	59.2 _{5.2}
LAGoNN _{cheap}	38.2 _{3.2}	55.3 _{1.8}	61.0 _{1.2}	54.4 _{6.7}

Table 27

Method	Amazon Counterfactual			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	73.9 _{2.5}	80.0 _{1.0}	80.1 _{2.3}	79.1 _{2.1}
SetFit	76.5 _{1.6}	77.0 _{2.4}	74.7 _{0.5}	76.5 _{1.0}
LAGoNN _{exp}	76.6 _{2.2}	78.0 _{2.1}	76.3 _{4.9}	78.2 _{1.0}
SetFit _{lite}	76.5 _{1.6}	80.4 _{3.8}	83.5 _{0.8}	80.3 _{2.8}
LAGoNN _{lite}	78.6 _{2.2}	80.8 _{1.9}	83.1 _{0.7}	81.0 _{1.7}
RoBERTa _{freeze}	73.9 _{2.5}	76.6 _{1.4}	78.5 _{0.7}	76.4 _{1.7}
kNN	54.5 _{3.1}	64.2 _{1.9}	66.6 _{1.3}	64.7 _{3.5}
Log Reg	76.5 _{1.6}	80.6 _{0.5}	81.2 _{0.3}	80.0 _{1.4}
LAGoNN	78.6 _{2.2}	81.2 _{1.4}	81.6 _{1.1}	80.8 _{0.9}
Probe	52.3 _{2.0}	64.1 _{1.8}	67.2 _{1.4}	63.1 _{4.3}
LAGoNN _{cheap}	47.3 _{3.4}	60.7 _{1.5}	65.2 _{1.4}	59.5 _{5.2}

Table 28

Method	Amazon Counterfactual			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	73.6 _{2.1}	78.6 _{3.9}	82.4 _{1.1}	78.9 _{2.2}
SetFit	73.8 _{4.4}	69.8 _{4.0}	64.1 _{4.6}	69.6 _{3.6}
LAGoNN _{exp}	76.0 _{3.0}	73.4 _{2.6}	72.3 _{2.9}	72.5 _{3.4}
SetFit _{lite}	73.8 _{4.4}	80.4 _{1.8}	82.4 _{0.8}	78.3 _{4.3}
LAGoNN _{lite}	76.0 _{3.0}	80.0 _{1.3}	82.5 _{0.9}	79.2 _{3.2}
RoBERTa _{freeze}	73.6 _{2.1}	76.8 _{1.6}	77.9 _{1.0}	76.5 _{1.3}
kNN	41.7 _{3.4}	57.9 _{3.3}	58.3 _{3.3}	56.8 _{5.1}
Log Reg	73.8 _{4.4}	79.2 _{1.9}	80.1 _{1.0}	78.6 _{1.8}
LAGoNN	76.0 _{3.0}	80.1 _{2.0}	81.4 _{1.1}	79.8 _{1.4}
Probe	52.4 _{3.4}	64.7 _{2.5}	67.5 _{0.4}	63.4 _{4.4}
LAGoNN _{cheap}	48.1 _{3.4}	62.0 _{2.0}	65.3 _{0.8}	60.5 _{5.0}

Table 29

Method	Toxic Conversations			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	7.9 _{0.5}	21.2 _{3.7}	33.8 _{5.5}	21.9 _{9.3}
SetFit	8.8 _{1.2}	18.1 _{3.4}	24.7 _{4.1}	17.6 _{5.5}
LAGoNN _{exp}	8.9 _{1.7}	17.4 _{6.6}	26.4 _{5.2}	17.9 _{6.0}
SetFit _{lite}	8.8 _{1.2}	15.9 _{4.8}	18.0 _{3.9}	14.9 _{3.2}
LAGoNN _{lite}	8.9 _{1.7}	16.1 _{5.9}	19.8 _{6.0}	15.5 _{3.7}
RoBERTa _{freeze}	7.9 _{0.5}	12.8 _{2.4}	19.1 _{3.2}	13.5 _{3.5}
kNN	7.9 _{0.0}	8.7 _{0.4}	8.7 _{0.2}	8.5 _{0.3}
Log Reg	8.8 _{1.2}	13.1 _{2.5}	16.3 _{3.0}	13.0 _{2.6}
LAGoNN	8.9 _{1.7}	13.8 _{3.9}	17.1 _{4.8}	13.4 _{2.6}
Probe	13.1 _{2.8}	24.6 _{2.6}	30.1 _{2.1}	23.9 _{5.6}
LAGoNN _{cheap}	11.3 _{2.2}	21.7 _{2.7}	27.4 _{2.3}	21.3 _{5.3}

Table 30

Method	Toxic Conversations			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	24.1 _{5.6}	43.1 _{3.4}	52.1 _{2.5}	42.4 _{8.2}
SetFit	21.8 _{6.6}	44.5 _{4.1}	51.4 _{1.9}	42.1 _{9.3}
LAGoNN _{exp}	22.7 _{9.8}	49.1 _{5.6}	53.4 _{2.3}	45.6 _{9.8}
SetFit _{lite}	21.8 _{6.6}	41.4 _{4.4}	44.8 _{3.1}	39.0 _{7.0}
LAGoNN _{lite}	22.7 _{9.8}	47.0 _{6.3}	50.2 _{5.4}	43.7 _{8.6}
RoBERTa _{freeze}	24.1 _{5.6}	31.2 _{4.4}	34.0 _{4.0}	30.5 _{3.1}
kNN	11.5 _{2.5}	14.7 _{4.0}	15.3 _{3.2}	14.6 _{1.1}
Log Reg	21.8 _{6.6}	26.7 _{5.3}	30.2 _{4.0}	26.6 _{2.7}
LAGoNN	22.7 _{9.8}	27.6 _{8.9}	30.3 _{8.7}	27.4 _{2.4}
Probe	23.3 _{2.7}	33.0 _{2.8}	37.1 _{1.8}	32.5 _{4.2}
LAGoNN _{cheap}	20.5 _{3.2}	31.1 _{3.2}	35.6 _{1.8}	30.5 _{4.6}

Table 31

Method	Toxic Conversations			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	34.2 _{3.4}	45.5 _{1.9}	52.4 _{3.3}	45.7 _{5.6}
SetFit	33.6 _{2.9}	47.2 _{2.2}	46.6 _{3.3}	44.3 _{4.3}
LAGoNN _{exp}	36.6 _{4.2}	48.2 _{2.7}	49.9 _{3.7}	48.0 _{4.4}
SetFit _{lite}	33.6 _{2.9}	52.6 _{2.0}	55.1 _{1.6}	48.8 _{7.3}
LAGoNN _{lite}	36.6 _{4.2}	56.1 _{1.5}	57.7 _{1.4}	52.3 _{6.8}
RoBERTa _{freeze}	34.2 _{3.4}	38.4 _{2.1}	39.5 _{1.8}	38.0 _{1.5}
kNN	19.4 _{1.9}	21.5 _{3.4}	22.4 _{2.9}	21.6 _{0.8}
Log Reg	33.6 _{2.9}	39.2 _{2.9}	41.6 _{2.7}	38.6 _{2.4}
LAGoNN	36.6 _{4.2}	42.7 _{3.7}	45.0 _{3.5}	42.0 _{2.5}
Probe	29.0 _{2.7}	36.1 _{1.2}	39.1 _{1.5}	35.5 _{3.3}
LAGoNN _{cheap}	26.1 _{2.7}	34.3 _{1.3}	37.5 _{1.8}	33.6 _{3.6}

Table 32

Method	Hate Speech Offensive			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	50.6 _{3.0}	65.2 _{3.9}	70.3 _{1.2}	64.2 _{5.3}
SetFit	54.4 _{4.3}	66.3 _{1.8}	68.9 _{2.0}	64.3 _{4.5}
LAGoNN _{exp}	57.0 _{5.2}	67.0 _{4.4}	69.8 _{2.1}	64.9 _{4.6}
SetFit _{lite}	54.4 _{4.3}	65.5 _{3.0}	65.9 _{3.5}	63.5 _{3.9}
LAGoNN _{lite}	57.0 _{5.2}	66.6 _{2.6}	66.6 _{1.9}	64.3 _{4.1}
RoBERTa _{freeze}	50.6 _{3.0}	54.1 _{1.6}	55.3 _{2.3}	54.1 _{1.3}
kNN	55.6 _{4.8}	57.3 _{2.3}	58.8 _{3.6}	57.4 _{1.1}
Log Reg	54.4 _{4.3}	57.0 _{3.9}	58.2 _{3.8}	57.2 _{1.1}
LAGoNN	57.0 _{5.2}	58.2 _{4.1}	58.3 _{3.4}	58.3 _{0.6}
Probe	46.5 _{2.2}	57.8 _{1.7}	60.3 _{1.2}	56.5 _{4.5}
LAGoNN _{cheap}	47.1 _{1.3}	56.5 _{2.2}	59.5 _{2.5}	55.6 _{3.8}

Table 35

Method	Toxic Conversations			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	32.3 _{1.1}	42.7 _{1.8}	54.1 _{3.4}	43.8 _{6.3}
SetFit	35.7 _{3.4}	32.6 _{6.2}	37.4 _{2.7}	36.5 _{1.9}
LAGoNN _{exp}	40.4 _{4.4}	40.2 _{6.6}	39.8 _{7.5}	40.0 _{1.2}
SetFit _{lite}	35.7 _{3.4}	52.7 _{2.5}	53.9 _{2.2}	46.8 _{7.8}
LAGoNN _{lite}	40.4 _{4.4}	52.9 _{2.6}	54.0 _{2.3}	48.3 _{6.4}
RoBERTa _{freeze}	32.3 _{1.1}	39.2 _{1.5}	41.0 _{0.6}	38.5 _{2.4}
kNN	17.4 _{0.8}	23.7 _{2.6}	24.3 _{2.7}	23.1 _{2.0}
Log Reg	35.7 _{3.4}	44.5 _{2.9}	46.1 _{2.8}	43.6 _{2.9}
LAGoNN	40.4 _{4.4}	46.6 _{2.7}	48.1 _{2.2}	46.1 _{2.2}
Probe	29.5 _{2.4}	35.9 _{0.9}	40.2 _{0.9}	36.1 _{3.5}
LAGoNN _{cheap}	26.8 _{2.7}	34.5 _{1.3}	38.5 _{0.8}	34.4 _{3.7}

Table 33

Method	Hate Speech Offensive			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	61.9 _{3.4}	70.8 _{1.0}	72.5 _{1.4}	69.9 _{3.2}
SetFit	64.3 _{4.2}	70.6 _{2.4}	72.4 _{0.5}	69.8 _{2.8}
LAGoNN _{exp}	63.8 _{4.9}	71.0 _{2.1}	72.3 _{1.0}	70.0 _{3.0}
SetFit _{lite}	64.3 _{4.2}	70.3 _{2.2}	71.2 _{2.1}	69.3 _{2.3}
LAGoNN _{lite}	63.8 _{4.9}	70.7 _{1.4}	71.4 _{1.0}	69.4 _{2.5}
RoBERTa _{freeze}	61.9 _{3.4}	63.2 _{4.1}	64.1 _{4.5}	63.2 _{0.6}
kNN	64.3 _{4.0}	63.3 _{2.9}	63.9 _{2.5}	63.7 _{0.4}
Log Reg	64.3 _{4.2}	67.3 _{3.2}	67.6 _{2.3}	66.9 _{1.1}
LAGoNN	63.8 _{4.9}	65.0 _{5.3}	66.7 _{5.9}	65.3 _{0.9}
Probe	55.6 _{1.7}	63.8 _{0.8}	66.1 _{0.3}	63.2 _{3.0}
LAGoNN _{cheap}	56.0 _{3.6}	62.2 _{1.4}	66.0 _{0.9}	62.3 _{2.9}

Table 36

Method	Hate Speech Offensive			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	30.2 _{1.4}	43.5 _{2.5}	51.2 _{2.2}	44.3 _{7.4}
SetFit	30.3 _{0.8}	44.0 _{1.3}	51.1 _{2.0}	43.8 _{6.5}
LAGoNN _{exp}	30.3 _{0.7}	40.7 _{2.9}	49.1 _{4.4}	42.2 _{6.2}
SetFit _{lite}	30.3 _{0.8}	43.4 _{2.5}	45.5 _{3.4}	41.6 _{4.6}
LAGoNN _{lite}	30.3 _{0.7}	40.9 _{3.4}	41.5 _{4.8}	39.1 _{3.6}
RoBERTa _{freeze}	30.2 _{1.4}	33.5 _{3.1}	34.4 _{3.4}	33.1 _{1.4}
kNN	31.5 _{1.2}	35.9 _{2.7}	37.4 _{2.0}	35.8 _{1.7}
Log Reg	30.3 _{0.8}	38.4 _{2.5}	41.1 _{1.5}	37.8 _{3.3}
LAGoNN	30.3 _{0.7}	35.7 _{2.6}	39.1 _{2.4}	35.6 _{2.7}
Probe	29.0 _{0.2}	34.7 _{1.5}	40.1 _{2.1}	35.1 _{3.8}
LAGoNN _{cheap}	29.0 _{0.1}	36.9 _{1.8}	40.5 _{2.1}	36.2 _{3.7}

Table 34

Method	Hate Speech Offensive			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	59.7 _{3.5}	66.9 _{1.2}	69.2 _{1.8}	66.4 _{2.7}
SetFit	60.7 _{1.3}	66.3 _{1.6}	67.5 _{0.9}	65.9 _{2.2}
LAGoNN _{exp}	61.5 _{1.7}	66.4 _{1.4}	67.7 _{0.9}	66.1 _{1.8}
SetFit _{lite}	60.7 _{1.3}	66.3 _{2.0}	66.5 _{0.9}	65.1 _{1.7}
LAGoNN _{lite}	61.5 _{1.7}	67.1 _{1.1}	67.3 _{0.8}	66.0 _{1.7}
RoBERTa _{freeze}	59.7 _{3.5}	60.4 _{2.7}	63.1 _{2.3}	61.0 _{1.3}
kNN	60.7 _{1.3}	59.6 _{2.8}	59.5 _{2.5}	59.5 _{0.5}
Log Reg	60.7 _{1.3}	62.5 _{0.7}	63.4 _{1.0}	62.3 _{1.0}
LAGoNN	61.5 _{1.7}	62.8 _{1.5}	64.2 _{1.0}	63.0 _{0.9}
Probe	54.9 _{1.4}	58.5 _{0.9}	60.9 _{0.4}	58.7 _{1.7}
LAGoNN _{cheap}	54.2 _{2.3}	58.6 _{0.6}	60.6 _{0.5}	58.5 _{1.8}

Table 37

Method	Liar			
<i>Extreme</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	32.0 _{2.7}	34.7 _{2.9}	35.1 _{4.3}	33.7 _{1.0}
SetFit	31.2 _{3.8}	30.4 _{3.1}	31.8 _{2.9}	31.5 _{0.7}
LAGONN _{exp}	30.6 _{4.7}	30.3 _{2.0}	31.3 _{2.0}	31.1 _{0.6}
SetFit _{lite}	31.2 _{3.8}	32.7 _{3.8}	33.5 _{4.2}	32.7 _{0.8}
LAGONN _{lite}	30.6 _{4.7}	31.8 _{3.9}	32.4 _{2.7}	31.6 _{0.6}
RoBERTa _{freeze}	32.0 _{2.7}	32.8 _{4.5}	34.2 _{5.0}	33.2 _{0.7}
kNN	27.0 _{0.5}	27.3 _{0.8}	27.9 _{0.8}	27.4 _{0.3}
Log Reg	31.2 _{3.8}	33.7 _{5.1}	35.7 _{5.1}	34.3 _{1.6}
LAGONN	30.6 _{4.7}	32.0 _{4.6}	33.7 _{5.4}	32.6 _{0.9}
Probe	30.7 _{2.0}	30.6 _{3.9}	31.7 _{2.9}	31.1 _{0.4}
LAGONN _{cheap}	30.7 _{2.0}	30.5 _{3.8}	31.4 _{2.6}	31.0 _{0.4}

Table 38

Method	Liar			
<i>Imbalanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	31.4 _{3.2}	35.8 _{2.6}	40.0 _{4.3}	36.2 _{2.4}
SetFit	32.3 _{4.5}	35.9 _{3.1}	36.4 _{2.2}	35.2 _{1.1}
LAGONN _{exp}	32.3 _{4.6}	35.7 _{3.4}	36.5 _{2.3}	35.7 _{1.4}
SetFit _{lite}	32.3 _{4.5}	35.6 _{2.7}	37.4 _{2.6}	35.8 _{1.6}
LAGONN _{lite}	32.3 _{4.6}	35.2 _{2.4}	36.6 _{2.7}	35.5 _{1.3}
RoBERTa _{freeze}	31.4 _{3.2}	34.1 _{2.6}	35.6 _{3.2}	34.0 _{1.4}
kNN	27.0 _{0.2}	28.5 _{1.0}	29.0 _{1.0}	28.7 _{0.7}
Log Reg	32.3 _{4.5}	36.5 _{3.1}	38.5 _{3.4}	36.3 _{2.0}
LAGONN	32.3 _{4.6}	34.9 _{2.2}	36.9 _{2.5}	35.3 _{1.4}
Probe	30.7 _{3.0}	32.8 _{1.8}	35.0 _{1.6}	33.5 _{1.5}
LAGONN _{cheap}	30.4 _{3.0}	32.9 _{1.8}	35.4 _{1.7}	33.5 _{1.7}

Table 39

Method	Liar			
<i>Moderate</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	33.9 _{3.1}	38.4 _{2.7}	43.9 _{2.2}	39.5 _{3.0}
SetFit	33.0 _{2.6}	37.2 _{1.8}	38.7 _{1.5}	37.4 _{1.6}
LAGONN _{exp}	34.1 _{3.4}	38.7 _{2.3}	39.0 _{1.8}	37.8 _{1.5}
SetFit _{lite}	33.0 _{2.6}	38.5 _{1.3}	40.4 _{2.0}	38.2 _{2.1}
LAGONN _{lite}	34.1 _{3.4}	38.4 _{2.0}	39.6 _{1.5}	37.9 _{1.6}
RoBERTa _{freeze}	33.9 _{3.1}	35.3 _{2.6}	36.8 _{2.2}	35.4 _{1.0}
kNN	29.2 _{0.8}	29.7 _{1.5}	30.0 _{0.6}	29.8 _{0.3}
Log Reg	33.0 _{2.6}	37.2 _{3.9}	39.4 _{3.5}	37.0 _{1.8}
LAGONN	34.1 _{3.4}	37.0 _{3.1}	38.6 _{3.0}	36.8 _{1.3}
Probe	31.6 _{1.1}	34.7 _{2.5}	37.0 _{2.5}	34.9 _{1.7}
LAGONN _{cheap}	31.4 _{0.9}	35.3 _{2.3}	37.6 _{2.0}	35.3 _{1.9}

Table 40

Method	Liar			
<i>Balanced</i>	1 st	5 th	10 th	Average
RoBERTa _{full}	33.8 _{2.1}	39.4 _{2.4}	43.5 _{1.7}	40.2 _{3.2}
SetFit	34.4 _{2.3}	36.7 _{1.7}	37.0 _{1.3}	36.5 _{1.1}
LAGONN _{exp}	33.8 _{1.8}	34.2 _{2.7}	37.2 _{1.9}	36.2 _{1.4}
SetFit _{lite}	34.4 _{2.3}	38.7 _{2.3}	40.3 _{2.8}	38.0 _{2.1}
LAGONN _{lite}	33.8 _{1.8}	37.6 _{2.0}	39.4 _{2.8}	37.2 _{1.9}
RoBERTa _{freeze}	33.8 _{2.1}	36.6 _{1.6}	38.6 _{1.5}	36.7 _{1.5}
kNN	30.1 _{0.4}	31.3 _{2.1}	30.6 _{1.1}	30.9 _{0.4}
Log Reg	34.4 _{2.3}	38.3 _{2.5}	40.0 _{2.0}	37.9 _{1.6}
LAGONN	33.8 _{1.8}	38.3 _{1.3}	40.6 _{0.6}	38.1 _{2.0}
Probe	32.1 _{1.9}	35.2 _{1.4}	37.2 _{2.5}	35.2 _{1.7}
LAGONN _{cheap}	31.9 _{1.9}	36.0 _{1.0}	37.5 _{2.5}	35.7 _{1.8}

Table 41

A.4 Examples of LAGONN modified text

WARNING: Some of the examples below are of an offensive nature. Please view with caution.

In this section, we provide examples of how LAGONN_{exp} modifies test text from the datasets we studied under the BOTH configuration. We choose this configuration because the information it appends from a NN in the training data to a test instance encapsulates both the LABEL and the TEXT configuration. LAGONN_{exp} was trained under a balanced distribution and five examples per label were chosen randomly on the first, fifth, and tenth step to demonstrate how the same test instance might be decorated with different training examples as the training data grow. We recognize that some the images below are difficult to see and have made the .csv files available with our code and data files. Note that MPNET’s separator token is </s>, not [SEP].

Test Modified	Gold Label
What rapper still relevant and popular today has the best rhyme schemes? </> </s> <insincere question 3.859471321105957> What would be a good nickname for Trump, Donald Dumbck, and President Spankovich? </> </s> <invalid question 4.124274253845315> What are after class 12 courses in commerce stream to choose from? I have completed my class 12 [expected 90+] and aim to do business (not aim to do job).	valid question
Which books do you suggest to someone who get a free time and will help him stay motivated? </> </s> <invalid question 3.9508353637695312> What are the best online courses to learn data science? </> </s> <insincere question 4.300448417663574> What are the more steps in Career Oriented Education?	valid question
How will you feel if someone talks badly about Kunt? </> </s> <insincere question 3.5063605308532715> Why are the UK government and the media (especially the BBC and the Guardian) demonising ordinary British people, manipulating buzz words like &racist-right&, &Islamophobia&, &racist& to suppress legitimate outrage at Muslim grooming gangs? </> </s> <valid question 3.609603751879883> How do Israelis and Palestinian view Nasser Yassin?	valid question
Why is equine HIPP inherited? </> </s> <valid question 4.066534996032715> Can you share some of the pics of hotel of Indian Gandhi medical college, Shimla, Himachal Pradesh? </> </s> <insincere question 4.23175760650635> I am an experienced programmer and in my high school my teacher tried to make me use python so I said, "No, Trust me, python is just a language for beginners, thereby making it not for me." I got sent out. Did I do anything wrong?	valid question
How do the Valerie Stevens leather jackets achieve their quality during the manufacturing process? </> </s> <valid question 3.9721384048461914> How are the Lancaster leather sofas manufactured? </> </s> <insincere question 4.3559441566467285> I am an experienced programmer and in my high school my teacher tried to make me use python so I said, "No, Trust me, python is just a language for beginners, thereby making it not for me." I got sent out. Did I do anything wrong?	valid question
Is Ariana Grande really as mean and bitchy as she seems? </> </s> <insincere question 3.572277545928955> Why is Alla Bhatt so dumb? </> </s> <valid question 3.924571950966797> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question
Do you guys know that aliens are real and all those satellites we send up in space work as a sort of tracking device for them so in a few years it will be too late for Earth? </> </s> <insincere question 3.5094439029693604> Have you noticed how conservatives are capturing the English language and modifying the definitions of political words? </> </s> <valid question 3.6931655197343555> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question
Is it politically incorrect to say female privilege, but it is a more accurate term to say, white female privilege? </> </s> <insincere question 3.323280096054077> Why are the UK government and the media (especially the BBC and the Guardian) demonising ordinary British people, manipulating buzz words like &racist-right&, &Islamaphobia&, &racist& to suppress legitimate outrage at Muslim grooming gangs? </> </s> <valid question 3.986680269241333> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question
On Mother's Day, is it reasonable to reflect there is some truth in the unfashionable notion than women are more driven by emotion and men more driven by reason? </> </s> <insincere question 3.499204158782959> Why are the UK government and the media (especially the BBC and the Guardian) demonising ordinary British people, manipulating buzz words like &racist-right&, &Islamaphobia&, &racist& to suppress legitimate outrage at Muslim grooming gangs? </> </s> <valid question 3.771740198135378> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question
If the U.S. president is a role model, is it acceptable for children to say "go fuck yourself," and use the word "pussy" to describe women? </> </s> <insincere question 3.497847080230715> Why are the UK government and the media (especially the BBC and the Guardian) demonising ordinary British people, manipulating buzz words like &racist-right&, &Islamaphobia&, &racist& to suppress legitimate outrage at Muslim grooming gangs? </> </s> <valid question 3.845909357070923> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question

Figure 5: Insincere Questions, step 1.

Test Modified	Gold Label
What rapper still relevant and popular today has the best rhyme schemes? </> </s> <insincere question 3.87190794947632> What would be a good nickname for Trump, Donald Dumbck, and President Spankovich? </> </s> <invalid question 4.028958797454834> Why does Dancing with the Stars not include Bachata as one their dance styles?	valid question
Which books do you suggest to someone who get a free time and will help him stay motivated? </> </s> <valid question 3.6081225872039795> What is a good degree to get at community college if you want to explore different subjects and figure out your career path? </> </s> <insincere question 3.8502604961395264> What are the more steps in Career Oriented Education?	valid question
How will you feel if someone talks badly about Kunt? </> </s> <valid question 3.535556163757324> How do I stop feeling bad after a girl had a crush on me? </> </s> <insincere question 3.689171075820923> Why Indian girls go crazy about marrying Shri. Rahul Gandhi ji?	valid question
Why is equine HIPP inherited? </> </s> <insincere question 3.6035702228546143> Can female animals with male humans sex? </> </s> <valid question 3.79418632054901123> How long do guinea pigs live for?	valid question
How do the Valerie Stevens leather jackets achieve their quality during the manufacturing process? </> </s> <valid question 3.747288217081299> How are the Lancaster leather sofas manufactured? </> </s> <insincere question 3.944884770690292> Why don't all Trump supporters buy only made in USA goods, e.g. many of them have their cars of Asian/European companies, shop in places where more than 70% of items are not made in USA, eat multi-national cuisine or otherwise stop their hypocrisy?	valid question
Is Ariana Grande really as mean and bitchy as she seems? </> </s> <insincere question 3.3252298831939697> Why is Alla Bhatt so dumb? </> </s> <valid question 3.7413415908813477> How do I stop feeling bad after a girl had a crush on me?	insincere question
Do you guys know that aliens are real and all those satellites we send up in space work as a sort of tracking device for them so in a few years it will be too late for Earth? </> </s> <insincere question 3.0673365592956543> Isn't it obvious now that walking on the moon by the Americans was a hoax, because walking on the bright side of the moon, even in a space suit would be fatal? </> </s> <valid question 3.1978228092193604> Why do we weunch satellites?	insincere question
Is it politically incorrect to say female privilege, but it is a more accurate term to say, white female privilege? </> </s> <insincere question 2.9176881217956035> How does the privilege of being attractive compare to the privilege of being White in the US? </> </s> <valid question 3.11248117248535> Is the media wrong for enforcing gender stereotypes?	insincere question
On Mother's Day, is it reasonable to reflect there is some truth in the unfashionable notion than women are more driven by emotion and men more driven by reason? </> </s> <insincere question 3.102353811264038> Do women look down on men who are single, even if the man is more successful in other aspects of his life? </> </s> <valid question 3.1890125274658203> Why are some women uninterested in sex?	insincere question
If the U.S. president is a role model, is it acceptable for children to say "go fuck yourself," and use the word "pussy" to describe women? </> </s> <insincere question 3.163693904876709> Is it wrong to take your retarded son to a hooker for his 21st birthday? </> </s> <valid question 3.456286907196045> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question

Figure 6: Insincere Questions, step 5.

Test Modified	Gold Label
What rapper still relevant and popular today has the best rhyme schemes? </> </s> <valid question 3.7103171348571777> What is the oldest fashion trends running yet? </> </s> <insincere question 3.871907949447632> What would be a good nickname for Trump, Donald Dumbck, and President Spankovich?	valid question
Which books do you suggest to someone who get a free time and will help him stay motivated? </> </s> <valid question 3.1401429176330566> How can I stay motivated when learning something new? </> </s> <insincere question 3.735560417175293> I'm hungry and i'm too lazy too get out of bed, should I get a psychologist or ask you questions?	valid question
How will you feel if someone talks badly about Kunt? </> </s> <insincere question 3.4893462657928467> Does Tamil Isai Soudanar support Vijayendra for disrespecting the Tamil Anthem? </> </s> <valid question 3.535556163757324> How do I stop feeling bad after a girl had a crush on me?	valid question
Why is equine HIPP inherited? </> </s> <insincere question 3.3252298831939697> Why is Alla Bhatt so dumb? </> </s> <valid question 3.7413415908813477> How do I stop feeling bad after a girl had a crush on me?	valid question
How do the Valerie Stevens leather jackets achieve their quality during the manufacturing process? </> </s> <valid question 2.747288217081299> How are the Lancaster leather sofas manufactured? </> </s> <insincere question 3.9087233543395996> Are Newport cigarettes designed to selectively destroy black people's DNA?	valid question
Is Ariana Grande really as mean and bitchy as she seems? </> </s> <valid question 3.183567762374878> Is the girl who used to be quite rude and would run through boyfriends very fast. But now that school started again, she seems to have gotten a lot nicer throughout Summer. Is she faking her politeness, and is it worth pursuing her? </> </s> <insincere question 3.325866220240845971438> These boxer/briefs are very soft, very comfortable, and fit like high-end underwear the likes of which you might get at, oh, say, Calvin Klein for example, but for about half the price.	insincere question
Do you guys know that aliens are real and all those satellites we send up in space work as a sort of tracking device for them so in a few years it will be too late for Earth? </> </s> <insincere question 3.0673365592956543> Isn't it obvious now that walking on the moon by the Americans was a hoax, because walking on the bright side of the moon, even in a space suit would be fatal? </> </s> <valid question 3.1978228092193604> Why do we weunch satellites?	insincere question
Is it politically incorrect to say female privilege, but it is a more accurate term to say, white female privilege? </> </s> <insincere question 2.9176881217956035> How does the privilege of being attractive compare to the privilege of being White in the US? </> </s> <valid question 3.11248117248535> Is the media wrong for enforcing gender stereotypes?	insincere question
On Mother's Day, is it reasonable to reflect there is some truth in the unfashionable notion than women are more driven by emotion and men more driven by reason? </> </s> <insincere question 2.9901626110076904> Do you agree that females think with their brains and males with their testicles? </> </s> <valid question 3.1890125274658203> Why are some women uninterested in sex?	insincere question
If the U.S. president is a role model, is it acceptable for children to say "go fuck yourself," and use the word "pussy" to describe women? </> </s> <insincere question 2.994286298751831> Why do feminists let their daughters have sex with their boyfriend's at home? </> </s> <valid question 3.456286907196045> Do you agree with Congressman Steve King's comments on immigrant children in detention centers?	insincere question

Figure 7: Insincere Questions, step 10.

Test Modified	Gold Label
Clinging to the wall, doesn't flop around when a bag is pulled out, the mess of bags falling out is gone. </> </s> <not-counterfactual 3.6492726802825928> Hopes that it will keep it's shape after washing. </> </s> <counterfactual 4.012346267700195> ""Hlad I reviewed this immediately I would have given this product five stars because it worked.""	not-counterfactual
I like these jeans they sit low enough without being inappropriate when you sit or bend over. </> </s> <counterfactual 3.402600049972534> "But oddly enough, the bottoms are a little too loose in the waist [37] and could have used another inch or two in the inseam (I normally take a 35"" or 36"" in jeans, depending on the brand if this helps)."" </> </s> <not-counterfactual 3.420140845971438> These boxer/briefs are very soft, very comfortable, and fit like high-end underwear the likes of which you might get at, oh, say, Calvin Klein for example, but for about half the price.	not-counterfactual
He was very professional and wish all transactions I make through Amazon were this good. </> </s> <not-counterfactual 3.439908618027> I wish I had had him as an instructor at college. </> </s> <not-counterfactual 4.054030895233154> I worried that it would be cheap or not fit or...whatever... But WOW!	not-counterfactual
Well written with a twist I didn't expect. </> </s> <not-counterfactual 3.3257973194122134> "The crossover from the characters from one novel to others keeps me interested, after all, I do hate to miss a 'Dee-Ann or Eggie'" appearance."" </> </s> <counterfactual 3.682003021240234> ""Hlad I reviewed this immediately I would have given this product five stars because I worked.""	not-counterfactual
Doesn't feel like the quality lev's I am used to. </> </s> <not-counterfactual 3.2773308753967285> However, the fabric is not that great. It's cheap scratchy cotton. </> </s> <counterfactual 3.746659755706787> "The blanket is nice and soft but it is white, so if it doesn't light up it isn't much use!	not-counterfactual
If we had wall studs, I believe the enclosed hardware would have been sufficient. </> </s> <counterfactual 3.4338643550872803> I wish the storage compartment was a little bigger and opened up instead of sliding on and off. </> </s> <not-counterfactual 3.978530887890625> I worried that it would be cheap or not fit or...whatever... But WOW!	counterfactual
If this ever turns into a film, I hope they do it justice! </> </s> <not-counterfactual 3.5291523933410645> "The crossover from the characters from one novel to others keeps me interested, after all, I do hate to miss a 'Dee-Ann or Eggie'" appearance."" </> </s> <counterfactual 3.7511432172086792> ""Hlad I reviewed this immediately I would have given this product five stars because I worked.""	counterfactual
If you don't want a prominent display this rack is too large for most bed or living rooms, it is wider and taller than my tall Brooklyn wardrobe style dresser which was the largest piece in the room until this shoe rack. </> </s> <not-counterfactual 3.865670680999756> "It also validates ""(the incorrect)"" assumption that we are alone in the feelings we suppress when we sense the complete garbage that is thrown out into society."" </> </s> <counterfactual 4.063361167907715> The blanket is nice and soft but it is white, so if it doesn't light up it isn't much use!	counterfactual
I wish I could have seen all of the places he recommend! </> </s> <counterfactual 3.56270662823975> I wish I had had him as an instructor at college. </> </s> <not-counterfactual 4.141315937042236> I worried that it would be cheap or not fit or...whatever... But WOW!	counterfactual
I wish I could replace just that small stupid piece, since there's nothing wrong with the rest of the hose assembly. </> </s> <counterfactual 3.6057372093200684> I wish the storage compartment was a little bigger and opened up instead of sliding on and off. </> </s> <not-counterfactual 4.06487131187744> I worried that it would be cheap or not fit or...whatever... But WOW!	counterfactual

Figure 8: Amazon Counterfactual, step 1.

Test Modified	Gold Label
Clinging to the wall, doesn't flop around when a bag is pulled out, the mess of bags falling out is gone. </> </s> <not-counterfactual 3.161406993865967> And the dvd cases were tightly packed to ensure they didn't move around. </> </s> <counterfactual 3.308583974858257> The case is small, cord seems to always want to stay linked and coiled, plug should be angled and not straight... which are all items that others have pointed out.	not-counterfactual
I like these jeans they sit low enough without being inappropriate when you sit or bend over. </> </s> <counterfactual 2.606198310852051> "But oddly enough, the bottoms are a little too loose in the waist [37] and could have used another inch or two in the inseam (I normally take a 35"" or 36"" in jeans, depending on the brand if this helps)."" </> </s> <not-counterfactual 3.2680045413970947> A tad loose but I rather have it fit this way than too tight.	not-counterfactual
He was very professional and wish all transactions I make through Amazon were this good. </> </s> <not-counterfactual 3.3291680812835693> This new speaker was just what the doctor ordered and I couldn't be more pleased. </> </s> <counterfactual 3.4589436054229736> Had the person handling the shipping of this item been at all concerned with the use of the product at the end of the mailing process, the slightest bit of care could have been taken to ensure it's proper delivery.	not-counterfactual
Well written with a twist I didn't expect. </> </s> <not-counterfactual 2.65165852050662> The book had some interesting twists that I did use coming and I look forward to reading part two of this series. </> </s> <counterfactual 2.8373162746429443> Fun read Could have been a little longer with more detail.	not-counterfactual
Doesn't feel like the quality lev's I am used to. </> </s> <counterfactual 2.728787182006836> It has the same great comfortable & flattering features plus the great denim texture that Lee has perfected--smoothing and stretchy without the excessive cling--but I think it must have been designed for people who have a greater surplus of belly than I. </> </s> <not-counterfactual 2.856729545864668> Will keep but won't be that casual sexy top you always want to turn to.	not-counterfactual
If we had wall studs, I believe the enclosed hardware would have been sufficient. </> </s> <not-counterfactual 2.6638145446777344> It was a little tricky to find the center of the studs using my stud finder but once I felt comfortable with the lines I had drawn, I drilled the pilot holes and bolted this thing to the wall. </> </s> <counterfactual 2.879924774699922> The only thing I would have like for it to have a hole in the middle so I can put the stopper in without removing the mat.	counterfactual
If this ever turns into a film, I hope they do it justice! </> </s> <not-counterfactual 2.671574354171753> I read this book because of the motion picture that is coming out soon. </> </s> <counterfactual 3.1458709239959717> Was a good story, though there could have been more to it.	counterfactual
If you don't want a prominent display this rack is too large for most bed or living rooms, it is wider and taller than my tall Brooklyn wardrobe style dresser which was the largest piece in the room until this shoe rack. </> </s> <counterfactual 2.73537683486938448> I bought this mount because I wanted one that would sit on three studs instead of two because my TV is quite heavy and I would have had a hard time centering it on my wall if I didn't have the wide hanging rail that this one has. </> </s> <not-counterfactual 2.87861712241213> Good for under the bed shoe storage. If the wife wants to use it.	counterfactual
I wish I could have seen all of the places he recommend! </> </s> <counterfactual 2.799947738647461> I wish I had had him as an instructor at college. </> </s> <not-counterfactual 3.3013432025908426> And as the ole man isn't any version of slender it was good that he got to try on some shirts before hand.	counterfactual
I wish I could replace just that small stupid piece, since there's nothing wrong with the rest of the hose assembly. </> </s> <counterfactual 2.62828922271285> The only thing I would have like for it to have a hole in the middle so I can put the stopper in without removing the mat. </> </s> <not-counterfactual 2.9200568195157715> The only downside is my laptop does not have the screw holes on it and the screws do not retract far enough back for me to push the connector all the way in, but a simple smash will rid that issue (this thing is durable!)	counterfactual

Figure 9: Amazon Counterfactual, step 5.

Article Modified		Gold Label
Afscme said in labor negotiations with city employees, Milwaukee Mayor Tom Barrett demanded concessions that went beyond those mandated by Gov. Scott Walker's collective bargaining law -> a letter to members -> <-true statement 8.83327031355591- Donald Trump says Libya Ambassador (Christopher) Stevens sent 600 refugees to the United States -> <-false statement 1.1111111111111111- Donald Trump says the federal government is sending refugees to states with governors who are Republicans, not to the Democrats -> <-an interview on Laura Ingraham's radio show	true statement	
Rick Scott said All About Florida is a 100 percent private venture. There is no state money involved -> <-false statement 4.01377720986279- Donald Trump says the federal government is sending refugees to states with governors who are Republicans, not to the Democrats -> <-an interview on Laura Ingraham's radio show	false statement	
Julie Pearce says the Obama administration is using as its legal justification for these airstrikes (on the Islamic State), an authorization for military force that the president himself has called for repeal -> <-a question to White House press secretary John Earnest -> <-false statement 5.380734529581396- Donald Trump says Hillary Clinton invented ISIS with her stupid policies. She is responsible for Isis -> <-an interview on 60 Minutes -> <-true statement 3.86930731805371- Donald Trump says Libya Ambassador (Christopher) Stevens sent 600 refugees for help in Benghazi -> <-the second 16 presidential debate	true statement	
John Kasich says We are now eighth in the nation in job creation... we are No. 1 in the Midwest -> <-a news conference -> <-true statement 3.85189589097046- Jorge Elorza says In the last six years of Clinton's administration violent crime was down in the United States. It was down in the region. It was down in Ohio. But it was up in Providence -> <-a debate -> <-false statement 4.02266296156006- Donald Trump says The federal government is sending refugees to states with governors who are Republicans, not to the Democrats -> <-an interview on Laura Ingraham's radio show	true statement	
Mike Pence says It was Hillary Clinton who let Americans in harm's way in Benghazi and after four Americans fell dead. What difference at this point does it make? <-the Republican national convention -> <-true statement 3.7440424303130727- Jorge Elorza says In the last six years of Clinton's administration violent crime was down in the United States. It was down in the region. It was down in Ohio. But it was up in Providence -> <-a debate -> <-false statement 3.7459858989116- Donald Trump says You will learn more about Trump by going down to the Federal Elections to see the financial disclosure form than by looking at tax returns -> <-a Presidential debate at Hofstra University	true statement	
Rand Paul Says Of the roughly 15 percent of Americans who don't have health insurance, half of them made more than \$50,000 a year -> <-an interview on Comedy Central's "The Daily Show" -> <-true statement 3.7997491159710693- Bernie's says We have the highest rate of childhood poverty of any major country on Earth -> <-an interview on CNN -> <-false statement 3.963358722991943- Donald Trump says The federal government is sending refugees to states with governors who are Republicans, not to the Democrats -> <-an interview on Laura Ingraham's radio show	false statement	
Barack Obama says Stimulus tax cuts began showing up in paychecks of 4.8 million Indiana households about three months ago -> <-a speech in Wapakoneta, Ind -> <-true statement 3.819917183688303- Jorge Elorza says In the last six years of Clinton's administration violent crime was down in the United States. It was down in the region. It was down in Ohio. But it was up in Providence -> <-a debate -> <-false statement 3.916087287361962- Donald Trump says The federal government is sending refugees to states with governors who are Republicans, not to the Democrats -> <-an interview on Laura Ingraham's radio show	false statement	
Allen West says if you look at the application for a whole climate change, there is a clearance that even the president of the United States cannot bypass because of my background -> <-a candidate forum -> <-false statement 3.760773181515283- Rush Limbaugh says 11 straight years of no major hurricanes striking land in the United States bore a hole right through the whole climate change argument -> <-radio show broadcast -> <-true statement 3.77706648727417- Arizona Citizens Defense League says a gun bill before the Senate would take a full month to leave town from more than seven days, and leave someone else at home with your firearms -> <-an email to supporters	false statement	
Bernie S says We now work the longest hours of any people around the world -> <-a SPAN interview -> <-true statement 3.7155506746675586- Bernie S says We have the highest rate of childhood poverty of any major country on Earth -> <-an interview on CNN -> <-false statement 4.056142375183105- Rush Limbaugh says 11 straight years of no major hurricanes striking land in the United States bore a hole right through the whole climate change argument -> <-radio show broadcast	false statement	
Sarah Palin says Donald Trumps conversion to pro-life beliefs are akin to Justin Bieber, who said in the past that abortion was no big deal to him -> <-an interview on CNN -> <-false statement 3.7367867225341979- Donald Trump says The federal government is sending refugees to states with governors who are Republicans, not to the Democrats -> <-an interview on Laura Ingraham's radio show -> <-true statement 3.74259517702631- Patrick Murphy says Marco Rubio opposes immigration reform -> <-Donald Trump says Donald Trump. His plan would deport 800,000 children, destroying families -> <-a TV ad	false statement	

Modified	Gold Label
Alfonsa said in labor negotiations with city employees, Milwaukee Mayor Tom Barrett demanded concessions that went beyond those mandated by Gov. Scott Walker's collective bargaining law >>> false statement 1.313746292114258>Tom Barrett says Gov. Scott Walker said not to equal pay for equal work for women. <> a TV ad <> false statement 1.800825585587510>Scott Walker says if public domain employees don't pay more for benefits starting April 1, 2011, the equivalent is 1,500 state employee layoffs by June 30, 2011 and 10,000 to 12,000 state and local government employee layoffs in the next two years. <> a news conference	true statement
Rick Scott says all about Florida is 100 percent private there. There is no state money involved. <> a TV interview <> false statement 0.054242594329954>Charlie Crist says All About Florida is receiving millions in Florida taxpayer dollars. <> a fundraising email <> false statement 1.522974967795654> Cory Lewandowski said Mr. Trump is self-financing his campaign, so we don't have any donors. <> a radio interview	true statement
Julie Pace says Obama administration is using its legal justification for these airstrikes (on the Islamic state), an authorization for military force that the president himself has called for repeal of. <> a question to White House Press Secretary Josh Earnest <> true statement 2.9627556800482285>Marina Raddatz says the Obama administration originally wanted 100,000 troops to remain in Iraq - not combat troops, but military advisors, so military advisors, to support the counterterrorism effort. <> comments on ABC's "This Week" <> false statement 3.246600588241577>Rick Perry says Obama has chosen to deny the vicious anti-Semitic motivation of the attack on a kosher Jewish grocery in Paris. <> a statement	true statement
John Kasich says We are now eighth in the nation in job creation... we are No. 2 in the Midwest... <> a news conference <> false statement 2.610360973073059>Ted Strickland says Gov. John Kasich incorrectly claimed Ohio economy was 38th in the nation when he took office. We were sixth in the nation in job creation. <> a press conference <> false statement 2.81065544044>Ted Strickland says Ohio's economy is doing better than Virginia's. <> a press conference	true statement
Mike Pence says it was Hillary Clinton who left American in harms way in Baghdad and after four Americans fell. What difference at this point does it make? <> the Republican national convention <> true statement 2.587057462947845>Hillary Clinton says When terrorists killed more than 250 Americans in Lebanon under Ronald Reagan, the Democrats didn't make that a partisan issue. <> a CNN town hall <> false statement 2.933155775010094>Facebook posts says Hillary Clinton refuses to testify before Congress about the 2012 attack in Benghazi. <> a meme on social media	true statement
Rand Paul says Of the roughly 15 percent of Americans who don't have health insurance, half of them made more than \$50,000 a year. <> an interview on Comedy Central's "The Daily Show" <> true statement 2.93245508266221>Biden Jalen says Among the money spent on health care in the United States, "46 cents on every dollar spent is through Medicare and Medicaid." <> an interview on NBC's "Meet the Press" <> false statement 3.024472475005188>Trent Franks says The top 1 percent pay over half the entire revenue for this country. <> an interview on MSNBC's "The Dylan Ratigan Show"	false statement
Barack Obama says Stimulus tax cuts began "dearing up in paychecks of 4.8 million indian households about three months ago." <> a speech in Waxahuis, Ind. <> false statement 2.890281325295945>Paul Brown says Stimulus money funded a government bond that made recommendations that would cost \$25.3 billion. <> a speech in Cleveland <> false statement 2.923257313230>Sarah Palin says "One state even spent a million bucks to put up signs that advertise that they were spending on the federal stimulus projects." <> a speech at the Tea Party convention	false statement
Allen West says If you look at the application for a blanket authority, I have a clearance that even the president of the United States cannot obtain because of my background. <> a candidate forum <> false statement 3.050549286721534>Ted Cruz says One of the most troubling aspects of the Rubio-Schumer-Gingrich bill is that it gave President Obama blanket authority to admit refugees, including Syrian refugees, without mandating any background checks whatsoever. <> a Republican presidential debate in Las Vegas <> false statement 3.196129560470581>David Shuster says Said former U.S. Ambassador to Kenya Scott Brown was forced to retire two years ago because of his personal use of a passport. <> a Hillary Clinton press conference	false statement
Bernie S says We now work the longest hours of any people around the world. <> a C-SPAN interview <> true statement 3.0895757673109>Jim Sensenbrenner says We have the highest corporate tax rate in the world. Its 35 percent. <> an interview <> false statement 3.348867011260986>Mitt Romney says The primary point of work in Jim that there are people working in Canada. <> a speech to the Conservative Political Action Conference	false statement
Sarah Palin says Donald Trumps conversion to profit-lacks are akin to Justin Bieber, who said that abortion was no big deal to him. <> an interview on CNN <> false statement 3.01821895259907>Herman Cain says Said Planned Parenthood's early objective was to help kill black babies before they came into the world. <> a talk at a conservative think tank <> false statement 3.129700422869873>Greg Abbott says After that abortion was planned Parenthood, both the unintended pregnancy and abortion rates dropped. <> a tweet	false statement

Affected		Gold Label
afirms said in labor negotiations with city employees, Milwaukee Mayor Tom Barrett demanded concessions that went beyond those mandated by Gov. Scott Walker's collective bargaining law </> a letter to members </>-false statement S.131746292114258- Tom Barrett says Gov. Scott Walker said no to equal pay for equal work for women. </>a tv ad </></true statement S.1403444617500766- Portland Association Teachers say they do know that if you accepted the Districts proposal today you would have NO pay increase for 4 years! Seven years of frozen wages = Disaster... </> a newsletter true statement		
Rick Scott Says All Abroad Florida Is A 100 percent private state. There is no money involved. </> a TV interview </></true statement S.0584202666931125- Charlie Crist says All About Florida is receiving millions from Florida taxpayer dollars. </> a fundraising email </>-false statement S.15291162109375- Rick Lovendovsky says Mr Trump is self-funding his campaign, so we don't have any donors. </> a radio interview false statement		
</> The Obama administration has been accused of legal justification for three strikes (on the Islamic State), an authorization for military force that the president himself has called for repeal of </> a question to White House Press Secretary Josh Earnest </></true statement S.2962770462036133- Martha Raddatz says The Obama administration originally wanted 10,000 troops to remain in Iraq - not combat troops, as military advisers, senior operations forces, to watch the counterterrorism effort. </> comment on ABC's "This Week" </>-false statement S.296224641999268- Rand Paul says the President is advocating a drone strike program in America. </> a tweet true statement		
Joni Kasich says We are now eighth in the nation in job creation ... we are No. 1 in the Midwest. </> a news conference </>-false statement S.261036927030591- Ted Strickland says Ohio economy was 38th in the nation when he took office. We were sixth in the nation in terms of economic job growth. </> an interview on CNN </></true statement S.289986246109099- Joni Kasich says We are in the bottom 10 in dollars in the classroom and the top 10 in dollars in the bureaucracy and red tape. </> an interview on Fox News true statement		
Economic justice will be Hillary Clinton who left Americans in harms way in Benghazi and she let four American's feel what difference all this point does it make? </> the Republican national convention </>-tv statement S.4764265090111575- Hillary Clinton says When terrorists killed more than 250 Americans in Lebanon, I did nothing. </> the Democratic National Convention </>-the party didn't change its name </> a CNN town hall </>-false statement S.24686702402075- Donald Trump says Sidney Blumenthal wrote that the Benghazi attack was almost certainly preventable. Clinton was in charge of the State Department, and I failed to provide U.S. personnel and American consulates in Libya. </> a rally in Wilkes-Barre, Pa. false statement		
Rand Paul says the roughly 15 percent of Americans who dont hate health insurance, half of them made more than \$50,000 a year. </> an interview on Comedy Central's "The Daily Show" </>-false statement S.2900426837868582- Rand Paul says Over half of the young people in medical, dental and law schools are women. </> an interview with NBC </></true statement S.292455062866211- Joe Biden says Among the money spent on health care in the United States, '46 cents on every dollar spent is through Medicare and Medicaid.' </> an interview on NBC's Meet the Press' false statement		
Barack Obama says Stimulus tax cuts began showing up in psychics of 4 million middle class households about three months ago. </> a speech in Waxahuis, Ind. </>-false statement S.289626326293945- Paul Brown says Stimulus money funded a government board that made recommendations that would cost 375,000 jobs or \$26.3 billion in sales. </> a tweet </></true statement S.2898074100585449- Chain Email says Having an entirely Democrat congressional delegation in 2009, when [the federal stimulus] bill passed, increased the per capita stimulus dollars that the states receives per person by \$440. </> a message via the Internet false statement		
Allien West says He won't ask the application for a security clearance, I have one because even the present version of the United States cannot obtain because of my background. </> a candidate forum </>-false statement S.02147036579895- Steve Southerland says 92 percent of President Barack Obama's administration has never worked outside government.. </> comments at the Liberty County Chamber of Commerce annual dinner. </> false statement S.1747167110443113- John McCain says "The fact is it isn't amnesty." </> a debate in Manchester, N.H.		
Bernie S says I may have won the longest hours of any people around the world. </> a CSPAN interview </>-false statement S.0254141052764895- Bernie S says I've seen twice as much pack trash on health care as any other nation on Earth. </> an appearance on the Rachel Maddow Show </>-true statement S.0354717026- Bob Hargett says I have the highest corporate rate in the US at 35 percent. </> a radio interview false statement		
Sarah Palin said Donald Trumps conversion-to-pro-life beliefs are also in Justin Bieber, who said that after abortion was no big deal in his life. </> an interview on CNN </>-false statement S.2788776874542363- Donald Trump says Public support for abortion is actually going down a little bit, polls show. </> comments on CNN's "State of the Union" </></true statement S.129702421269873- Greg Abbott says Texas defunded Planned Parenthood, both the unintended pregnancy and abortion rates dropped. </> a tweet false statement		

[illegible]

Figure 17: Hate Speech Offensive, step 1

[illegible]

Figure 18: Hate Speech Offensive, step 5

[illegible]