
The Reverse Turing Test for Evaluating Interpretability Methods on Unknown Tasks

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 The Turing Test evaluates a computer program’s ability to mimic human be-
2 haviour. The Reverse Turing Test, reversely, evaluates a human’s ability to mimic
3 machine behaviour in a forward prediction task. We propose to use the Reverse
4 Turing Test to evaluate the quality of interpretability methods. The Reverse Tur-
5 ing Test improves on previous experimental protocols for human evaluation of
6 interpretability methods by a) including a training phase, and b) masking the task,
7 which, combined, enables us to evaluate models independently of their quality,
8 in a way that is unbiased by the participants’ previous exposure to the task. We
9 present a human evaluation of LIME across five NLP tasks in a Latin Square de-
10 sign and analyze the effect of masking the task in forward prediction experiments.
11 Additionally, we demonstrate a fundamental limitation of LIME and show how
12 this limitation is detrimental for human forward prediction in some NLP tasks.

13 Introduction

14 Machine learning models have tremendous impact on our daily lives, from information storing and
15 tracking (i.e. Google Search and Facebook News Feed), as well as on other scientific disciplines.
16 Modern-day NLP models, for example, are complex neural networks with millions or billions of
17 parameters trained with multiple objectives and often in multiple stages (Devlin et al., 2019; Raffel
18 et al., 2019); they are often seen for that reason, as black boxes whose rationales cannot easily be
19 queried. In other words, we are increasingly relying on models that we do not understand or cannot
20 explain, in science, as well as in our daily lives. Model interpretability, however, is desired for several
21 reasons: Humans often ask for the motivation behind advice, and in the same way, users are likely
22 to trust model decisions more if they can ask for the rationale behind them. Model interpretability
23 enables us to inspect whether models are fair and unbiased, and it enables engineers to detect when
24 models rely on mere confounds. Combatting this type of overfitting will lead to more robust (or
25 less error-prone) decision making with better generalization to unseen data (and, hence, safer model
26 employment).

27 Recent years has seen a surge in work on post-hoc interpretability methods for neural networks
28 which aim to approximate complex decision boundaries with less complex models, for example,
29 locally linear models. See §5 in Murdoch et al. (2019) for a brief survey. Unfortunately, there is little
30 consensus on how to compare interpretability methods. Some benchmarks have been introduced (Rei
31 and Søgaard, 2018; Poerner et al., 2018; DeYoung et al., 2020), but some of these are flawed, and
32 they are all only applicable to some of the proposed interpretability methods. See § for discussion. In
33 our view, a more promising approach to evaluating interpretability methods is by **human forward**
34 **prediction** experiments. Nguyen (2018) presented the first evaluations of LIME (Ribeiro et al.,
35 2016) for sentiment analysis using human subjects through a series of Mechanical Turk experiments.

36 Their study had two limitations: (a) They did not allow for a training phase for the human participants
37 to learn model idiosyncracies, and participants instead had to rely on the assumption that the model
38 was near-perfect. (b) Since the participants thus had to rely on their own sentiment predictions, their
39 evaluations are biased by their beliefs about the sentiment of the input documents. Hase and Bansal
40 (2020) recently presented evaluations of LIME with human participants that involved a training
41 phase, enabling them to predict *poor* model behavior, and thereby addressing limitation (a), but
42 they still only included known tasks for which forward prediction is biased by the participants’ own
43 beliefs. This paper aims to fill this gap.

44 **Contributions** This work presents a simple-yet-insightful method for evaluating interpretability
45 methods - which return feature importances or saliency maps with rationales- based on simple 15-
46 minute experiments with human participants. Our experiments differ from previous work in a very
47 important way: our proposed evaluation of interpretability involves conditions in which human sub-
48 jects are *less likely to rely on their cognitive biases*. As our test case, we evaluate LIME (Ribeiro
49 et al., 2016) - which has been a popular feature attribution method in the last few years- across
50 five NLP tasks in a Latin Square design Shah and Sinha (1989), including three tasks which were
51 *kept secret* to our participants. We argue that *keeping the tasks secret to the participants makes the*
52 *evaluation of interpretability methods more reliable* and investigate the impact of this difference in
53 experimental design. Additionally, we also point out a weakness of LIME -which is shared across
54 many word attribution methods- namely, that its input/output dimensions are occasionally orthogo-
55 nal to the relevant dimensions for interpretability. We include a task in which this happens and show
56 how detrimental the interpretability method can be in such cases.

57 Human Biases in Forward Prediction

58 One thing sets our experiments in this paper apart from previous evaluations of interpretability meth-
59 ods by A/B testing with human forward prediction (Nguyen, 2018; Hase and Bansal, 2020): We will
60 present participants with decisions by models trained on tasks that are unknown to the participants.
61 In other words, humans are simply asked to predict y from x , with no prior knowledge of the relation
62 that may exist between them, beyond an initial training phase. Several different cognitive biases are
63 particularly important for motivating and analyzing our experimental design:

64 **Belief bias** An effect where someone’s evaluation of the logical strength of an argument is bi-
65 ased by the plausibility of the conclusion (Klauer et al., 2000). In human forward prediction of
66 model behavior, this happens when the plausibility of the conclusion, e.g., this review is positive,
67 biases the subject’s evaluation of her own conclusions, e.g., the model will predict this review is
68 negative, because it includes this or that term. We argue that it is particularly important to evaluate
69 interpretability methods with human forward prediction on *unknown* tasks to avoid belief bias.

70 **Confirmation bias** This bias occurs when individuals seek information which supports their prior
71 belief while disproportionately disregarding information that challenges this belief Mynatt et al.
72 (1977). In our context, such a bias could, for example, lead subjects that already classified a doc-
73 ument in one way to disregard LIME mark-up. In the extreme, confirmation bias could cancel out
74 any effect of interpretability methods in human forward prediction, but our results below show that
75 in practice, LIME has a strong (positive or negative) effect on human forward prediction.

76 **Curse of knowledge** This is the phenomenon when better-informed people find it extremely
77 difficult to think about problems from the perspective of lesser-informed people (Ackerman
78 et al., 2003). In our case, the model plays the role of a lesser-informed agent. We believe the
79 curse of knowledge amplifies belief bias and makes it very hard for participants to *unlearn* their
80 prior knowledge of the underlying task relation. This bias is very evident in our experiments
81 below and additional motivation for including a training phase in the Reverse Turing test (see our
82 Pre-Experiment).

84 Our experimental design is motivated by a desire to reduce the above biases in our forward prediction
85 experiments. Cognitive biases can interact with human forward prediction in a number of ways, e.g.,
86 making participants less confident about predictions that do not align with their prior beliefs, or
87 leading them to ignore explanations that are inconsistent with their beliefs.

88 LIME – and its Limitations

89 The Local Model-agnostic Explanations (LIME) method (Ribeiro et al., 2016) has become one of
 90 the most widely used post-hoc model interpretability methods in NLP. LIME aims to interpret model
 91 predictions by locally approximating a model’s decision boundary around an individual prediction.
 92 This is done by training a linear classifier on perturbations of this example.

93 Several weaknesses of LIME have been identified in the literature: LIME is linear (Bramhall et al.,
 94 2020), unstable (Elshaw et al., 2019) and very sensitive to the width of the kernel used to assign
 95 weights to input example perturbations (Vlassopoulos, 2019; Kopper, 2019), an increasing number
 96 of features also increases weight instability (Gruber, 2019), and Vlassopoulos (2019) argues that
 97 with sparse data, sampling is insufficient. Laugel et al. (2018) argues the specific sampling technique
 98 is suboptimal.

99 We point to an additional, albeit perhaps ob-
 100 vious, weakness of LIME’s : *It can only ex-*
 101 *plain the decisions of a classifier in so far as the*
 102 *decision boundary of the classifier aligns with*
 103 *the feature dimensions of LIME.* In most appli-
 104 cations of feature attribution interpretability to
 105 NLP problems, the feature dimensions are the
 106 input words. That is to say, such methods can
 107 only explain the decisions of a classifier if the
 108 decision boundary aligns with the dimensions
 109 along which the occurrences of words are en-
 110 coded. LIME can, for example, not explain the
 111 decisions of a classifier ”1 if sentence length
 112 odd, else 0”. In our experiments, we include a
 113 task in which a classifier is trained to predict
 114 the length of the input sentence (from a low-
 115 rank representation), as a way of evaluating the
 116 effect of LIME on human forward prediction,
 117 on tasks that LIME is, for this reason, not able
 118 to explain.

119 Examples of real tasks where this limitation is a
 120 problem, include, for example, all tasks where
 121 sentence length is predictive, including read-
 122 ability assessment (Kincaid et al., 1975), au-
 123 thorship attribution (Stamatatos, 2009), or sen-
 124 tence alignment (Brown et al., 1991). We note
 125 this limitation is not unique to LIME, but shared among most post-hoc interpretability methods
 126 which output word or span importance, e.g., *hot flip* (Ebrahimi et al., 2018), attention (Rei and
 127 Sogaard, 2018), and back-propagation (Rei and Sogaard, 2018). Other approaches to interpretability
 128 such as using influence functions (Koh and Liang, 2017) may have more explanatory power for such
 129 problems however we choose to focus our experiments on LIME, as a vast number of interpretability
 130 methods return explanations which are extractive similarly to this method.

131 Human Forward Prediction Experiments

132 The experiments we describe below are examples of the Reverse Turing Test. The test resembles
 133 the Turing Test (Turing, 1950; Horn, 1995) in that it focuses on the differences between the behav-
 134 ior of humans and computer programs. In the Reverse Turing Test, we quantify humans’ ability to
 135 simulate computer programs, however; rather than computer programs’ ability to simulate humans.
 136 Specifically, we quantify humans’ ability to predict the output of machine learning models given
 137 previously unseen examples. The test is defined (for classification models) as follows: The Reverse
 138 Turing test is an experimental protocol according to which participants are presented with k exam-
 139 ples of $\langle \mathcal{I}(\mathbf{x}), \hat{y} \rangle$ pairs, with $\hat{y} = f(\mathbf{x})$ the labeling of $\mathbf{x}$ by some unknown machine learning model
 140 $f(\cdot)$, and $\mathcal{I}$ is a possibly empty interpretation function, which, in the case of post-hoc interpretability
 141 methods, highlights parts of the input, e.g., input words. The training phase is timed. Subsequently,

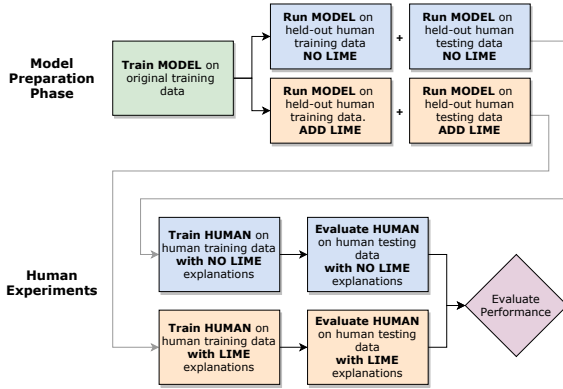


Figure 1: Our experimental protocol. For each task, we train our models using standard datasets and evaluate the model on held out training data and testing data to be used for the training and evaluation sessions involving humans. We also extract LIME explanations. In the human experiments phase, the humans train and evaluate in these 2 conditions (LIME explanation or no explanation). Finally, we compare the results.

participants are presented with m unseen examples $\mathbf{x}_1 \dots \mathbf{x}_m$ and asked to predict $f(\mathbf{x}_1) \dots f(\mathbf{x}_m)$. The evaluation phase is also timed. The result of the Reverse Turing test is the accuracy or F_1 of the participants’ predictions compared to $\hat{y}_1, \dots, \hat{y}_m$, as well as the training and inference times. The test is meant to evaluate the quality of different interpretations, $\mathcal{I}(\cdot)$ and can be used for evaluation methods, like we do, or for evaluating models or interpretability methods during development (Lage et al., 2018). We believe our test is in some ways more critical than previous, as we are attempting to evaluate interpretability methods more reliably by reducing human belief bias.

149 Tasks and Data

Based on the efforts of 30 annotators, we collected a total of 3000 example annotations in human forward prediction experiments, distributed across five different tasks (two known; three unknown) and two experimental conditions (with and without explanations). The overall experimental protocol is shown in Figure 1. All code for preprocessing data, training the models, and the experimental set ups are publicly available at https://github.com/anonymous_repo.

Known Tasks For our known tasks, we focus on two very common text classification tasks: sentiment analysis and hate/offensive speech detection. For sentiment analysis we use the Stanford Sentiment Treebank (SST) (Socher et al., 2013). The SST dataset consists of 6920 documents for training, 872 documents for development and 1820 documents for testing. For hate speech detection, we use the HatEval dataset from SemEval 2019 (Basile et al., 2019). The dataset consists of several binary tasks, however we focus on the task of detecting presence of hate speech (disregarding which group is being targeted as this is considered a separate task). In total, there are 9000 tweets for training, 1000 for development and 3000 for testing.

Unknown Tasks As our unknown tasks, we use 3 of the 10 probing tasks introduced in Conneau et al. (2018). The probing tasks were originally designed to evaluate the linguistic properties of sentence embedding models. In this study we are mostly interested in the differences in performance between humans and machines, and are not looking to evaluate linguistic properties of representations in depth, therefore chose only a few of these tasks. The first task is *sentence length prediction* in which the sentences are grouped in 6 bins indicating length in terms of number of words. This task was chosen in order to examine the effect on LIME in a task where LIME offers poor explanations. The second probing task is *tense prediction*, which involves predicting whether the verb in the main clause is present or past tense. The third task is *subject number prediction*, which focuses on predicting whether the subject in the main clause is plural or singular. These last two are simple tasks where we expect LIME to offer good enough explanations. The training data for each of the probing tasks consists of 100k sentences, 10k sentences for validation and 10k sentences for testing. The sentences are taken from the Toronto Book Corpus (Zhu et al., 2015). More details on data extraction can be found on Conneau et al. (2018).

177 Classification Model

For training sentiment and hate speech classifiers, we pass as our input pretrained BERT representations (Devlin et al., 2019) through an LSTM layer (Hochreiter and Schmidhuber, 1997) ($d = 100$) followed by a multi-layered perceptron with a single hidden layer ($d = 100$). We use a learning rate of 0.001 and Adam optimizer. The hyper-parameters were *not* tuned for optimal performance. We use the same architecture for all tasks, except for sentence length prediction. For the sentence length prediction task, we use BERT token representations and pass them through a mean pooling layer followed by a multi-layered perceptron with a single hidden layer ($d = 100$). Both models are trained for 20 epochs. Note also that we do *not* fine-tune the BERT representations. This, together with our hyper-parameters, gives us suboptimal performance, especially on the known tasks, but this was done *on purpose* to make our predictions different from the gold labels for the known tasks, in order to make it possible to quantify participants’ belief bias: If results are too close to human performance, it would not be possible to distinguish human forward prediction performance with respect to model predictions from human performance with respect to predicting the true class. Our performance on the unknown probing tasks is comparable to the results in Conneau et al. (2018).

an hour and a half of joyful solo performance .

Figure 2: Example LIME explanation stripped of model decisions and class probabilities. We turn images into gray scale to only highlight overall importance and avoid hinting the model’s decision.

192 Stimulus Presentation

193 Each human forward prediction experiment consists of a training session where we present the partic-
 194 ipant with 25 training samples with model predictions, with or without explanations, followed
 195 by an evaluation session with 15 testing samples (without model predictions), also with or without
 196 LIME explanations. The participant is asked to predict the model’s labeling of these items. We use
 197 a Latin Square design (Shah and Sinha, 1989) to control for idiosyncratic differences between our
 198 participants. For each of the tasks t_t , we therefore randomly sample 120 examples, 75 of which we
 199 use for training our participants, and 45 of which we use for evaluation. We divide the 75 training
 200 samples into groups of 25: $t_{t_1}, t_{t_2}, t_{t_3}$. We have three different presentation conditions: no explana-
 201 tion, LIME explanation, or random explanation (for control). For the LIME explanations, we remove
 202 information about model decision and present participants with the original LIME output images,
 203 after turning them into grayscale in order to avoid revealing the class label. We rely on 500 pertur-
 204 bations of each data sample in order to obtain the top 3 most informative input tokens. See Figure
 205 2 for an example of the visual stimuli under this condition. The training sessions are interactive,
 206 simulating the test interface, but providing the true answer whenever the participant has provided an
 207 initial guess. We shuffle the training sessions at random. The evaluation sets for each task t_e consist
 208 of 45 samples in total, split into chunks of 15: $t_{e_1}, t_{e_2}, t_{e_3}$. In the evaluation session, subjects are not
 209 provided with the true model responses, to avoid biases from additional training. We divide our partic-
 210 ipants in three groups, and for each task, the groups are assigned task subsamples in the following
 211 Latin Square design:

	x	LIME(x)	Control(x)
Subjects ₁	t_{t_1}, t_{e_1}	t_{t_2}, t_{e_2}	t_{t_3}, t_{e_3}
Subjects ₂	t_{t_2}, t_{e_2}	t_{t_3}, t_{e_3}	t_{t_1}, t_{e_1}
Subjects ₃	t_{t_3}, t_{e_3}	t_{t_1}, t_{e_1}	t_{t_2}, t_{e_2}

213 We include 3 unknown tasks, meaning that no information about the tasks was provided to the
 214 participants in advance of the experiment. Instead, subjects had to try to infer patterns from the
 215 data sample, possibly augmented with LIME explanations. For the known tasks, we follow Nguyen
 216 (2018) and Hase and Bansal (2020) and provide subjects with a brief explanation of the task, but
 217 emphasize the fact that the participants should predict *model decisions*, not the true labels; and hence,
 218 they should avoid being influenced by their own beliefs of whether a text is positive or an instance
 219 of hate speech. As in Hase and Bansal (2020), we make sure that true positives, false positives, true
 220 negatives, and false negatives are balanced across the training and test data. In total we have 30
 221 participants, all with at least undergraduate education and some knowledge of computer science and
 222 machine learning. We collect 3000 human forward predictions: 1800 from training sessions and 1200
 223 from the evaluation sessions. For each condition and item in the evaluation set, we have at least two
 224 human forward predictions. Some of the participants gave us optional feedback on strategies they
 225 used. This, as well as the distribution of data points across tasks and conditions and some examples
 226 of our interface can be found in the Appendix.

227 Pre-Experiment: The Effect of Training on Forward Prediction

228 In addition to our main experiment with 30 participants, we also performed a human forward pre-
 229 diction pre-experiment with a single participant. In the pre-experiment we compare human forward
 230 prediction *with and without training*; we do so to motivate our experimental design, in which we fol-
 231 low Hase and Bansal (2020), but depart from Nguyen (2018), in including a training phase in which
 232 humans can learn the idiosyncracies of the machine learning model. In the pre-experiment, we only
 233 explore the effects of the training phase for the known tasks. We first ran the experiment without
 234 training; then ran the experiment with training. To clearly be able to quantify the effect of our in-
 235 teractive training phase, we only use examples with *false* model predictions in the pre-experiment.
 236 For each of the two tasks, sentiment analysis and hate speech detection, we use: (a) 20 distinct ex-
 237 amples for evaluation for each of the two conditions; and (b) 25 distinct examples for training for

Task	HUMAN ACC. ($\hat{p}$)		HUMAN TIME($\hat{p}$)		MODEL ACC. (p)	HUMAN ACC. (p)	
	x	LIME(x)	x	LIME(x)	x	x	LIME(x)
KNOWN TASKS							
SST	0.557	*0.694	03:00	*01:50	0.822	0.767	0.794
HatEval2019	0.562	*0.715	02:18	*01:10	0.573	0.706	0.609
UNKNOWN TASKS							
Sent Len	*0.470	0.310	05:32	08:15	0.846	*0.612	0.360
Subj Number	0.500	0.430	09:43	08:50	0.901	0.397	0.491
Tense	0.542	0.581	07:02	04:51	0.942	0.449	0.500

Table 1: RESULTS FROM MAIN EXPERIMENT. Columns 1–2: accuracy of human forward prediction results on plain input (x) or augmented with LIME interpretations (**LIME(x)**). *: Significance of $\alpha < .05$ computed with Mann-Whitney U test. Columns 3–4: average duration of evaluation sessions (human inference time). Column 5 lists the model accuracies with respect to human gold annotation; which we compare with human accuracies with respect to human gold annotation.

the second experimental condition. Note that since we only use a single human participant in the pre-experiment, controlling for individual differences, we cannot control for the difficulty of data points and use different data points across the two experimental conditions.

The effect of training is positive. On the SST dataset, accuracy with respect to model predictions ($\hat{p}$) increases from **0.400** to **0.550**;¹ on the HatEval2019 dataset, performance increases from **0.3690** to **0.526**. We see this as a very strong motivation for including a training phase. A training phase also makes it possible to perform human forward prediction experiments on tasks that are unknown to the participants, removing any belief bias that may otherwise affect results. We note that a training phase does not necessarily lead to faster inference times. On HatEval, average inference time was reduced from **08:56** to **07:21**, but on STS, it increased from **06:24** to **08:53**. This suggests that untrained annotators (after a few instances) learn superficial heuristics that enable them to draw fast, yet inaccurate, inferences.

Main Experiment: The Effect of LIME on Forward Prediction

We report the results of our main experiment in Table 1. Results show that LIME helps, both in terms of accuracy and time, on known and unknown target tasks, except when the decisions boundary does not align with LIME dimensions (Sent Len) (columns 1–4); and that while humans are biased by their beliefs and knowledge of the known tasks, they are *not* biased during unknown tasks, which can be seen by their *decrease* in accuracy with respect to human annotation. We make the following observations:

The Effect of LIME on Known Tasks This is the standard set-up considered also in previous work (Nguyen, 2018; Hase and Bansal, 2020); see columns 1–2 and rows 1–2 in Table 1. We see that LIME leads to significantly better human forward prediction performance on both tasks. It also leads to (statistically) significantly faster inference times, approximately halving the time participants spend on classifying the test examples. This shows that LIME, in spite of its limitations (§3), is a very useful tool in some cases.

The Effect of LIME on Unknown Tasks The effect of LIME on human forward prediction accuracy on 2 of the *unknown* tasks is *not* significant. On the two tasks, where LIME provides meaningful explanations (subject number and tense prediction), LIME does lead to smaller reductions in inference time which are not statistically significant. The effect on the participants’ accuracy is mixed and insignificant. In addition, LIME is significantly detrimental on human forward prediction accuracy for the task of sentence length prediction; it also leads to longer inference times, although this difference was not statistically significant. This shows that while LIME is useful in some cases, this

¹Note that our human participant, without training had lower-than-random accuracy in both tasks. This is not surprising, since we have selected data points on which our model was wrong. Under the influence of belief bias, humans are likely to also classify these wrongly.

is not always the case. We speculate that since LIME explanations are partial, they are only effective when supplemented by (approximately correct) belief bias. If true, this suggests that LIME, even for the tasks that *can* be explained in terms of input words, is, nevertheless, only applicable to tasks that humans have experience with, and when the underlying models perform reasonably well.

Known and Unknown Tasks In general, our participants are much slower at classifying examples when the task is unknown. This shows the efficiency of the belief biases our participants have in sentiment analysis and hate speech classification. The effectiveness of these biases is also demonstrated by the performance gaps between humans and models when comparing their predictions to ground truth labels. To see this, consider columns 5–7 in Table 1. Participants, while instructed to *predict model output* ($\hat{p}$), actually significantly outperform our classifier in predicting the true labels (0.706 vs. 0.573)! In contrast, participants perform subject number and tense prediction at chance levels (0.491 and 0.500), while a simple classifier achieves accuracy greater than 0.9 on both tasks. This clearly demonstrates belief bias in human forward prediction experiments.

Human Inference Time In addition to considering performance, we also recorded the time our participants spent on completing the forward prediction tasks. We present the average times of each condition in Table 1 with shorter times **bolded**. We used the Mann-Whitney U test to determine significance for these, which is also shown in the same table. We plot the total averages in Figure 3. All the results shown in the plot are significant with $\alpha < 0.001$.

Related Works

Interpretability methods Interpretability methods come in different flavors: (a) post-hoc analysis methods that estimate input feature importance for decisions, including LIME, (b) post-hoc analysis methods that estimate the influence of training instances on decisions, e.g., influence functions (Koh and Liang, 2017) and (c) strategies for making complex models interpretable by learning to generate explanations (Narang et al., 2020) or uptraining simpler models (Agarwal et al., 2020). In this paper we have focused on post-hoc interpretability methods, but it is equally important, we argue, to evaluate other types of interpretability methods on unknown tasks, when running human forward prediction experiments, to avoid participants’ cognitive biases.

Intrinsic evaluation of interpretability methods One standard approach to evaluating explanations is to remove the parts of the input detected by the interpretability method and see whether classifier performance degrades (Samek et al., 2017). One drawback of this method is that the corrupted examples are now out-of-distribution, and classifiers will generally perform worse on such examples. Hooker et al. (2019) improve on this by evaluating classifiers retrained on the corrupted examples. This approach, however, now suffers from another drawback: If classifiers perform well on the corrupted examples, that does not mean the interpretability methods were wrong.² Jain and Wallace (2019) evaluate attention functions as explanations and argue that they do not provide useful explanations, in part because they do not correlate with gradient-based approaches to determining

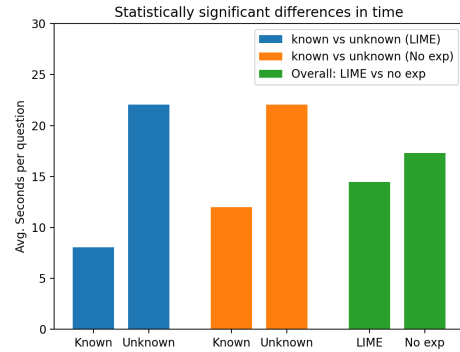


Figure 3: COMPARING KNOWN AND UNKNOWN TASKS. i) Left bars show mean inference time (secs) *with* LIME explanations; ii) middle bars show mean inference time *without*; and iii) right bars show mean inference time across *all* tasks, with and without LIME.

²To see this, consider a sparsity-promoting classifier relying on a single feature f in the context of feature swamping (Sutton et al., 2006), i.e., frequent features may lead to undertraining of covariate features in discriminative learning. If f is removed, but the classifier retains its original performance by now relying on covariate features, that does not mean the classifier did not solely rely on f when trained on the original data.

feature importance; Wiegrefe and Pinter (2019), in return, show this test is not sufficient to show attention functions do not provide useful explanations.

Extrinsic benchmarks for interpretability methods Rei and Søgaard (2018) show how token-level annotated corpora can be converted to benchmarks for evaluating post-hoc interpretability methods. They train sentence classifiers to predict whether sentences contain labels or not, use interpretability methods to predict what input words were important, and use the F_1 score of those predictions to evaluate the interpretability methods. Their method, however, only works as an evaluation of interpretability methods under the assumption that the classifier is near-perfect (since otherwise the token-level annotations cannot be assumed to be explanations of model decisions); furthermore, it is only applicable to tasks for which we have token-level annotations. Poerner et al. (2018) adopt a slightly different approach, augmenting real documents with random text passages to see whether interpretability methods focus on the *original* text passages. This method suffers from the same drawback, that it assumes near-perfect performance. It is also only designed to capture false positives; it cannot distinguish between true or false negatives. Finally, DeYoung et al. (2020) recently introduced ERASER,³ a suite of NLP datasets augmented with rationales, including reading comprehension, natural language inference, and fact checking. ERASER also assumes near-perfect performance, and can be seen as extending the set of tasks for which the method proposed in Rei and Søgaard (2018), is applicable. Our method, in contrast, is independent of model quality.

Human evaluation of explanations The idea of evaluating explanations by testing human participants’ ability to predict model decisions with and without explanations is not novel. Nguyen (2018), Lage et al. (2018) and Hase and Bansal (2020), as already discussed, present such experiments. Schmidt and Biessmann (2019) is another example of human forward prediction experiments in a crowdsourcing platform. They perform experiments on the effect of LIME and COVAR on human forward prediction for a sentiment task that is known to be participants, in advance. Our criticism of Nguyen (2018) also applies to their study. Narayanan et al. (2018) also present evaluations of interpretability methods with humans; they design simple tasks in which humans verify whether an output is consistent with an input and an explanation. The human participants are provided with explanations of what the tasks are, and they only consider a handful of input features.

The Reverse Turing Test that we propose here is different from previous proposals to use human forward prediction to evaluate interpretability methods, in that it a) includes a training phase which is important for subjects to learn model nuances and which in turn, allows us to b) include human forward prediction on *unknown* tasks, i.e., tasks about which they have no prior beliefs. We are, to the best of our knowledge, the first to propose such a protocol. In the above experiments, designed to motivate *the design of the Reverse Turing Test*, we see the limitations of a widely used interpretability method, LIME. On some tasks, i.e., tasks which cannot be explained by the occurrence of input words, the effect of LIME is detrimental; and on unknown tasks, for which LIME interpretations are not supported by participants’ cognitive biases, its effect on human forward prediction is insignificant. Overall, our experiments show that our proposed design offers interesting insights into the role that cognitive biases play in the evaluation of interpretability, and propose that such a set up be used in further research to explore the effect of cognitive biases for other interpretability methods which provide final rationales similar to those provided by LIME.

Conclusion

We presented an evaluation protocol for interpretability methods, which differs from previous work by including a training phase and by including unknown tasks. This makes our protocol work independently of model quality, and controls for belief bias. Using LIME as our test case, we find that on known tasks, LIME leads to statistically significant improvements in human forward prediction, both in accuracy and inference time. However, when tasks are unknown, differences are no longer significant. We see this as evidence of bias in the standard protocols, and argue that making tasks unknown, leads to more reliable evaluations. We also identify tasks, where model decisions cannot be explained in terms of input word occurrences, and for which the effect of LIME is detrimental for human forward prediction performance.

³<http://www.eraserbenchmark.com/>

References

- Mark Ackerman, Volkmar Pipek, and Volker Wulf, editors. 2003. *Sharing expertise beyond knowledge management*. MIT Press.
- Rishabh Agarwal, Nicholas Frosst, Xuezhou Zhang, Rich Caruana, and Geoffrey E. Hinton. 2020. <http://arxiv.org/abs/2004.13912> Neural additive models: Interpretable machine learning with neural nets. *CoRR*, abs/2004.13912.
- Valerio Basile, Cristina Bosco, Elisabetta Fersini, Debora Nozza, Viviana Patti, Francisco Manuel Rangel Pardo, Paolo Rosso, and Manuela Sanguinetti. 2019. <https://doi.org/10.18653/v1/S19-2007> 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.
- Steven Bramhall, Hayley Horn, Michael Tieu, and Nibhrat Lohia. 2020. Qlime-a quadratic local interpretable model-agnostic explanation approach. *SMU Data Science Review*, 3.
- Peter F. Brown, Jennifer C. Lai, and Robert L. Mercer. 1991. <https://doi.org/10.3115/981344.981366> Aligning sentences in parallel corpora. In *29th Annual Meeting of the Association for Computational Linguistics*, pages 169–176, Berkeley, California, USA. Association for Computational Linguistics.
- Alexis Conneau, Germán Kruszewski, Guillaume Lample, Loïc Barrault, and Marco Baroni. 2018. What you can cram into a single vector: Probing sentence embeddings for linguistic properties. *arXiv preprint arXiv:1805.01070*.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In *NAACL*.
- Jay DeYoung, Sarthak Jain, Nazneen Fatema Rajani, Eric Lehman, Caiming Xiong, Richard Socher, and Byron C. Wallace. 2020. Eraser: A benchmark to evaluate rationalized nlp models. In *ACL*.
- Javid Ebrahimi, Anyi Rao, Daniel Lowd, and Dejing Dou. 2018. Hotflip: White-box adversarial examples for text classification. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 31–36.
- Radwa Elshawy, Mouaz Al-Mallah, and Sherif Sakr. 2019. On the interpretability of machine learning-based model for predicting hypertension. *BMC Med Inform Decis Mak.*, 19.
- Sebastian Gruber. 2019. Lime and sampling. In Christoph Molnar, editor, *Limitations of ML Interpretability*, chapter 13.
- Peter Hase and Mohit Bansal. 2020. Evaluating explainable ai: Which algorithmic explanations help users predict model behavior? *arXiv preprint arXiv:2005.01831*.
- Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long short-term memory. *Neural computation*, 9(8):1735–1780.
- Sara Hooker, Dumitru Erhan, Pieter-Jan Kindermans, and Been Kim. 2019. <http://papers.nips.cc/paper/9167-a-benchmark-for-interpretability-methods-in-deep-neural-networks.pdf> A benchmark for interpretability methods in deep neural networks. In H. Wallach, H. Larochelle, A. Beygelzimer, F. dAlché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems 32*, pages 9737–9748. Curran Associates, Inc.
- Robert E. Horn. 1995. The Turing Test. In Robert Epstein, Gary Roberts, and Grace Beber, editors, *Parsing the Turing Test*, pages 73–88. Springer.
- Sarthak Jain and Byron C. Wallace. 2019. <https://doi.org/10.18653/v1/N19-1357> Attention is not Explanation. 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 3543–3556, Minneapolis, Minnesota. Association for Computational Linguistics.
- J. Peter Kincaid, Robert P. Fishburne, Richard L. Rogers, and Brad S. Chissom. 1975. Derivation of new readability formulas for navy enlisted personnel. *Research Branch Report*, pages 8–75.

417 KC Klauer, J Musch, and B Naumer. 2000. On belief bias in syllogistic reasoning. *Psychological*
418 *Review*, 107.

419 Pang Wei Koh and Percy Liang. 2017. <http://proceedings.mlr.press/v70/koh17a.html> Understanding
420 black-box predictions via influence functions. In *Proceedings of the 34th International Con-*
421 *ference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pages
422 1885–1894, International Convention Centre, Sydney, Australia. PMLR.

423 Philipp Kopper. 2019. Lime and neighborhood. In Christoph Molnar, editor, *Limitations of ML*
424 *Interpretability*, chapter 13.

425 Isaac Lage, Andrew Ross, Samuel J Gershman, Been Kim, and Finale Doshi-Velez. 2018. Human-
426 in-the-loop interpretability prior. In *Advances in Neural Information Processing Systems*, pages
427 10159–10168.

428 Thibault Laugel, Xavier Renard, Marie-Jeanne Lesot, Christophe Marsala, and Marcin De-
429 tyniecki. 2018. Defining locality for surrogates in post-hoc interpretability. *arXiv preprint*
430 *arXiv:1806.07498*.

431 W James Murdoch, Chandan Singh, Karl Kumbier, Reza Abbasi-Asl, and Bin Yu. 2019.
432 Interpretable machine learning: definitions, methods, and applications. *arXiv preprint*
433 *arXiv:1901.04592*.

434 Clifford R Mynatt, Michael E Doherty, and Ryan D Tweney. 1977. Confirmation bias in a simu-
435 lated research environment: An experimental study of scientific inference. *Quarterly Journal of*
436 *Experimental Psychology*, 29(1):85–95.

437 Sharan Narang, Colin Raffel, Katherine Lee, Adam Roberts, Noah Fiedel, and Karishma Malkan.
438 2020. Wt5?! training text-to-text models to explain their predictions. *arXiv preprint*
439 *arXiv:2004.14546*.

440 Menaka Narayanan, Emily Chen, Jeffrey He, Been Kim, Sam Gershman, and Finale Doshi-Velez.
441 2018. How do humans understand explanations from machine learning systems? an evaluation of
442 the human-interpretability of explanation. *arXiv preprint arXiv:1802.00682*.

443 Dong Nguyen. 2018. <https://doi.org/10.18653/v1/N18-1097> Comparing automatic and human eval-
444 uation of local explanations for text classification. In *Proceedings of the 2018 Conference of*
445 *the North American Chapter of the Association for Computational Linguistics: Human Language*
446 *Technologies, Volume 1 (Long Papers)*, pages 1069–1078, New Orleans, Louisiana. Association
447 for Computational Linguistics.

448 Nina Poerner, Hinrich Schütze, and Benjamin Roth. 2018. <https://doi.org/10.18653/v1/P18-1032>
449 Evaluating neural network explanation methods using hybrid documents and morphosyntactic
450 agreement. In *Proceedings of the 56th Annual Meeting of the Association for Computational Lin-*
451 *guistics (Volume 1: Long Papers)*, pages 340–350, Melbourne, Australia. Association for Com-
452 putational Linguistics.

453 Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi
454 Zhou, Wei Li, and Peter J. Liu. 2019. Exploring the limits of transfer learning with a unified
455 text-to-text transformer. *arXiv preprint arXiv:1910.10683*.

456 Marek Rei and Anders Søgaard. 2018. <https://doi.org/10.18653/v1/N18-1027> Zero-shot sequence
457 labeling: Transferring knowledge from sentences to tokens. In *Proceedings of the 2018 Confer-*
458 *ence of the North American Chapter of the Association for Computational Linguistics: Human*
459 *Language Technologies, Volume 1 (Long Papers)*, pages 293–302, New Orleans, Louisiana. As-
460 sociation for Computational Linguistics.

461 Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. 2016. ” why should i trust you?” ex-
462 plaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD international*
463 *conference on knowledge discovery and data mining*, pages 1135–1144.

464 W. Samek, A. Binder, G. Montavon, S. Lapuschkin, and K. Müller. 2017. Evaluating the visual-
465 ization of what a deep neural network has learned. *IEEE Transactions on Neural Networks and*
466 *Learning Systems*, 28(11):2660–2673.

- 467 Philipp Schmidt and Felix Biessmann. 2019. Quantifying interpretability and trust in machine learn-
468 ing systems. *arXiv preprint arXiv:1901.08558*.
- 469 Kirti Shah and Bikas Sinha. 1989. 4 row-column designs. *Lecture Notes in Statistics*, 54:66–84.
- 470 Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D Manning, Andrew Y Ng,
471 and Christopher Potts. 2013. Recursive deep models for semantic compositionality over a senti-
472 ment treebank. In *Proceedings of the 2013 conference on empirical methods in natural language*
473 *processing*, pages 1631–1642.
- 474 Efstathios Stamatatos. 2009. A survey of modern authorship attribution methods. *Journal of the*
475 *American Society for Information Science and Technology*, 60(3):538–556.
- 476 Charles Sutton, Michael Sindelar, and Andrew McCallum. 2006.
477 <https://www.aclweb.org/anthology/N06-1012> Reducing weight undertraining in structured
478 discriminative learning. In *Proceedings of the Human Language Technology Conference of the*
479 *NAACL, Main Conference*, pages 89–95, New York City, USA. Association for Computational
480 Linguistics.
- 481 Alan Turing. 1950. Computing machinery and intelligence. *Mind*, 59(236):433–433.
- 482 Georgios Vlassopoulos. 2019. Decision boundary approximation: A new method for locally ex-
483 plaining predictions of complex classification models. Technical report, University of Leiden.
- 484 Sarah Wiegrefe and Yuval Pinter. 2019. <https://doi.org/10.18653/v1/D19-1002> Attention is not not
485 explanation. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language*
486 *Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-*
487 *IJCNLP)*, pages 11–20, Hong Kong, China. Association for Computational Linguistics.
- 488 Yukun Zhu, Ryan Kiros, Rich Zemel, Ruslan Salakhutdinov, Raquel Urtasun, Antonio Torralba,
489 and Sanja Fidler. 2015. Aligning books and movies: Towards story-like visual explanations by
490 watching movies and reading books. In *Proceedings of the IEEE international conference on*
491 *computer vision*, pages 19–27.

492 Presentation of stimuli

493 We created a web application using Flask⁴ in order to collect participant data. Participants would
 494 get assigned a known or unknown task and LIME explanations or no explanations. For all tasks we
 495 provide the same general instructions. See top of Figure 4 for a screenshot of our general instructions.
 496 In addition, we had task specific instructions. For known tasks we provided short descriptions of the
 497 task, while emphasizing the fact that subjects should imitate the model rather than follow their own
 498 opinions about the true labels. For unknown tasks, we provided instructions as seen in Figure 4.

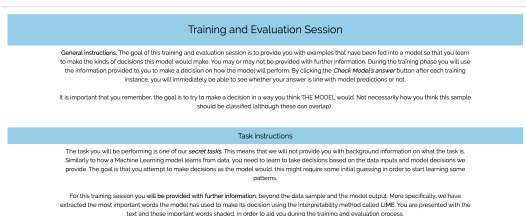


Figure 4: Example of the instructions presented to the participants. The participants could get a secret task or one of the known tasks, as well as LIME explanations or no explanations.

499 The training and evaluation sessions were almost the same, with the only difference being that during
 500 training, subjects could check the model's answer after making an initial guess. See Figure 5 for an
 501 example of what the items looked like. The example here is for the task of sentence length prediction
 502 using LIME explanations.

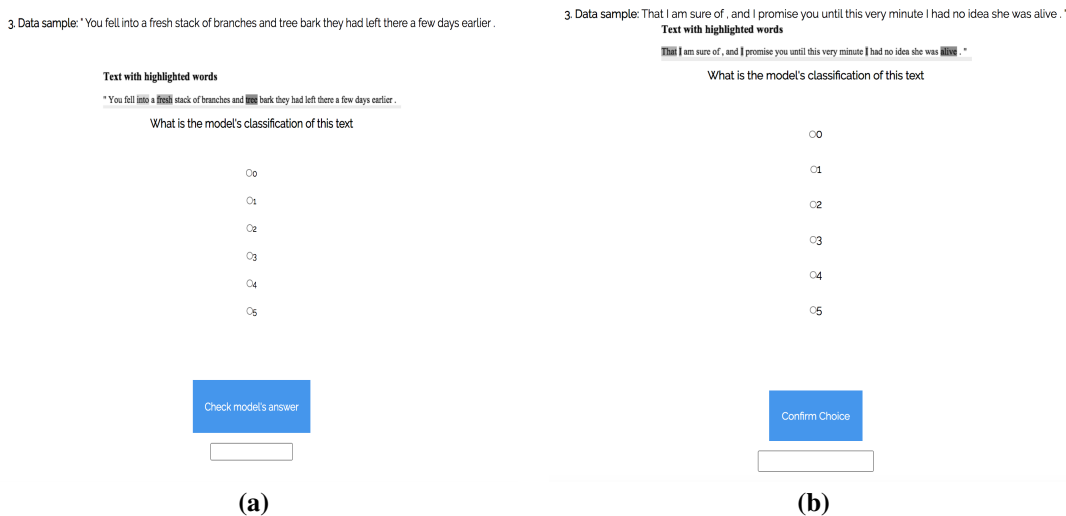


Figure 5: (a) Example of item in the training session for sentence length prediction. Note that the participants are able to check the model answer (b) Example of item in the evaluation session for sentence length prediction. Here the participants are no longer able to check the model answer

503 Subject Feedback

504 As an optional part of our tests, subjects provided some insight into the strategies they came up with
 505 or troubles they had when solving a task. We only had this feedback from some of the participants,
 506 which can be found in Table

⁴<https://flask.palletsprojects.com/en/1.1.x/>

2.

Explanation	Task	Strategy
none lime	Binary Binary	Tried to identify what kinds of data the ML model fails 1. Read the sentence. 2. Paid attention to the shaded words: if the overall sentiment of these words was clear I assumed the model would classify them accordingly. Otherwise I tried to consider how easy it would be for the model to understand the compositional meaning of the sentence assuming it will make mistakes at phenomena involving ironies or comparsions to proper names etc.
none	Binary	logical
none lime none	Hateval Hateval Hateval	Keywords, the sentimental polarity of the sentence only look at highlighted words logical
lime none	Sent Len Sent Len	Haven't got the faintest idea. I was very lost in this task. I could not find topics in the sentences so I tried to focus whether sentences contained similar words guessing that these would be mapped to the same class...
lime none	Tense Tense	no clue 1st Person 1 Person vs 2nd Person Multiple participants

Table 2: Feedback on strategies found by participants. Writing a strategy was not mandatory therefore we do not have written feedback from every participant.