

NOISYICL: A Little Noise in Model Parameters Calibrates In-context Learning

Anonymous ACL submission

Abstract

In-Context Learning (ICL) is suffering from unsatisfactory performance and under-calibration due to high prior bias and unfaithful confidence. Some previous works fine-tuned language models for better ICL performance with enormous datasets and computing costs. In this paper, we propose NOISYICL, simply perturbing the model parameters by random noises to strive for better performance and calibration. Our experiments on two models and 12 downstream datasets show that NOISYICL can help ICL produce more accurate predictions. Our further analysis indicates that NOISYICL enables the model to provide more fair predictions, and also with more faithful confidence. Therefore, we believe that NOISYICL is an effective calibration of ICL. Our experimental code is uploaded to Github¹.

1 Introduction

Language Models (LMs) have demonstrated the ability of In-Context Learning (ICL), where LMs learn tasks from few-shot input-label demonstrations in the form of natural language, without explicit parameter updates (Dong et al., 2022).

Nevertheless, ICL still underperforms the end-to-end models (Mosbach et al., 2023). A recent study shows that vanilla LMs are biased towards the knowledge acquired during pre-training, which is deemed harmful to ICL (Fei et al., 2023). There has been some effort in fine-tuning or calibrating LMs towards ICL tasks (Min et al., 2022; Zhao et al., 2021; Wei et al., 2021; Fei et al., 2023; Lu et al., 2022), which focus on reducing the gap between the knowledge gained from pre-training and ICL tasks, producing significant improvements in the ICL performance. However, fine-tuning these enormous LMs on additional data incurs a considerably high computational cost.

¹Not available during anonymous review.

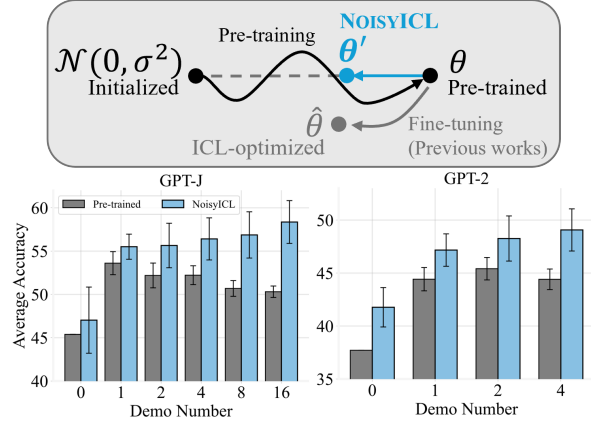


Figure 1: **Upper:** A sketch of NOISYICL: Unlike previous works which fine-tuned LMs towards ICL tasks, we perturb LMs by random noise sampled from the normal distribution $\mathcal{N}(0, \sigma^2)$ with intensity λ , then perform ICL. **Lower:** The average accuracy of ICL with and without NOISYICL w.r.t. the number of demos.

Inspired by Wu et al. (2022), which shows the benefits of noise for fine-tuning end-to-end LMs, we hypothesize that introducing noise to the model parameters of pre-trained LMs can fit LMs towards ICL with a lower computational cost. In this paper, we propose NOISYICL, which simply adds noise to model parameters and then performs ICL on the noised models, as shown in Fig. 1, to investigate the benefit of noise for ICL.

As shown in Fig. 1 and Table 1, our experiments on two models and 12 classification datasets show that adding appropriate noise into model parameters improves the performance of ICL by around 10% with obviously lower computational cost.

To investigate the reason for such performance improvement, our further experiments hypothesize that NOISYICL acts as a model calibration. In detail, we find that: **1.** NOISYICL neutralizes prediction bias among general tokens and labels, and **2.** NOISYICL leads the model to produce more faithful confidence.

Table 1: Accuracy and Macro-F1 results ($mean_{std}$, $k = 4$). A better result is in **green**. λ : The optimal intensity of noise searched on the validation set, **Acc.**: Accuracy, **MF1**: Macro-F1, **ECE₁**: the Expected Calibrated Error described in §3.3; **w/o**: Not using NOISYICL, **w/**: Using NOISYICL; Datasets abbreviation described in Appendix. **A**; **Val.**: the results on the validation sets, **Test**: the results on the testing sets.

	Dataset		PS	HS	SE'14R	SE'14L	RTE	MRPC	Ethos	FP	SST2	TEE	TES	TEH	Mean	
	λ		0.2	0.2	0.1	0.1	0.012	0.2	0.004	0.08	0.004	0.01	0.08	0.002	—	
GPT-J	Acc. (%)	w/o	37.46 _{1.36}	71.97 _{1.19}	36.48 _{0.84}	34.82 _{0.21}	49.88 _{1.03}	43.18 _{1.48}	53.44 _{0.56}	62.36 _{0.17}	81.48 _{0.97}	45.68 _{0.48}	49.32 _{1.07}	50.04 _{1.68}	51.34	
		w/	59.82 _{5.82}	82.95 _{1.36}	51.54 _{7.44}	45.14 _{4.06}	51.00 _{1.35}	59.82 _{6.15}	56.00 _{0.89}	62.42 _{0.43}	82.52 _{1.03}	46.35 _{0.53}	50.90 _{1.32}	50.35 _{1.06}	61.23	
		Val.	33.09 _{1.76}	81.30 _{0.79}	33.67 _{1.31}	34.26 _{1.33}	46.52 _{1.11}	50.97 _{1.43}	57.74 _{1.38}	61.87 _{0.16}	79.15 _{1.17}	44.72 _{0.62}	50.20 _{0.85}	53.16 _{1.07}	52.22	
		Test	55.96 _{3.43}	87.16 _{5.29}	46.79 _{4.47}	41.77 _{3.80}	47.64 _{1.27}	53.43 _{5.23}	57.07 _{1.53}	61.90 _{0.30}	78.82 _{0.75}	45.17 _{0.57}	48.11 _{1.38}	53.12 _{1.10}	56.41	
	MF1 (%)	w/o	24.11 _{0.87}	26.75 _{1.38}	32.53 _{1.46}	28.62 _{0.68}	47.37 _{1.14}	43.18 _{1.48}	53.24 _{0.69}	25.91 _{0.26}	81.24 _{1.04}	25.65 _{0.54}	32.87 _{1.81}	49.01 _{1.81}	39.20	
		w/	20.80 _{1.13}	24.22 _{0.63}	46.88 _{6.70}	43.25 _{4.12}	48.61 _{1.38}	48.38 _{2.69}	55.89 _{0.86}	27.81 _{0.74}	82.33 _{1.06}	26.31 _{0.80}	38.19 _{5.33}	49.37 _{0.92}	42.67	
		Val.	21.86 _{0.86}	26.61 _{1.61}	30.71 _{1.47}	33.17 _{1.52}	44.94 _{1.30}	49.56 _{1.48}	57.24 _{1.45}	26.59 _{0.52}	77.44 _{1.36}	23.57 _{0.86}	34.23 _{1.25}	53.16 _{1.07}	39.92	
		Test	22.87 _{1.57}	24.22 _{0.89}	43.88 _{3.90}	42.17 _{4.22}	45.93 _{1.31}	46.51 _{1.56}	56.46 _{1.54}	27.24 _{1.22}	77.11 _{0.84}	23.78 _{0.84}	32.68 _{2.17}	53.11 _{1.10}	41.33	
	ECE ₁ (% , ↓)	w/o	28.70 _{1.78}	14.39 _{1.13}	31.20 _{0.69}	38.56 _{0.36}	29.35 _{0.95}	27.79 _{1.50}	23.20 _{0.62}	23.90 _{0.17}	11.45 _{0.93}	42.79 _{0.36}	17.57 _{0.81}	28.45 _{1.23}	26.44	
		w/	12.77 _{2.94}	9.87 _{2.56}	13.96 _{7.26}	21.34 _{6.75}	28.11 _{0.95}	8.79 _{3.41}	20.33 _{0.99}	21.43 _{3.83}	42.03 _{0.88}	10.28 _{5.37}	28.18 _{1.37}	19.16		
		Val.	30.33 _{1.42}	9.41 _{0.78}	34.45 _{1.39}	36.78 _{1.56}	32.42 _{1.10}	28.29 _{1.47}	18.97 _{1.45}	23.70 _{0.49}	10.75 _{1.02}	44.09 _{0.79}	16.56 _{0.56}	24.50 _{0.99}	25.85	
		Test	9.65 _{2.83}	5.97 _{3.45}	17.53 _{5.87}	23.85 _{5.43}	31.49 _{1.45}	20.24 _{5.56}	19.28 _{1.91}	22.77 _{3.17}	10.62 _{0.84}	43.81 _{0.76}	14.43 _{2.40}	24.67 _{0.92}	20.36	
GPT-2	Acc. (%)	λ	0.014	0.1	0.08	0.06	0.002	0.08	0.2	0.04	0.014	0.04	0.04	0.2	—	
		w/o	45.44 _{1.50}	39.26 _{0.95}	44.65 _{0.51}	42.83 _{0.93}	51.41 _{0.76}	69.06 _{0.19}	43.20 _{0.41}	53.03 _{1.35}	59.63 _{0.28}	29.67 _{0.77}	38.55 _{1.43}	41.82 _{0.51}	46.54	
		w/	49.17 _{1.43}	64.06 _{3.41}	46.05 _{0.66}	46.99 _{1.77}	52.11 _{1.50}	59.65 _{3.26}	51.64 _{1.67}	58.26 _{2.02}	61.50 _{1.23}	31.17 _{0.94}	42.44 _{1.86}	51.48 _{4.38}	51.21	
		Test	39.62 _{1.33}	43.62 _{1.17}	40.03 _{1.26}	39.41 _{1.29}	50.89 _{1.04}	63.12 _{0.38}	45.93 _{0.51}	52.27 _{0.81}	59.47 _{1.41}	27.43 _{0.59}	29.48 _{1.29}	41.66 _{0.57}	44.41	
	MF1 (%)	w/o	40.95 _{1.42}	67.25 _{4.00}	45.28 _{3.50}	36.90 _{0.18}	51.01 _{0.91}	61.20 _{0.68}	50.90 _{2.63}	56.19 _{1.18}	61.54 _{1.51}	33.38 _{1.16}	34.01 _{1.15}	50.28 _{2.44}	49.07	
		Val.	w/o	25.62 _{1.81}	17.88 _{0.37}	37.69 _{0.08}	37.71 _{0.88}	50.82 _{0.67}	41.45 _{0.49}	34.36 _{0.41}	35.02 _{2.00}	57.74 _{0.42}	20.33 _{0.49}	33.74 _{1.86}	30.98 _{0.50}	35.28
		w/	26.05 _{1.59}	24.20 _{0.85}	32.40 _{0.25}	30.78 _{1.69}	51.57 _{1.59}	48.90 _{1.11}	49.38 _{0.17}	35.35 _{1.00}	61.28 _{1.18}	22.14 _{0.91}	34.46 _{1.13}	46.20 _{3.91}	38.56	
		Test	w/o	24.70 _{1.17}	18.24 _{0.52}	35.78 _{1.49}	38.79 _{1.25}	49.37 _{1.31}	39.90 _{0.94}	35.56 _{0.75}	35.78 _{1.11}	59.46 _{1.42}	18.89 _{0.62}	29.36 _{1.29}	34.43 _{0.79}	35.02
	ECE ₁ (% , ↓)	w/	24.49 _{1.20}	23.74 _{1.10}	34.03 _{1.20}	27.41 _{1.97}	49.56 _{0.99}	41.46 _{1.26}	47.25 _{3.97}	36.67 _{2.14}	61.15 _{1.50}	24.06 _{1.39}	31.48 _{3.32}	49.05 _{1.35}	37.52	
		Val.	w/o	5.94 _{0.81}	36.18 _{0.96}	14.79 _{0.92}	15.51 _{1.25}	30.35 _{1.12}	19.39 _{0.40}	47.50 _{0.37}	8.33 _{1.12}	2.10 _{0.88}	30.85 _{1.06}	14.17 _{1.04}	51.18 _{0.62}	23.02
		w/	5.74 _{1.40}	11.64 _{2.21}	16.36 _{4.41}	17.55 _{1.50}	29.51 _{1.54}	21.23 _{1.51}	24.01 _{0.97}	8.20 _{1.78}	2.99 _{1.10}	30.92 _{1.99}	13.52 _{1.18}	23.93 _{5.48}	17.13	
		Test	w/o	9.29 _{1.37}	28.14 _{1.21}	17.90 _{1.04}	15.04 _{1.29}	31.42 _{1.13}	23.02 _{0.54}	45.21 _{0.78}	8.00 _{0.92}	1.89 _{0.55}	32.76 _{0.80}	22.30 _{1.33}	48.69 _{0.51}	23.64
	w/	8.16 _{1.12}	9.01 _{2.57}	18.69 _{1.81}	25.35 _{2.98}	31.07 _{1.22}	26.06 _{1.64}	25.52 _{4.12}	6.88 _{1.71}	3.12 _{0.80}	29.73 _{2.28}	19.15 _{4.69}	21.63 _{2.11}	18.70		

Our contribution can be summarized as:

- We propose NOISYICL, which simply adds noise into LMs and then executes ICL (§2). Our experiment shows that NOISYICL obtains a better ICL performance (§3.2).
- We show that adding noise effectively calibrates LMs to reduce prediction bias and unfaithful confidence in ICL (§3.3).

2 NOISYICL

Here we introduce the basic form of ICL and our perturbation method named NOISYICL.

In-context Learning. Suppose a supervised classification dataset $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^n$, where x_i is an input, and $y_i \in \mathbb{U}$ is the corresponding label in a label space $\mathbb{U}$. For each query x_q to be predicted, we sample a demo sequence $G = \{(x_{a_j}, y_{a_j})\}_{j=1}^k$ from $\mathcal{D}$, where k is the number of demos, and construct a prompt input in natural language form $s = f(G, x_q)$ with a pattern f . Then, we input s into the LM $P_\theta(\cdot)$ with parameters θ and get an output token distribution $P_\theta(\cdot|s)$. We choose the label token l with the maximum probability among the label space as the prediction $\hat{y}_q$, that is:

$$\hat{y}_q = \operatorname{argmax}_{l \in \mathbb{U}} P_\theta(l|s). \quad (1)$$

Notice that we only construct prompts to drive the model to predict answers generatively, without any parameter updates.

NOISYICL. For each parameter matrix θ_i in the LM $P_\theta(\cdot)$ used for ICL, we simply do an interpolation between θ_i and a noise matrix sampled from an isotropic normal distribution $\mathcal{N}(0, \sigma^2)$ with intensity λ , that is:

$$\theta'_i = (1 - \lambda)\theta_i + \lambda\mathcal{N}(0, \sigma^2), \quad (2)$$

where λ and σ are model or task-specific hyperparameters. Then we perform the aforementioned ICL with the interpolated LM $P_{\theta'}(\cdot)$.

3 Experiments and Results

We investigate the effectiveness and calibration abilities of NOISYICL. We find that NOISYICL can improve ICL performance by an average of 10% (§3.2) and effectively calibrate the model's prediction bias and unfaithful confidence (§3.3).

3.1 Settings

Data. We use 12 commonly used classification datasets in our experiments. Since some of the datasets do not provide valid splitting, we randomly split these datasets into validation sets and test sets (see Appendix A for further details).

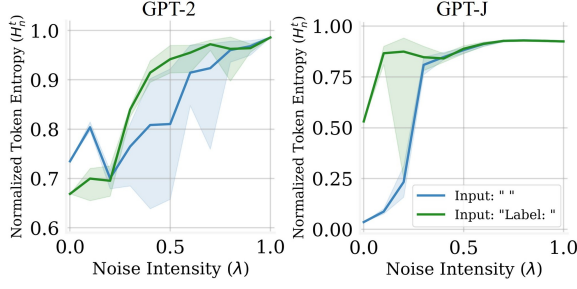


Figure 2: The correlation between the normalized token entropy H_n^t and the noise intensity λ with empty inputs. When the noise gets stronger, the H_n^t becomes higher, which indicates a fairer output.

Models. We use GPT-2 (parameters: 137M) (Radford et al., 2019) and GPT-J (parameters: 6053M) (Wang and Komatsuzaki, 2021). The model checkpoints are loaded from huggingface².

Hyperparameters. We fix σ , the standard deviation of the noise distribution, to 0.02, which is the same as the initialization distribution of both models. We believe the noise intensity λ should vary w.r.t. datasets and models. Therefore, we determine the most suitable λ by a simple search method for each dataset and model on the validation set (details in Appendix D). The selected intensities are shown in Table 1, which are concentrated in $(0, 0.3]$. Moreover, we confirm that the λ keeps relatively stable w.r.t. k given a dataset and model (Appendix D).

Other details. We default to use four demos and a simple prompt template as shown in Appendix B. Every labeled data is treated as the test query once, and for each query, we do 2 tries. Each experiment is repeated 10 times with different noise matrices.

3.2 NOISYICL Improves ICL Performance

We show accuracy and macro-F1 on the 12 downstream datasets with and without NOISYICL on the optimal λ in Table 1, and averaged results in Fig. 1 (see Appendix C for results of other numbers of demos (k)).

The results show that NOISYICL produces an average performance improvement of around 10%. This phenomenon preliminarily confirms our hypothesis: NOISYICL fits the pre-trained LMs towards ICL.

However, such gains vary depending on the dataset and model. In some combinations of

²huggingface.co/gpt2, and huggingface.co/EleutherAI/gpt-j-6b

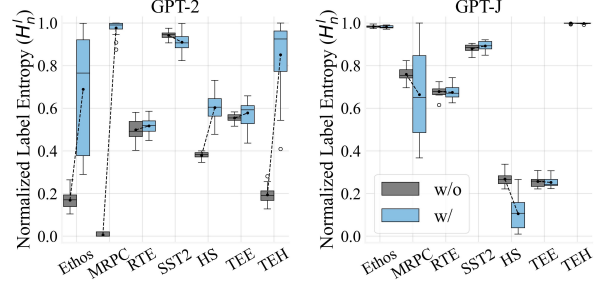


Figure 3: The normalized label entropy H_n^l on both models and 7 datasets with and without appropriate-noised NOISYICL. In most cases, the H_n^l with NOISYICL (w/) is greater than without NOISYICL (w/o).

datasets and models, significant performance improvement cannot be observed. We speculate the reason is that NOISYICL does not provide new knowledge from new training examples, and these datasets are too hard for the model intrinsically no matter whether NOISYICL is used.

3.3 NOISYICL Is A Calibration

This section finds that NOISYICL conducts the following aspects of calibrations:

- **Lower prediction bias.** When no valid query is given, the predicted labels should have balanced frequencies. However, predictions conducted by under-calibrated LMs usually have significant bias, which is harmful to ICL (Fei et al., 2023; Zhao et al., 2021; Wei et al., 2023; Lu et al., 2022). Eliminating these biases can be seen as an aspect of calibration.
- **More faithful confidence.** In classification tasks, the outputted probability of the predicted label is called confidence. Suitable confidence should faithfully reflect the accuracy of outputs, that is, a prediction with greater accuracy should be assigned with greater confidence (Corbière et al., 2019). It has been proven that faithful confidence improves the stability of the model (Guo et al., 2017; Grabinski et al., 2022). Making the confidence more faithful is also an aspect of calibration (Guo et al., 2017; Tian et al., 2023).

Specifically, we find that NOISYICL induces the model to predict labels with more fair and faithful confidence.

NOISYICL reduces prediction bias. We investigate the following two types of prediction bias. (See Appendix E for further calculation details.)

(1) **Intrinsic token bias.** LMs are pre-trained with natural corpus, where different tokens have various frequencies. LMs learn such frequencies and act as biases among different tokens in prediction, which are deemed harmful to ICL (Fei et al., 2023). We believe that such token bias can be reduced by NOISYICL.

To quantify this, we calculate normalized token entropy H_n^t : the entropy of the predicted probability distribution among the whole vocabulary given an empty input (we use “”, and “Label: ”), w.r.t. various noise intensities.

The results are shown in Fig. 2. A clear positive correlation between the noise intensity and H_n^t can be observed, meaning the model gives a fairer output when more noise is given.

(2) **Overall label bias.** In classification tasks, when no query is given, equal probabilities should be assigned to labels fairly, which benefits ICL (Lu et al., 2022). We believe that NOISYICL promotes such fairness.

To quantify this, we use normalized label entropy H_n^l : the entropy of the predicted probability distribution among the label space given some demos and an empty query. Notice that it is natural for a model to predict a “neutral” label when no query is given, so label bias on a dataset with a “neutral” label cannot be fairly measured with H_n^l . Therefore, we ignore these datasets with “neutral” labels (details in Appendix A).

The results are shown in Fig. 3, indicating that in most cases, NOISYICL calibrates the overall bias, especially on GPT-2.

NOISYICL promotes faithful confidence. The Expected Calibration Error (ECE_p) (Naeini et al., 2015) is a widely-used indicator for the faithfulness of confidence in classification tasks:

$$ECE_p = \mathbb{E}(|\max(\hat{z}) - \mathbb{E}_{y=\arg\max_i \hat{z}_i}(1)^p|)^{\frac{1}{p}}, \quad (3)$$

where $\hat{z}$ is the predicted probability vector on the label space by a classification model, the final prediction ($\arg\max_i \hat{z}_i$) can be obtained with a confidence ($\max \hat{z}$), and the ground-truth label is y .

Let $p = 1$, we use the ECE_1 to investigate the faithfulness of the ICL output. The details of the calculation are shown in Appendix F. A lower ECE_1 means more faithful confidence, that is, the confidence becomes a more accurate prediction of accuracy (Corbière et al., 2019). Fig. 4 shows a case study of GPT-2 on the hate_speech18

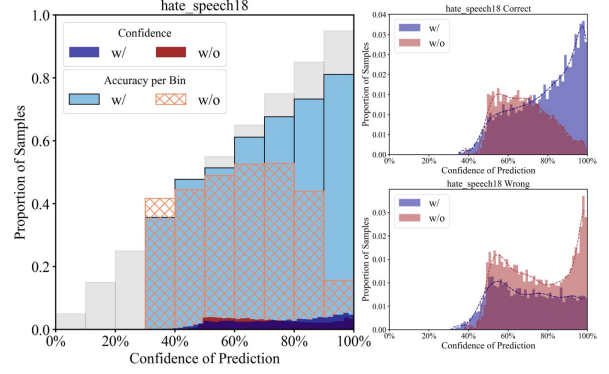


Figure 4: **Left:** Reliability diagrams (sparse bars) and global confidence distribution (dense bars) of GPT-2 on hate_speech18 with (w/, $ECE_1 = 9.01\%$) and without (w/o, $ECE_1 = 28.14\%$) NOISYICL. The predictions are divided into bins according to confidence, and we visualize the accuracy of each bin as a histogram. The grey bars are ideal, which is closer to the one with NOISYICL. **Right upper:** Confidence distribution on correct predictions. Relatively right-shifted with NOISYICL. **Right lower:** Confidence distribution on wrong predictions. Relatively left-shifted with NOISYICL. ($k = 4$)

dataset, indicating that the prediction is more faithful with NOISYICL (full visualized data are in Appendix G). We show the results with and without appropriate-noised NOISYICL for ECE_1 on the 12 datasets in Table 1.

In most cases, the ECE_1 is reduced by around 25% by NOISYICL, which means the confidence is more faithful with NOISYICL. This suggests that NOISYICL can drive the model to produce more faithful confidence, that is, less over-confidence in wrong predictions, and less under-confidence in correct predictions, which suggests that NOISYICL solves the confidence calibration.

Based on the above two investigations, we believe that NOISYICL is an effective calibration for the ICL scenario. This can be a reasonable explanation for the observed performance improvement in §3.2.

4 Conclusion

In this paper, we propose NOISYICL, which simply adds random noise to the parameters of LMs to fit them from pre-trained knowledge to ICL. We show that NOISYICL can improve the ICL performance and also calibrate the model for fairer outputs and more faithful confidence.

5 Limitations

As mentioned before, unlike the fine-tuning on additional ICL-style datasets (Min et al., 2022; Zhao et al., 2021; Wei et al., 2021), NOISYICL does not provide new knowledge for the model, so the calibrated model can not discover tasks that are not potentially included in the pre-training data (Gu et al., 2023). Moreover, a naive search for the best noise intensity is not efficient and satisfactory. An effective detection of λ should be developed. Currently, we can only confirm that for one dataset and model, the λ keeps relatively stable, especially when a large k is given (Appendix. D).

Future Works. Besides fixing the limits, future works can focus on where and how the noise should be introduced. In Transformer-based models, different layers have different abilities (Wang et al., 2023; Jawahar et al., 2019; Kobayashi et al., 2023). So, treating these layers differently may be an effective improvement of NOISYICL. An isotropic normal distribution is too simple, noise sampling methods should be discussed.

Moreover, adding noise to model parameters can be an erasing of pre-training (Ilharco et al., 2022), so, the search for λ is the search for the best checkpoint of pre-training. With these checkpoints, we can determine (Han et al., 2023; Han and Tsvetkov, 2022) which data is disadvantageous to ICL, and what knowledge is essential to ICL, to better reveal the essence of ICL.

References

- 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.
- Luisa Bentivogli, Peter Clark, Ido Dagan, and Danilo Giampiccolo. 2009. The fifth pascal recognizing textual entailment challenge. *TAC*, 7:8.
- Charles Corbière, Nicolas Thome, Avner Bar-Hen, Matthieu Cord, and Patrick Pérez. 2019. Addressing failure prediction by learning model confidence. *Advances in Neural Information Processing Systems*, 32.
- Ido Dagan, Oren Glickman, and Bernardo Magnini. 2005. The pascal recognising textual entailment challenge. In *Machine learning challenges workshop*, pages 177–190. Springer.
- Ona de Gibert, Naiara Perez, Aitor García-Pablos, and Montse Cuadros. 2018. *Hate Speech Dataset from a White Supremacy Forum*. In *Proceedings of the 2nd Workshop on Abusive Language Online (ALW2)*, pages 11–20, Brussels, Belgium. Association for Computational Linguistics.
- Bill Dolan and Chris Brockett. 2005. Automatically constructing a corpus of sentential paraphrases. In *Third International Workshop on Paraphrasing (IWP2005)*.
- Qingxiu Dong, Lei Li, Damai Dai, Ce Zheng, Zhiyong Wu, Baobao Chang, Xu Sun, Jingjing Xu, and Zhifang Sui. 2022. A survey for in-context learning. *arXiv preprint arXiv:2301.00234*.
- Yu Fei, Yifan Hou, Zeming Chen, and Antoine Bosselut. 2023. *Mitigating label biases for in-context learning*. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 14014–14031, Toronto, Canada. Association for Computational Linguistics.
- Danilo Giampiccolo, Bernardo Magnini, Ido Dagan, and William B Dolan. 2007. The third pascal recognizing textual entailment challenge. In *Proceedings of the ACL-PASCAL workshop on textual entailment and paraphrasing*, pages 1–9.
- Julia Grabinski, Paul Gavrikov, Janis Keuper, and Margret Keuper. 2022. Robust models are less overconfident. *Advances in Neural Information Processing Systems*, 35:39059–39075.
- Yuxian Gu, Li Dong, Furu Wei, and Minlie Huang. 2023. Pre-training to learn in context. *arXiv preprint arXiv:2305.09137*.
- Chuan Guo, Geoff Pleiss, Yu Sun, and Kilian Q Weinberger. 2017. On calibration of modern neural networks. In *International conference on machine learning*, pages 1321–1330. PMLR.
- R Bar Haim, Ido Dagan, Bill Dolan, Lisa Ferro, Danilo Giampiccolo, Bernardo Magnini, and Idan Szpektor. 2006. The second pascal recognising textual entailment challenge. In *Proceedings of the Second PASCAL Challenges Workshop on Recognising Textual Entailment*, volume 7, pages 785–794.
- Xiaochuang Han, Daniel Simig, Todor Mihaylov, Yulia Tsvetkov, Asli Celikyilmaz, and Tianlu Wang. 2023. Understanding in-context learning via supportive pretraining data. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 12660–12673.
- Xiaochuang Han and Yulia Tsvetkov. 2022. Orca: Interpreting prompted language models via locating supporting data evidence in the ocean of pretraining data. *arXiv preprint arXiv:2205.12600*.

356	Gabriel Ilharco, Marco Tulio Ribeiro, Mitchell Wortsman, Ludwig Schmidt, Hannaneh Hajishirzi, and Ali Farhadi. 2022. Editing models with task arithmetic. In <i>The Eleventh International Conference on Learning Representations</i> .	409
357		410
358		411
359		412
360		413
361	Ganesh Jawahar, Benoît Sagot, and Djamé Seddah. 2019. What does bert learn about the structure of language? In <i>ACL 2019-57th Annual Meeting of the Association for Computational Linguistics</i> .	414
362		415
363		
364		
365	Goro Kobayashi, Tatsuki Kuribayashi, Sho Yokoi, and Kentaro Inui. 2023. Transformer language models handle word frequency in prediction head . In <i>Findings of the Association for Computational Linguistics: ACL 2023</i> , pages 4523–4535, Toronto, Canada. Association for Computational Linguistics.	416
366		417
367		418
368		
369		
370		
371	Yao Lu, Max Bartolo, Alastair Moore, Sebastian Riedel, and Pontus Stenetorp. 2022. Fantastically ordered prompts and where to find them: Overcoming few-shot prompt order sensitivity. In <i>Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 8086–8098.	419
372		420
373		421
374		422
375		423
376		
377		
378	P. Malo, A. Sinha, P. Korhonen, J. Wallenius, and P. Takala. 2014. Good debt or bad debt: Detecting semantic orientations in economic texts. <i>Journal of the Association for Information Science and Technology</i> , 65.	424
379		425
380		426
381		
382		
383	Sewon Min, Mike Lewis, Luke Zettlemoyer, and Hannaneh Hajishirzi. 2022. Metaicl: Learning to learn in context. In <i>Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 2791–2809.	427
384		428
385		429
386		430
387		431
388		432
389	Saif Mohammad, Felipe Bravo-Marquez, Mohammad Salameh, and Svetlana Kiritchenko. 2018. Semeval-2018 task 1: Affect in tweets. In <i>Proceedings of the 12th international workshop on semantic evaluation</i> , pages 1–17.	433
390		
391		
392		
393		
394	Ioannis Mollas, Zoe Chrysopoulou, Stamatis Karlos, and Grigorios Tsoumakas. 2020. Ethos: an online hate speech detection dataset .	434
395		435
396		436
397	Marius Mosbach, Tiago Pimentel, Shauli Ravfogel, Dietrich Klakow, and Yanai Elazar. 2023. Few-shot fine-tuning vs. in-context learning: A fair comparison and evaluation. <i>arXiv preprint arXiv:2305.16938</i> .	437
398		438
399		439
400		
401	Mahdi Pakdaman Naeini, Gregory Cooper, and Milos Hauskrecht. 2015. Obtaining well calibrated probabilities using bayesian binning. In <i>Proceedings of the AAAI conference on artificial intelligence</i> , volume 29.	440
402		441
403		442
404		443
405	Jane Pan, Tianyu Gao, Howard Chen, and Danqi Chen. 2023. What in-context learning "learns" in-context: Disentangling task recognition and task learning. <i>arXiv preprint arXiv:2305.09731</i> .	444
406		445
407		446
408		447
		448
	Maria Pontiki, Dimitris Galanis, John Pavlopoulos, Harris Papageorgiou, Ion Androutsopoulos, and Suresh Manandhar. 2014. SemEval-2014 task 4: Aspect based sentiment analysis . In <i>Proceedings of the 8th International Workshop on Semantic Evaluation (SemEval 2014)</i> , pages 27–35, Dublin, Ireland. Association for Computational Linguistics.	449
		450
		451
		452
		453
		454
		455
		456
	Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners.	457
		458
		459
		460
		461
	Sara Rosenthal, Noura Farra, and Preslav Nakov. 2017. Semeval-2017 task 4: Sentiment analysis in twitter. In <i>Proceedings of the 11th international workshop on semantic evaluation (SemEval-2017)</i> , pages 502–518.	462
		463
	Emily Sheng and David Uthus. 2020. Investigating societal biases in a poetry composition system. <i>arXiv preprint arXiv:2011.02686</i> .	464
		465
		466
	Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D Manning, Andrew Y Ng, and Christopher Potts. 2013. Recursive deep models for semantic compositionality over a sentiment treebank. In <i>Proceedings of the 2013 conference on empirical methods in natural language processing</i> , pages 1631–1642.	467
		468
		469
	Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher D Manning. 2023. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. <i>arXiv preprint arXiv:2305.14975</i> .	470
		471
		472
	Alex Wang, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel R Bowman. 2018. Glue: A multi-task benchmark and analysis platform for natural language understanding. In <i>International Conference on Learning Representations</i> .	473
		474
		475
	Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A 6 Billion Parameter Autoregressive Language Model. https://github.com/kingoflolz/mesh-transformer-jax .	476
		477
		478
	Lean Wang, Lei Li, Damai Dai, Deli Chen, Hao Zhou, Fandong Meng, Jie Zhou, and Xu Sun. 2023. Label words are anchors: An information flow perspective for understanding in-context learning . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 9840–9855, Singapore. Association for Computational Linguistics.	479
		480
		481
		482
		483
		484
		485
		486
		487
		488
		489
		490
		491
		492
		493
		494
		495
		496
		497
		498
		499
		500

Zhou, Tengyu Ma, et al. 2023. Symbol tuning improves in-context learning in language models. *arXiv preprint arXiv:2305.08298*.

Chuhan Wu, Fangzhao Wu, Tao Qi, and Yongfeng Huang. 2022. Noisy tune: A little noise can help you finetune pretrained language models better. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 680–685.

Zihao Zhao, Eric Wallace, Shi Feng, Dan Klein, and Sameer Singh. 2021. Calibrate before use: Improving few-shot performance of language models. In *International Conference on Machine Learning*, pages 12697–12706. PMLR.

A Datasets

The datasets used in this paper are shown in Table 2.

*To construct inputs of appropriate length, we remove data with a length exceeding 500 from the Ethos, and the number of the remaining data is 980.

Dataset Splitting. For each dataset, we first shuffle it with random seed 42. Then, we choose the 512 data at the tail as the testing data, and the 512 data at the head (if the total number of data is less than 1024, we choose the rest of the data.) as the validation data.

B Prompt Patterns

In this paper, we use a minimum prompt template. For each task category, we design a template as shown below.

For single-sentence classification datasets (x, y) , we use:

Input: $\langle x \rangle$, Label: $\langle y \rangle \setminus n$
 ...
 Input: $\langle x \rangle$, Label:

For aspect-based sentiment classification datasets $((x, a), y)$, we use:

Input: $\langle x \rangle$, Aspect: $\langle a \rangle$, Label: $\langle y \rangle \setminus n$
 ...
 Input: $\langle x \rangle$, Aspect: $\langle a \rangle$, Label:

For double-sentence classification datasets $((x_1, x_2), y)$. we use:

Input: $\langle x_1 \rangle$, Text 2: $\langle x_2 \rangle$, Label: $\langle y \rangle \setminus n$
 ...
 Input: $\langle x_1 \rangle$, Text 2: $\langle x_2 \rangle$, Label:

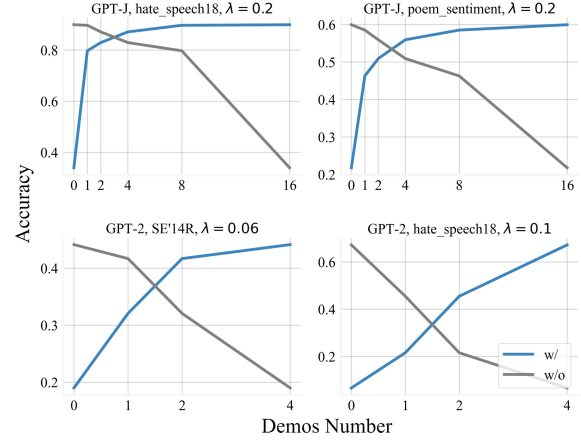


Figure 5: The relationship between demos quantity and accuracy in some cases. NOISYICL can make the model learn from the demos correctly.

Input: This is delicious!	Label: Positive
Input: I don't like it.	Label: Negative
Input: Boring.	Label: Negative
Input: It is fantastic!	Label: Positive
Input:	Label:

Figure 6: An example of inputs in the normalized label entropy calculation.

C Complete Results: Accuracy, Macro-F1, and ECE_1 With Various Demo Numbers (k)

We use various numbers of demos in the experiments shown in the §3.2.

The results of zero-shot ($k = 0$) are shown in Table 3.

The results of 1-shot ($k = 1$) are shown in Table 4.

The results of 2-shot ($k = 2$) are shown in Table 5.

The results of 8-shot ($k = 8$) are shown in Table 6.

The results of 16-shot ($k = 16$) are shown in Table 7.

Notice that in Table 6 and Table 7, some experiments are not feasible due to the length of sequences.

In these results, NOISYICL still outperforms the baseline, which proves that NOISYICL is generally applicable with various k .

D Searching and Stability of λ

We select a set of candidates of λ as $\{0, 0.002, 0.004, 0.006, 0.008, 0.01, 0.012, 0.014, 0.016, 0.018, 0.02, 0.04, 0.06, 0.08, 0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9, 1.0\}$, test the performance of

Table 2: Datasets and Abbreviations used in this paper. **Abbr.**: Abbreviation, **Data#**: The total number of data, **Label#**: The number of classes, **Neutral?**: Does the dataset have a neutral label?, **Major. ~ Minor.(%)**: The proportion of majority labels ~ the proportion of minority labels.

Dataset	Abbr.	Data#	Label#	Neutral?	Major. ~ Minor.(%)
<i>single-sentence classification:</i>					
poem_sentiment (Sheng and Uthus, 2020)	PS	1101	4	✓	62.2 ~ 5.5
hate_speech18 (de Gibert et al., 2018)	HS	10944	4	×	86.9 ~ 1.5
Ethos* (binary) (Mollas et al., 2020)	—	980	2	×	56.6 ~ 43.4
financial_phrasebank (all agree) (Malo et al., 2014)	FP	2264	3	✓	61.4 ~ 13.4
GLUE-SST2 (Wang et al., 2018; Socher et al., 2013)	SST2	68221	2	×	55.8 ~ 44.2
tweet_eval_emotion (Mohammad et al., 2018)	TEE	5052	4	×	43 ~ 9
tweet_eval_sentiment (Rosenthal et al., 2017)	TES	59899	3	✓	45.3 ~ 15.5
tweet_eval_hate (Basile et al., 2019)	TEH	12970	2	×	58 ~ 42
<i>aspect-based sentiment classification:</i>					
SemEval 2014-Task 4 Restaurants (Pontiki et al., 2014)	SE'14R	4722	3	✓	61.2 ~ 17.6
SemEval 2014-Task 4 Laptops (Pontiki et al., 2014)	SE'14L	2951	3	✓	45.3 ~ 15.5
<i>double-sentence classification:</i>					
GLUE-RTE (Wang et al., 2018; Dagan et al., 2005)	RTE	2767	2	×	50.2 ~ 49.8
(Haim et al., 2006; Giampiccolo et al., 2007; Bentivogli et al., 2009)					
GLUE-MRPC (Wang et al., 2018; Dolan and Brockett, 2005)	MRPC	4076	2	×	67.4 ~ 32.6

NOISYICL with this set of λ , and select the one with the best accuracy as the optimal λ .

Moreover, we find that given a dataset and model, λ remains relatively stable w.r.t the demo number k , especially when k is large. As shown in Fig. 7 and Fig. 8, we calculate $|\lambda_i - \lambda_j|$, the distance between the optimal λ of different demo numbers as the heatmap, and the normalized remaining range $R(\lambda_i)/\max(\lambda)$, where the $R(\lambda_i)$ is the range of all the optimal λ_c where the demo number $c \geq i$, and the $\max(\lambda)$ is the maximum of all the optimal λ on the dataset. It characterizes the dispersion of λ when the demo number is greater than the given one i .

The results show that while the k is increasing, the distance and the normalized remaining range quickly zero out, which supports our hypothesis: NOISYICL is a bridge from the pre-training to ICL. For that during this process, the demo number should have a minimal impact.

E Calculation of H_n^t and H_n^l

Normalized Token Entropy. We calculate the normalized token entropy H_n^t of the predicted probability distribution of tokens with an empty input as:

$$H_n^t = -\frac{\sum_{l \in \mathbb{V}} P_\theta(l|x_\emptyset) \ln P_\theta(l|x_\emptyset)}{\ln |\mathbb{V}|}, \quad (4)$$

where the $P_\theta(\cdot)$ is an LM with a vocabulary space $\mathbb{V}$ and size $|\mathbb{V}|$, the $x_\emptyset$ is an empty input, such as “”, or “Label: ”. The model outputs a global probability distribution of tokens when the $x_\emptyset$ is

given, and we calculate the normalized entropy on the distribution.

Normalized Label Entropy. We calculate the normalized label entropy H_n^l of the model prediction probability distribution on the label space with 4 demos and an empty queue from a dataset (Fig. 6 is an example) as:

$$H_n^l = -\frac{\sum_{l \in \mathbb{U}} r_{P,\theta}(l|x_*) \ln r_{P,\theta}(l|x_*)}{\ln |\mathbb{U}|}, \quad (5)$$

where the $r_{P,\theta}(l|x_*)$ is the predicted frequency of label l given the aforementioned query-less only-demos input x_* . We use 512 tries to estimate this frequency. The $\mathbb{U}$ and $|\mathbb{U}|$ are the label space and label space size, respectively.

F Calculation of ECE_1

Given a prediction set $\mathcal{Y} = \{(\hat{y}_i, z_i, y_i)\}_{i=1}^n$ of predictions $\hat{y}_i$, predicted confidence $z_i \in [0, 1]$, and supervised label y_i , in the calculation of ECE_1 , we first divide the confidence space $(0, 1)$ into $m = 10$ equal bins $B_j = [0.1(j-1), 0.1j]_{j=1}^m$ according to the z_i as shown in Fig. 4. The amount of prediction in each bin is denoted as $|B_j|$. For each bin, we calculate the accuracy $\alpha_j = \frac{1}{|B_j|} \sum_{\hat{y}_i \in B_j, \hat{y}_i = y_1} 1$ of predictions in the j th bin, and assign a standard accuracy based on the average of confidence in the bin, that is, $\alpha_j^s = \frac{1}{|B_j|} \sum_{z_i \in B_j} z_i$. The ECE_1 can be described as:

$$ECE_1 = \sum_{j=1}^m \frac{|B_j|}{n} |\alpha_j - \alpha_j^s|, \quad (6)$$

Table 3: Accuracy and Macro-F1 results (% , $mean_{std}$, $k = 0$). A better result is in **green**. Notation is the same with the Table 1.

Dataset		PS	HS	SE'14R	SE'14L	RTE	MRPC	Ethos	FP	SST2	TEE	TES	TEH	Mean	
λ		0.1	0.1	0.1	0.1	0.06	0.02	0.004	0.02	0.1	0.2	0.008	0.3	—	
Acc. (%)	Val.	w/o	16.60 _{0.00}	37.30 _{0.00}	33.98 _{0.00}	38.87 _{0.00}	52.54 _{0.00}	65.82 _{0.00}	51.07 _{0.00}	61.13 _{0.00}	57.42 _{0.00}	39.84 _{0.00}	49.22 _{0.00}	43.36 _{0.00}	45.60
		w/	21.88 _{6.75}	50.70 _{4.33}	40.74 _{10.27}	45.94 _{2.72}	53.40 _{0.55}	67.77 _{0.55}	51.07 _{0.19}	61.37 _{0.60}	60.98 _{3.87}	39.57 _{3.50}	49.53 _{0.20}	54.14 _{5.91}	49.76
	Test	w/o	17.38 _{0.00}	29.88 _{0.00}	28.91 _{0.00}	39.45 _{0.00}	49.80 _{0.00}	59.18 _{0.00}	54.10 _{0.00}	61.91 _{0.00}	66.41 _{0.00}	41.60 _{0.00}	49.22 _{0.00}	46.88 _{0.00}	45.39
		w/	22.97 _{7.52}	38.30 _{10.89}	31.25 _{5.79}	40.84 _{3.85}	49.77 _{1.46}	59.38 _{0.79}	54.36 _{0.15}	61.37 _{0.65}	65.96 _{2.27}	41.15 _{3.79}	49.14 _{0.26}	49.86 _{8.37}	47.03
GPT-J (%)	Val.	w/o	11.32 _{0.00}	20.81 _{0.00}	35.19 _{0.00}	36.56 _{0.00}	46.42 _{0.00}	48.13 _{0.00}	48.58 _{0.00}	36.13 _{0.00}	56.98 _{0.00}	15.99 _{0.00}	29.81 _{0.00}	33.40 _{0.00}	34.94
		w/	14.87 _{2.10}	25.82 _{2.65}	39.90 _{8.62}	41.36 _{3.98}	50.10 _{1.31}	49.08 _{0.54}	48.58 _{0.25}	37.35 _{0.74}	59.80 _{4.50}	25.30 _{2.60}	30.36 _{0.60}	39.66 _{6.23}	38.62
	Test	w/o	12.70 _{0.00}	16.02 _{0.00}	27.01 _{0.00}	39.38 _{0.00}	45.56 _{0.00}	40.48 _{0.00}	50.99 _{0.00}	34.91 _{0.00}	64.45 _{0.00}	16.26 _{0.00}	29.81 _{0.00}	42.58 _{0.00}	35.85
		w/	16.93 _{4.12}	17.92 _{3.40}	30.20 _{7.55}	40.50 _{4.43}	49.06 _{2.03}	39.96 _{0.59}	51.20 _{0.12}	35.56 _{1.38}	61.74 _{5.39}	22.54 _{5.12}	29.69 _{0.59}	39.70 _{8.87}	36.25
ECE ₁ (%, ↓)	Val.	w/o	36.56 _{0.00}	22.03 _{0.00}	28.11 _{0.00}	17.80 _{0.00}	29.83 _{0.00}	7.24 _{0.00}	28.08 _{0.00}	13.45 _{0.00}	5.45 _{0.00}	48.09 _{0.00}	22.10 _{0.00}	46.21 _{0.00}	25.41
		w/	22.52 _{6.40}	16.20 _{3.75}	16.77 _{9.48}	8.39 _{3.82}	23.90 _{1.51}	9.20 _{0.54}	28.12 _{0.24}	12.75 _{0.37}	4.22 _{2.43}	18.66 _{2.88}	22.53 _{0.19}	21.99 _{12.15}	17.10
	Test	w/o	35.33 _{0.00}	26.92 _{0.00}	33.41 _{0.00}	16.92 _{0.00}	31.46 _{0.00}	13.80 _{0.00}	24.56 _{0.00}	15.09 _{0.00}	3.49 _{0.00}	46.49 _{0.00}	22.10 _{0.00}	36.32 _{0.00}	25.49
		w/	24.86 _{6.34}	23.67 _{6.86}	26.19 _{10.00}	12.16 _{5.22}	24.43 _{2.49}	13.65 _{0.60}	24.49 _{0.13}	13.34 _{0.90}	5.91 _{2.43}	25.21 _{14.36}	22.37 _{0.36}	24.95 _{16.80}	20.10
λ		0.04	0.002	0.1	0.1	0.012	0.06	0.1	0.06	0.02	0.02	0.002	0.1	—	
Acc. (%)	Val.	w/o	4.30 _{0.00}	15.43 _{0.00}	23.44 _{0.00}	39.84 _{0.00}	49.41 _{0.00}	69.53 _{0.00}	44.23 _{0.00}	43.55 _{0.00}	60.55 _{0.00}	32.03 _{0.00}	46.68 _{0.00}	41.99 _{0.00}	39.25
		w/	5.00 _{0.50}	15.23 _{0.00}	40.90 _{15.99}	40.35 _{5.80}	50.82 _{0.42}	63.52 _{2.09}	55.98 _{1.93}	58.13 _{2.45}	63.87 _{0.96}	38.71 _{0.57}	46.41 _{0.16}	49.38 _{3.75}	44.02
	Test	w/o	5.86 _{0.00}	10.74 _{0.00}	22.27 _{0.00}	35.35 _{0.00}	48.63 _{0.00}	63.48 _{0.00}	46.88 _{0.00}	45.51 _{0.00}	64.84 _{0.00}	31.25 _{0.00}	37.50 _{0.00}	40.23 _{0.00}	37.71
		w/	6.52 _{0.61}	9.39 _{0.20}	41.89 _{11.21}	37.19 _{1.48}	47.95 _{0.57}	58.81 _{1.26}	53.89 _{1.52}	46.45 _{0.29}	59.39 _{3.25}	41.48 _{0.61}	37.21 _{0.22}	61.04 _{1.07}	41.77
GPT-2 (%)	Val.	w/o	4.70 _{0.00}	9.58 _{0.00}	24.11 _{0.00}	36.57 _{0.00}	47.91 _{0.00}	41.63 _{0.00}	34.16 _{0.00}	26.03 _{0.00}	54.25 _{0.00}	22.91 _{0.00}	46.68 _{0.00}	29.85 _{0.00}	31.53
		w/	5.25 _{0.57}	9.46 _{0.01}	27.57 _{5.24}	31.40 _{3.63}	50.31 _{0.61}	50.06 _{1.78}	51.44 _{2.19}	33.29 _{0.25}	61.13 _{2.01}	27.20 _{0.36}	38.58 _{0.10}	41.68 _{2.24}	35.61
	Test	w/o	6.08 _{0.00}	6.50 _{0.00}	20.45 _{0.00}	34.84 _{0.00}	46.00 _{0.00}	38.83 _{0.00}	35.37 _{0.00}	28.04 _{0.00}	64.67 _{0.00}	22.73 _{0.00}	32.56 _{0.00}	29.22 _{0.00}	30.44
		w/	5.68 _{1.01}	5.78 _{0.11}	27.89 _{5.96}	27.69 _{3.78}	46.66 _{0.47}	49.68 _{0.71}	48.89 _{3.86}	29.66 _{0.32}	58.07 _{4.31}	29.69 _{0.45}	32.17 _{0.22}	54.16 _{2.63}	34.67
ECE ₁ (%, ↓)	Val.	w/o	53.30 _{0.00}	46.81 _{0.00}	33.37 _{0.00}	13.58 _{0.00}	18.57 _{0.00}	22.00 _{0.00}	42.42 _{0.00}	13.50 _{0.00}	6.37 _{0.00}	20.26 _{0.00}	5.92 _{0.00}	49.98 _{0.00}	27.17
		w/	58.20 _{2.84}	48.35 _{0.11}	17.67 _{9.90}	14.21 _{4.88}	17.31 _{0.55}	22.03 _{1.67}	17.86 _{2.16}	17.87 _{2.35}	5.43 _{1.65}	14.39 _{0.67}	6.90 _{0.26}	28.80 _{1.95}	22.42
	Test	w/o	52.61 _{0.00}	51.74 _{0.00}	32.92 _{0.00}	14.46 _{0.00}	19.68 _{0.00}	27.22 _{0.00}	39.90 _{0.00}	10.44 _{0.00}	8.01 _{0.00}	22.48 _{0.00}	12.02 _{0.00}	49.70 _{0.00}	28.43
		w/	57.49 _{2.32}	54.24 _{0.30}	17.41 _{10.77}	20.08 _{5.50}	20.17 _{0.46}	25.34 _{0.77}	21.52 _{3.83}	8.81 _{0.48}	3.94 _{2.29}	11.62 _{0.67}	12.48 _{0.25}	14.84 _{2.55}	22.33

In terms of implementation, we use the MulticlassCalibrationError module given by the TorchMetrics³ with n_bins=10, norm='l1'.

G Complete Results: Reliability Diagrams and Confident Distribution

Reliability diagrams: as shown in Fig. 9 - 14 for GPT-J, and Fig. 21 - 24 for GPT-2.

Confidence distributions: as shown in Fig. 15 - 20 for GPT-J, and Fig. 25 - 28 for GPT-2.

H Case Analysis: NOISYICL Furtherance Correct ICL

Moreover, we find that in some cases, unperturbed ICL can't benefit correctly from scaling the number of demos, while, NOISYICL can help the model correct this issue, as shown in Fig. 5. These unperturbed models exhibit an overfitting-like phenomenon and also low accuracies, while NOISYICL can relieve it.

We speculate the reason is the mismatch between the pre-trained knowledge and ICL inputs. This leads to a decrease in the model's in-context task learning (Pan et al., 2023) ability, while NOISYICL reduces such a gap between pre-trained data and ICL style data, which makes models extract information from ICL inputs better.

³lightning.ai/docs/torchmetrics/stable/classification/calibration_error.html#multiclasscalibrationerror

I License for Artifacts

Here we discuss the license of the artifacts used in this paper.

Models. GPT-2 is under the MIT license, and GPT-J is under the apache-2.0 license.

Datasets. We list the open-source license for the datasets used in this paper as follows:

- cc-by-4.0: PS, SE' 14R, SE' 14L, TEE, TES, TEH
- cc-by-sa-3.0: HS, SST2, RTE, MRPC, FP
- agpl-3.0: Ethos

Consistency of Usage. Models are used with their original usage. Due to the different data splitting of these datasets, to ensure the consistency of the experiment methods, we use a re-splitting method as described in Appendix A. However, the overall usage is consistent with their intended use.

Table 4: Accuracy and Macro-F1 results (% , $mean_{std}$, $k = 1$). A better result is in **green**. Notation is the same with the Tabel 1.

Dataset			PS	HS	SE'14R	SE'14L	RTE	MRPC	Ethos	FP	SST2	TEE	TES	TEH	Mean
λ			0.2	0.012	0.1	0.02	0.2	0.016	0.016	0.014	0.018	0.014	0.002	0.08	—
GPT-J	Acc. (%)	Val. w/o	39.43 _{0.99}	76.02 _{0.86}	57.85 _{1.45}	53.77 _{1.82}	50.41 _{1.15}	58.48 _{1.48}	50.32 _{1.81}	53.93 _{0.82}	68.20 _{0.52}	43.48 _{0.80}	43.83 _{2.45}	50.90 _{1.02}	52.22
		w/	53.24 _{4.76}	76.86 _{0.54}	64.28 _{2.15}	55.57 _{1.05}	51.76 _{1.55}	58.28 _{1.35}	51.00 _{1.77}	53.96 _{1.27}	71.56 _{1.07}	44.32 _{0.50}	46.37 _{1.09}	52.13 _{0.96}	56.61
	Test	w/o	35.08 _{1.74}	84.92 _{0.94}	53.66 _{1.12}	51.60 _{1.29}	50.52 _{1.31}	53.50 _{2.26}	50.40 _{1.37}	52.29 _{1.48}	73.89 _{0.84}	41.94 _{0.92}	44.80 _{1.26}	50.71 _{1.46}	53.61
		w/	46.31 _{3.18}	85.64 _{1.06}	58.26 _{1.49}	53.71 _{1.08}	51.84 _{0.51}	54.28 _{1.10}	50.88 _{2.16}	51.96 _{1.10}	73.85 _{1.67}	42.06 _{0.93}	45.23 _{1.75}	52.11 _{1.38}	55.51
	MF1 (%)	Val. w/o	26.01 _{0.70}	24.73 _{0.66}	50.78 _{1.61}	52.26 _{1.85}	50.36 _{1.18}	51.11 _{1.43}	49.04 _{1.83}	38.25 _{1.01}	67.39 _{0.58}	26.71 _{1.08}	38.75 _{2.41}	49.42 _{0.85}	43.73
		w/	23.18 _{1.34}	24.49 _{0.92}	52.49 _{2.89}	53.61 _{1.35}	48.28 _{4.78}	50.61 _{1.72}	49.88 _{1.81}	38.94 _{1.99}	71.02 _{1.16}	27.65 _{1.31}	41.43 _{1.61}	50.75 _{0.97}	44.36
	Test	w/o	25.30 _{1.20}	24.73 _{0.71}	47.91 _{1.24}	50.62 _{1.25}	50.17 _{1.19}	49.97 _{2.33}	49.81 _{1.33}	37.90 _{1.98}	72.25 _{1.04}	24.67 _{1.43}	40.10 _{1.28}	48.45 _{1.61}	43.49
		w/	24.76 _{1.72}	25.13 _{1.02}	47.53 _{2.14}	52.18 _{0.89}	47.10 _{4.38}	50.65 _{1.07}	50.36 _{2.14}	38.93 _{1.27}	72.22 _{1.96}	24.97 _{0.92}	40.45 _{1.73}	49.94 _{1.40}	52.42
	ECE_1 (% , $\downarrow$)	Val. w/o	23.09 _{1.25}	22.67 _{0.82}	11.25 _{1.73}	17.01 _{1.66}	49.33 _{1.14}	35.31 _{1.39}	47.42 _{1.82}	27.54 _{0.97}	5.33 _{1.23}	44.33 _{1.12}	28.71 _{2.16}	45.64 _{1.19}	29.80
		w/	16.64 _{1.79}	21.81 _{0.59}	13.33 _{2.62}	15.57 _{0.64}	26.82 _{5.71}	35.56 _{1.89}	46.63 _{1.82}	27.39 _{1.11}	3.72 _{1.19}	42.88 _{0.69}	26.36 _{1.40}	43.95 _{1.10}	26.72
	Test	w/o	28.18 _{1.95}	13.97 _{0.95}	12.05 _{1.52}	15.95 _{1.05}	49.20 _{1.34}	39.94 _{2.33}	47.35 _{1.40}	28.24 _{1.42}	4.60 _{0.95}	45.34 _{0.90}	27.63 _{0.85}	46.98 _{1.37}	29.95
		w/	20.85 _{5.28}	13.21 _{1.07}	15.45 _{1.84}	14.28 _{0.93}	25.64 _{3.69}	35.56 _{1.89}	46.76 _{2.10}	28.13 _{0.97}	5.53 _{1.44}	44.71 _{0.93}	27.08 _{1.93}	45.28 _{1.33}	26.87
λ			0.02	0.002	0.016	0.004	0.002	0.002	0.2	0.06	0.02	0.02	0.06	0.1	—
GPT-2	Acc. (%)	Val. w/o	42.42 _{0.97}	41.37 _{1.11}	45.68 _{1.09}	39.69 _{1.69}	50.62 _{1.35}	68.71 _{0.20}	47.26 _{0.72}	45.45 _{1.81}	57.79 _{1.23}	28.11 _{0.82}	41.45 _{1.36}	42.70 _{0.32}	45.94
		w/	45.55 _{3.18}	37.34 _{1.33}	48.75 _{1.80}	41.66 _{1.04}	50.18 _{2.11}	68.81 _{0.18}	53.35 _{2.99}	55.33 _{2.72}	59.26 _{1.38}	32.73 _{0.53}	45.25 _{1.34}	48.41 _{2.15}	48.88
	Test	w/o	35.07 _{1.44}	39.64 _{0.61}	43.24 _{1.43}	38.47 _{0.97}	50.22 _{1.27}	62.94 _{0.42}	48.65 _{0.59}	45.09 _{1.51}	59.76 _{1.67}	27.61 _{1.20}	37.98 _{1.24}	44.37 _{0.81}	44.42
		w/	40.08 _{2.91}	34.92 _{1.12}	44.36 _{0.95}	39.41 _{1.30}	50.58 _{2.31}	62.47 _{0.29}	51.69 _{2.95}	52.12 _{1.60}	60.64 _{1.78}	30.55 _{0.79}	43.88 _{1.53}	55.35 _{0.79}	47.17
MF1 (%)	Val.	w/o	26.44 _{0.73}	20.32 _{1.44}	32.96 _{0.78}	33.69 _{1.92}	50.59 _{1.35}	41.67 _{0.39}	43.18 _{0.82}	30.19 _{1.51}	56.71 _{1.25}	22.59 _{0.81}	36.97 _{1.57}	32.23 _{0.53}	35.63
		w	26.37 _{1.41}	18.95 _{1.29}	35.97 _{1.24}	35.61 _{1.25}	50.13 _{2.10}	42.28 _{0.48}	47.26 _{2.81}	34.47 _{0.77}	59.22 _{1.39}	22.87 _{1.07}	30.49 _{1.97}	51.68 _{2.03}	37.94
	Test	w/o	23.97 _{1.23}	17.45 _{0.85}	33.27 _{1.93}	36.32 _{1.03}	49.85 _{1.26}	39.41 _{0.66}	43.80 _{0.77}	30.92 _{1.38}	59.63 _{1.61}	22.32 _{1.38}	36.21 _{1.30}	38.79 _{1.06}	36.00
		w/	26.31 _{1.58}	15.87 _{1.14}	33.78 _{1.33}	37.11 _{1.53}	50.18 _{2.40}	39.10 _{0.41}	44.93 _{3.66}	34.58 _{1.06}	59.52 _{1.82}	21.69 _{1.08}	33.40 _{1.36}	50.57 _{1.50}	37.25
ECE_1 (% , $\downarrow$)	Val.	w/o	14.45 _{1.49}	25.96 _{1.14}	16.68 _{1.33}	25.55 _{1.99}	42.92 _{1.41}	19.22 _{0.49}	35.57 _{1.01}	21.11 _{1.76}	5.40 _{1.16}	31.62 _{1.17}	20.72 _{1.70}	46.15 _{0.38}	25.44
		w/	9.64 _{1.86}	30.37 _{1.36}	17.27 _{2.01}	23.65 _{0.78}	43.31 _{2.03}	17.87 _{0.49}	19.53 _{3.89}	14.80 _{0.62}	3.94 _{0.41}	30.66 _{0.71}	24.11 _{2.54}	26.44 _{2.47}	21.80
	Test	w/o	18.82 _{1.59}	26.53 _{0.96}	18.64 _{1.41}	24.56 _{0.96}	43.56 _{1.28}	24.09 _{0.58}	35.15 _{0.76}	22.13 _{1.42}	2.83 _{1.04}	32.31 _{1.34}	22.45 _{1.38}	39.96 _{0.77}	25.92
		w/	13.20 _{2.03}	32.02 _{1.47}	20.45 _{1.07}	23.83 _{1.22}	43.07 _{2.30}	23.01 _{0.87}	23.65 _{6.52}	14.48 _{1.15}	3.22 _{1.50}	33.43 _{1.60}	19.37 _{1.64}	23.89 _{1.50}	22.80

Table 5: Accuracy and Macro-F1 results (% , $mean_{std}$, $k = 2$). A better result is in **green**. Notation is the same with the Tabel 1.

Dataset			PS	HS	SE'14R	SE'14L	RTE	MRPC	Ethos	FP	SST2	TEE	TES	TEH	Mean	
λ			0.2	0.2	0.1	0.1	0.08	0.2	0.014	0.08	0.016	0.018	0.02	0.008	—	
GPT-J	Acc. (%)	Val. w/o	42.54 _{1.49}	72.03 _{0.66}	41.50 _{0.97}	40.88 _{0.49}	51.31 _{1.19}	55.74 _{1.24}	49.81 _{2.28}	61.84 _{0.60}	72.05 _{0.96}	44.47 _{0.62}	48.85 _{0.75}	49.08 _{1.16}	52.51	
		w/	57.11 _{4.74}	77.34 _{3.29}	58.57 _{5.39}	51.25 _{2.71}	50.76 _{2.09}	58.98 _{6.13}	51.37 _{1.15}	60.47 _{1.77}	73.38 _{0.70}	44.67 _{0.63}	49.14 _{0.78}	50.62 _{2.11}	56.97	
		Test w/o	39.23 _{0.98}	82.42 _{1.21}	36.69 _{1.30}	39.26 _{1.67}	48.13 _{2.60}	50.97 _{1.43}	51.41 _{1.93}	60.77 _{0.64}	72.60 _{1.66}	44.28 _{0.73}	48.79 _{1.00}	51.68 _{1.95}	52.19	
		w/	50.95 _{4.52}	82.99 _{5.51}	53.91 _{3.42}	49.30 _{3.09}	49.45 _{1.30}	53.43 _{5.23}	51.46 _{1.62}	59.51 _{1.22}	72.83 _{1.36}	44.21 _{0.53}	48.77 _{1.20}	51.02 _{1.84}	55.65	
	MF1 (%)	Val. w/o	25.17 _{1.10}	25.03 _{0.33}	37.42 _{1.13}	39.11 _{0.83}	50.44 _{1.23}	51.35 _{1.45}	49.78 _{2.28}	31.83 _{1.31}	71.32 _{1.07}	25.43 _{1.03}	36.69 _{1.16}	49.05 _{1.17}	41.05	
		w/	22.75 _{1.84}	24.15 _{0.56}	50.06 _{5.11}	49.26 _{2.50}	50.47 _{2.22}	48.95 _{0.99}	51.36 _{1.15}	35.39 _{3.22}	72.77 _{0.78}	25.16 _{0.84}	37.72 _{1.56}	50.56 _{2.12}	43.22	
		Test w/o	25.51 _{1.25}	24.95 _{0.85}	34.99 _{1.38}	39.75 _{1.63}	47.82 _{2.64}	49.56 _{1.48}	51.18 _{1.96}	30.66 _{0.98}	70.34 _{1.87}	24.05 _{0.97}	36.78 _{1.28}	51.49 _{1.96}	40.59	
		w/	22.98 _{1.69}	25.10 _{0.79}	47.51 _{2.96}	48.50 _{3.12}	49.35 _{1.25}	46.51 _{1.56}	51.19 _{1.62}	32.85 _{2.38}	70.54 _{1.76}	24.05 _{0.67}	36.63 _{0.53}	50.87 _{1.80}	42.17	
	ECE_1 (% , $\downarrow$)	Val. w/o	15.11 _{1.29}	21.52 _{0.58}	23.45 _{1.46}	27.07 _{0.89}	39.26 _{0.83}	23.82 _{1.59}	38.51 _{2.24}	22.54 _{0.77}	3.56 _{0.89}	44.33 _{0.66}	19.22 _{0.62}	40.33 _{1.30}	26.56	
		w/	13.69 _{3.31}	11.25 _{5.10}	8.03 _{4.49}	15.54 _{3.79}	39.53 _{1.65}	12.12 _{2.82}	37.12 _{0.98}	21.10 _{1.70}	4.80 _{0.74}	43.87 _{0.96}	18.64 _{1.68}	38.66 _{2.22}	22.03	
		Test w/o	18.59 _{1.01}	13.15 _{1.15}	27.87 _{1.35}	26.71 _{1.83}	41.77 _{2.91}	28.29 _{1.47}	36.62 _{1.76}	22.97 _{0.87}	4.77 _{0.84}	44.53 _{0.98}	19.92 _{1.25}	36.63 _{1.86}	26.82	
		w/	13.49 _{2.88}	6.52 _{2.25}	11.44 _{3.91}	15.80 _{2.76}	41.33 _{2.22}	20.24 _{5.56}	36.61 _{1.67}	21.88 _{1.75}	5.32 _{1.16}	44.40 _{0.77}	19.98 _{0.70}	37.48 _{1.92}	22.87	
GPT-2	λ			0.02	0.1	0.08	0.08	0.006	0.002	0.2	0.04	0.02	0.02	0.04	0.2	—
	Acc. (%)	Val. w/o	41.86 _{1.78}	48.24 _{1.37}	46.99 _{0.84}	42.42 _{1.10}	50.88 _{1.05}	68.55 _{0.16}	45.85 _{1.27}	49.10 _{1.15}	56.56 _{1.31}	29.06 _{0.71}	39.39 _{1.51}	42.81 _{0.63}	46.81	
		w/	47.23 _{3.35}	48.73 _{6.40}	52.85 _{5.89}	46.62 _{1.59}	51.80 _{1.48}	68.89 _{4.43}	52.84 _{2.43}	61.31 _{0.81}	59.45 _{1.20}	30.29 _{0.82}	45.70 _{1.46}	51.09 _{4.07}	51.40	
		Test w/o	38.26 _{1.37}	52.04 _{0.91}	42.52 _{1.15}	39.71 _{1.47}	49.67 _{0.80}	62.77 _{0.39}	48.26 _{1.00}	47.42 _{1.45}	58.96 _{1.07}	28.47 _{0.97}	32.11 _{1.37}	44.74 _{0.82}	45.41	
		w/	40.95 _{1.84}	45.52 _{3.99}	50.12 _{3.58}	40.31 _{0.90}	50.64 _{1.87}	57.34 _{0.41}	51.76 _{2.27}	59.41 _{1.11}	61.33 _{1.51}	28.49 _{1.04}	40.29 _{1.18}	53.01 _{3.31}	48.26	
	MF1 (%)	Val. w/o	25.20 _{1.70}	21.61 _{1.20}	36.33 _{1.10}	36.50 _{0.51}	50.38 _{1.03}	41.38 _{0.39}	41.44 _{1.75}	33.91 _{1.20}	55.30 _{1.30}	20.40 _{1.31}	35.73 _{1.63}	33.12 _{1.17}	35.94	
		w/	25.40 _{0.74}	21.77 _{1.89}	32.89 _{2.37}	28.58 _{2.09}	51.21 _{1.53}	42.70 _{0.82}	48.94 _{1.40}	31.66 _{2.16}	59.17 _{1.11}	19.12 _{0.69}	30.66 _{1.33}	46.38 _{1.14}	36.54	
		Test w/o	24.91 _{1.27}	20.54 _{0.43}	35.57 _{0.99}	38.02 _{1.35}	48.42 _{0.93}	39.80 _{0.37}	42.83 _{1.35}	33.29 _{1.18}	58.84 _{1.10}	20.96 _{1.15}	31.99 _{1.36}	40.79 _{1.13}	36.33	
		w/	25.63 _{1.20}	20.78 _{0.73}	31.81 _{1.94}	28.97 _{2.76}	49.44 _{1.83}	44.06 _{1.11}	45.99 _{0.55}	33.03 _{1.60}	59.12 _{1.13}	18.40 _{1.10}	33.02 _{1.48}	50.92 _{1.60}	36.64	
	ECE_1 (% , $\downarrow$)	Val. w/o	10.21 _{1.71}	25.07 _{1.21}	14.93 _{1.00}	21.00 _{1.10}	35.37 _{1.19}	20.19 _{0.40}	41.29 _{1.34}	15.87 _{1.07}	6.43 _{1.22}	30.94 _{0.72}	18.02 _{1.74}	46.90 _{0.60}	23.85	
		w/	5.92 _{1.26}	29.23 _{4.69}	12.70 _{1.65}	22.24 _{1.86}	34.69 _{2.33}	19.18 _{0.61}	21.02 _{2.77}	10.53 _{2.79}	3.97 _{1.86}	35.02 _{0.56}	16.78 _{1.71}	27.18 _{1.94}	19.87	
		Test w/o	12.48 _{1.44}	19.52 _{0.68}	17.95 _{1.41}	20.35 _{1.46}	36.73 _{0.99}	23.63 _{0.44}	39.93 _{1.27}	17.40 _{1.25}	4.08 _{1.05}	31.48 _{1.10}	24.73 _{1.44}	42.57 _{1.44}	24.24	
w/		9.58 _{1.56}	33.12 _{2.08}	16.31 _{1.69}	28.04 _{3.60}	36.06 _{1.95}	17.86 _{0.86}	24.58 _{0.04}	8.99 _{2.05}	2.09 _{0.84}	37.68 _{1.94}	15.45 _{1.96}	21.12 _{2.08}	20.91		

Table 6: Accuracy and Macro-F1 results (% , $mean_{std}$, $k = 8$). A better result is in **green**. Notation is the same with the Table 1. Some experiments can't be conducted due to the length of the input sequence.

Dataset		PS	HS	SE'14R	SE'14L	RTE	MRPC	Ethos	FP	SST2	TEE	TES	TEH	Mean	
λ		0.2	0.2	0.1	0.1	0.1	0.2	0.004	0.012	0.002	0.006	0.08	0.002	—	
GPT-J	Acc. (%)	w/o	27.32 _{0.84}	68.87 _{1.30}	35.61 _{0.57}	33.57 _{0.74}	50.12 _{0.29}	35.86 _{0.55}	57.29 _{1.69}	62.34 _{0.05}	83.96 _{0.86}	47.64 _{0.52}	51.13 _{0.66}	50.62 _{1.62}	50.37
		w/	62.01 _{4.95}	85.29 _{0.95}	45.84 _{11.13}	40.84 _{6.41}	50.49 _{0.89}	58.87 _{7.93}	57.95 _{1.35}	62.44 _{0.08}	83.89 _{0.79}	47.81 _{0.90}	51.35 _{1.67}	52.05 _{0.75}	58.24
		w/o	25.70 _{0.84}	77.10 _{1.75}	33.30 _{0.63}	30.80 _{0.71}	46.35 _{1.39}	39.35 _{0.84}	61.67 _{1.38}	61.76 _{0.06}	80.59 _{1.02}	46.37 _{0.50}	50.68 _{0.35}	54.57 _{1.41}	50.69
		w/	58.54 _{3.27}	89.74 _{3.97}	41.68 _{6.66}	35.66 _{4.65}	47.53 _{1.26}	54.57 _{8.48}	61.46 _{0.96}	61.78 _{0.08}	79.78 _{0.69}	46.42 _{0.63}	50.62 _{0.41}	54.68 _{0.92}	56.87
	MF1 (%)	w/o	16.07 _{0.47}	23.66 _{0.32}	30.15 _{1.14}	24.48 _{1.40}	45.66 _{0.49}	33.37 _{0.56}	56.68 _{1.81}	25.91 _{0.26}	83.80 _{0.91}	27.20 _{0.59}	33.72 _{1.47}	48.52 _{1.98}	37.44
		w/	20.91 _{1.79}	27.76 _{1.35}	41.70 _{10.30}	35.28 _{8.07}	45.89 _{3.44}	46.69 _{2.90}	57.45 _{1.37}	26.24 _{0.36}	83.72 _{0.82}	27.39 _{0.84}	36.84 _{5.55}	49.94 _{0.99}	41.65
		w/o	16.01 _{0.94}	26.26 _{1.29}	28.72 _{0.67}	27.21 _{0.78}	42.58 _{1.37}	33.58 _{0.98}	60.88 _{1.44}	25.60 _{0.25}	78.98 _{1.18}	24.50 _{0.70}	33.19 _{0.40}	54.36 _{1.42}	37.66
		w/	20.81 _{1.17}	24.29 _{0.56}	39.15 _{6.23}	34.28 _{5.55}	42.47 _{4.45}	42.45 _{4.20}	60.52 _{1.02}	25.67 _{0.31}	78.06 _{0.79}	24.80 _{0.92}	33.85 _{3.11}	54.45 _{0.96}	40.07
ECE ₁ (% , ↓)	Val.	w/o	42.98 _{0.85}	10.43 _{1.20}	32.79 _{0.53}	42.43 _{0.99}	21.59 _{0.34}	33.19 _{0.72}	13.63 _{1.44}	24.54 _{0.13}	12.61 _{0.98}	40.10 _{0.65}	13.99 _{0.17}	21.11 _{1.51}	25.78
		w/	11.81 _{4.56}	9.51 _{2.14}	19.27 _{11.45}	25.96 _{10.94}	23.13 _{2.74}	8.10 _{2.84}	12.47 _{1.21}	24.33 _{0.69}	12.80 _{0.75}	39.95 _{1.00}	7.25 _{1.79}	20.01 _{0.66}	17.88
	Test	w/o	44.92 _{0.80}	7.93 _{1.41}	36.15 _{0.98}	43.80 _{1.01}	25.13 _{1.62}	29.83 _{0.74}	9.35 _{1.45}	25.00 _{0.23}	10.52 _{0.94}	41.83 _{0.54}	14.13 _{0.43}	16.78 _{1.45}	25.45
		w/	8.86 _{3.57}	6.04 _{1.94}	22.26 _{9.50}	29.44 _{8.54}	28.55 _{6.11}	15.28 _{7.58}	9.77 _{0.86}	24.92 _{0.52}	9.88 _{0.66}	41.83 _{0.80}	9.24 _{2.47}	16.77 _{0.95}	18.57

Table 7: Accuracy and Macro-F1 results (% , $mean_{std}$, $k = 16$). A better result is in **green**. Notation is the same with the Table 1. Some experiments can't be conducted due to the length of the input sequence.

Dataset			PS	HS	SE'14R	SE'14L	RTE	FP	SST2	TEE	TES	TEH	Mean	
λ			0.2	0.2	0.1	0.1	0.2	0.1	0.004	0.008	0.08	0.002	—	
GPT-J	Acc. (%)	Val.	w/o	18.52 _{0.78}	70.02 _{0.52}	40.86 _{0.90}	34.59 _{0.41}	48.91 _{0.62}	62.42 _{0.10}	88.24 _{0.70}	49.22 _{0.47}	51.84 _{0.97}	54.59 _{0.87}	51.92
			w/	61.78 _{5.94}	85.98 _{0.61}	46.84 _{12.92}	43.01 _{8.67}	50.76 _{1.15}	62.66 _{0.35}	88.83 _{0.71}	49.30 _{0.53}	53.11 _{1.99}	55.20 _{0.91}	59.75
		Test	w/o	18.75 _{0.38}	74.97 _{1.43}	34.98 _{1.10}	30.00 _{0.58}	45.95 _{0.45}	61.82 _{0.09}	81.88 _{1.00}	47.11 _{0.35}	51.70 _{0.42}	55.90 _{0.81}	50.31
			w/	59.97 _{2.40}	89.96 _{4.15}	42.56 _{6.98}	38.55 _{7.65}	52.96 _{1.39}	61.92 _{0.49}	82.20 _{0.86}	47.06 _{0.65}	52.13 _{0.88}	56.29 _{1.21}	58.36
	MF1 (%)	Val.	w/o	10.33 _{0.46}	30.44 _{0.84}	32.79 _{0.78}	24.15 _{0.66}	38.73 _{0.71}	26.08 _{0.46}	88.22 _{0.70}	28.75 _{0.40}	36.01 _{1.76}	53.61 _{0.88}	36.91
			w/	20.00 _{0.71}	23.54 _{0.57}	41.11 _{11.38}	35.07 _{10.38}	42.85 _{7.50}	27.18 _{1.53}	88.81 _{0.72}	29.04 _{0.55}	40.22 _{6.60}	54.02 _{1.02}	40.18
		Test	w/o	9.61 _{0.26}	27.64 _{1.69}	28.80 _{0.97}	25.05 _{0.94}	35.57 _{0.69}	25.80 _{0.31}	80.33 _{1.12}	25.50 _{0.54}	35.44 _{0.90}	55.76 _{0.81}	34.95
			w/	19.70 _{1.28}	23.96 _{0.08}	37.93 _{9.63}	36.06 _{8.97}	41.17 _{5.54}	26.27 _{1.96}	80.61 _{1.12}	25.45 _{0.95}	37.20 _{3.58}	56.13 _{1.25}	38.45
ECE ₁ (%, ↓)	Val.	w/o	54.29 _{0.95}	4.01 _{0.77}	27.53 _{1.27}	40.90 _{0.52}	26.13 _{0.54}	26.34 _{0.14}	14.43 _{0.75}	38.97 _{0.58}	11.46 _{0.99}	13.66 _{0.72}	25.77	
		w/	11.75 _{5.60}	9.70 _{2.72}	18.56 _{10.81}	23.33 _{13.20}	21.01 _{13.24}	18.07 _{5.22}	15.49 _{0.58}	38.61 _{0.69}	6.12 _{2.85}	13.25 _{0.68}	17.59	
	Test	w/o	53.72 _{0.37}	4.51 _{0.87}	34.09 _{1.27}	43.06 _{0.62}	29.93 _{0.68}	26.90 _{0.19}	9.96 _{1.00}	40.84 _{0.45}	11.24 _{0.41}	13.45 _{0.86}	26.77	
		w/	8.33 _{4.47}	5.63 _{2.27}	20.29 _{12.35}	24.94 _{10.05}	17.53 _{7.30}	20.44 _{3.90}	9.99 _{0.89}	41.00 _{0.83}	5.19 _{2.52}	13.03 _{1.33}	16.64	

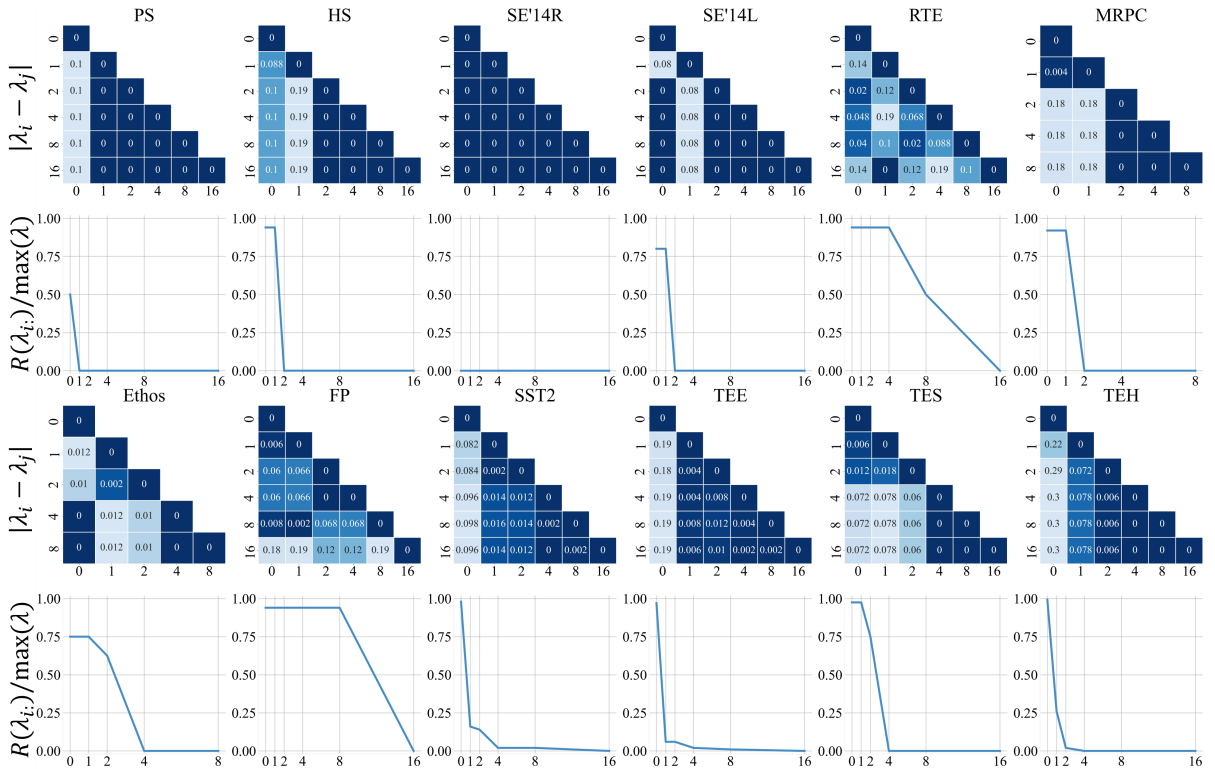


Figure 7: The stability of the optimal λ on various k for GPT-J. Upper figure: the distance between the optimal λ of different demo numbers; Lower figure: the normalized remaining range.

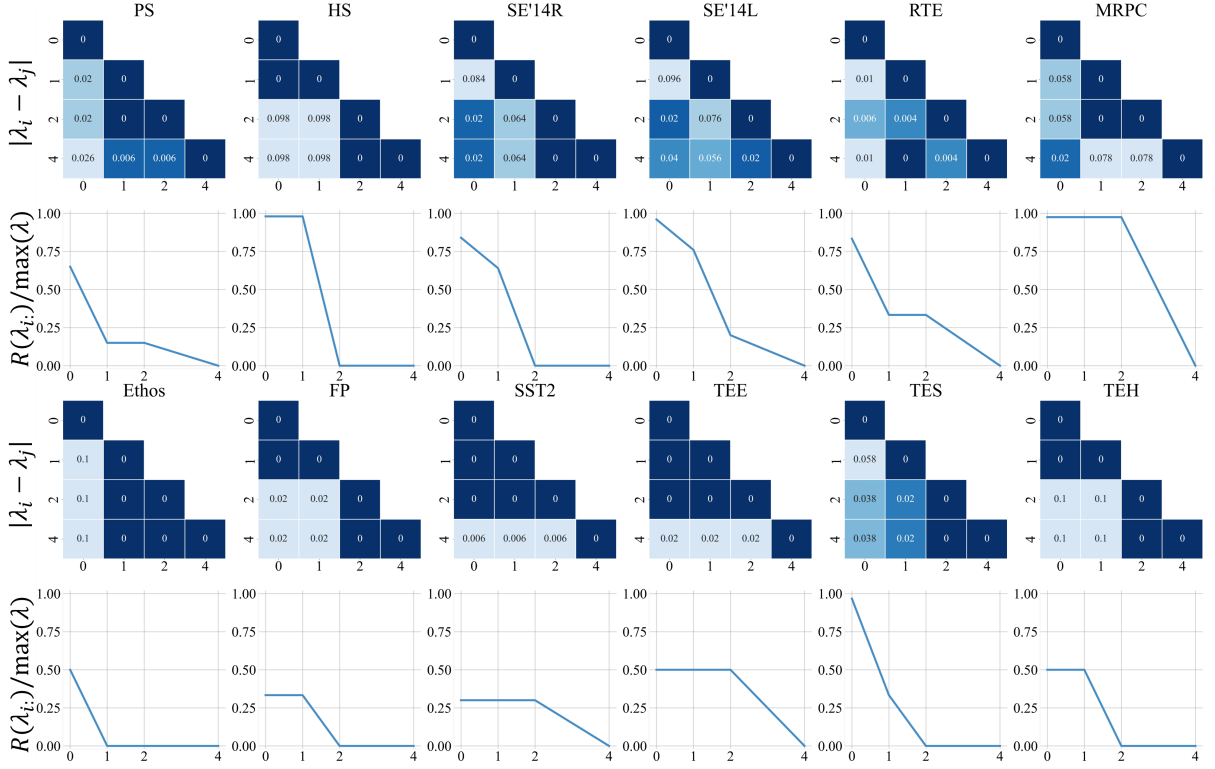


Figure 8: The stability of the optimal λ on various k for GPT-2. Upper figure: the distance between the optimal λ of different demo numbers; Lower figure: the normalized remaining range.

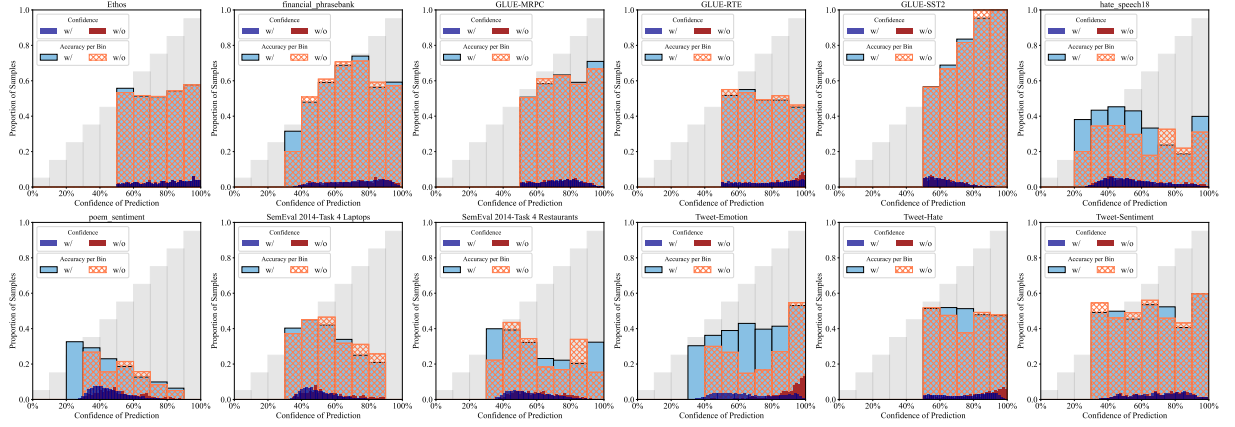


Figure 9: The reliability diagrams of GPT-J ($k = 0$).

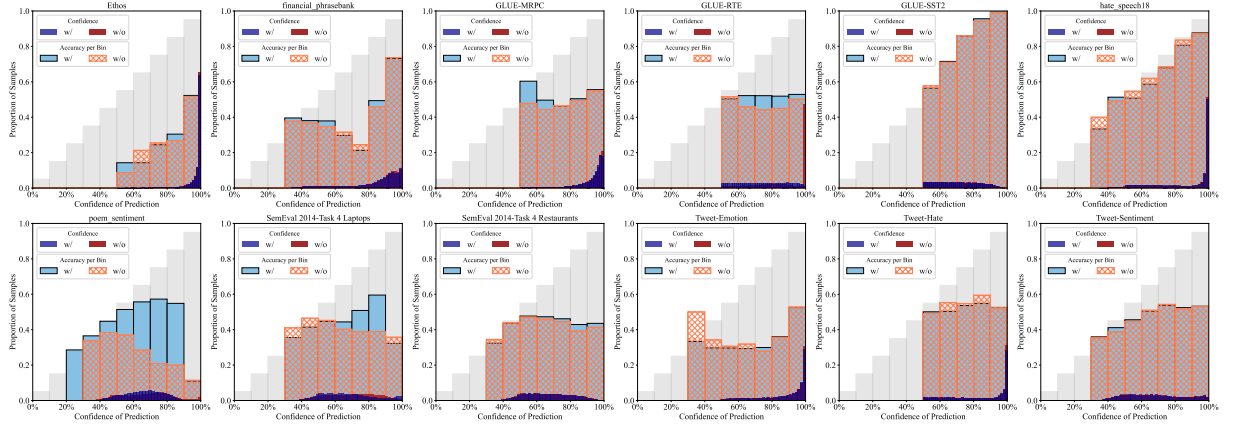


Figure 10: The reliability diagrams of GPT-J ($k = 1$).

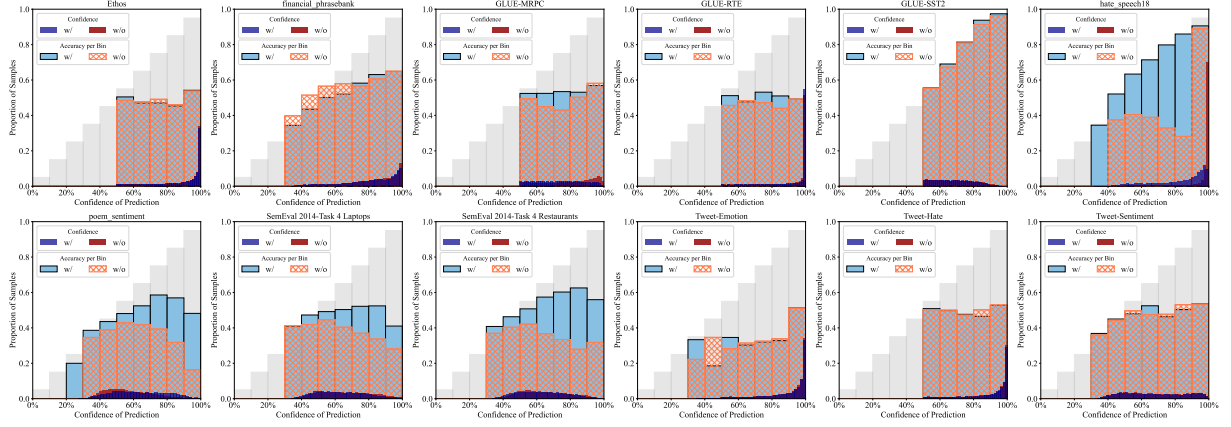


Figure 11: The reliability diagrams of GPT-J ($k = 2$).

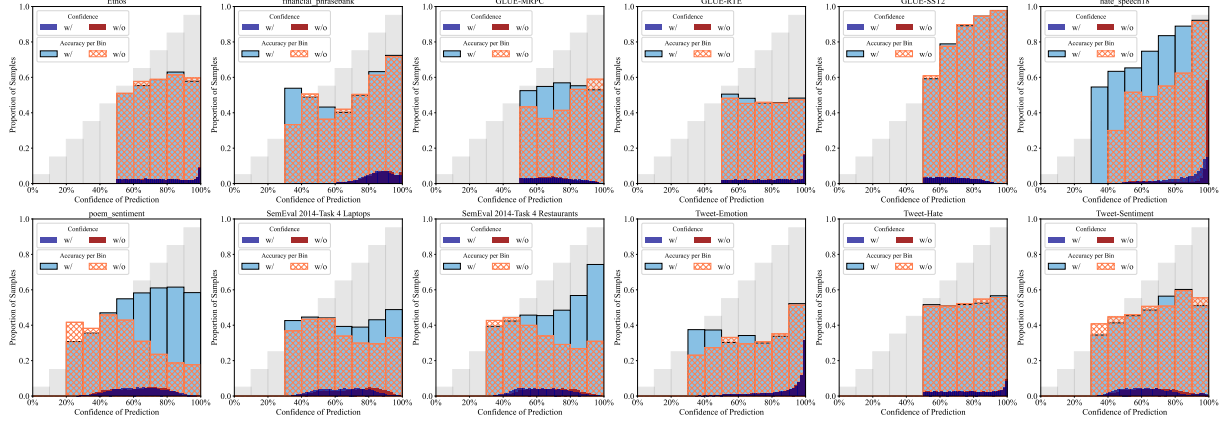


Figure 12: The reliability diagrams of GPT-J ($k = 4$).

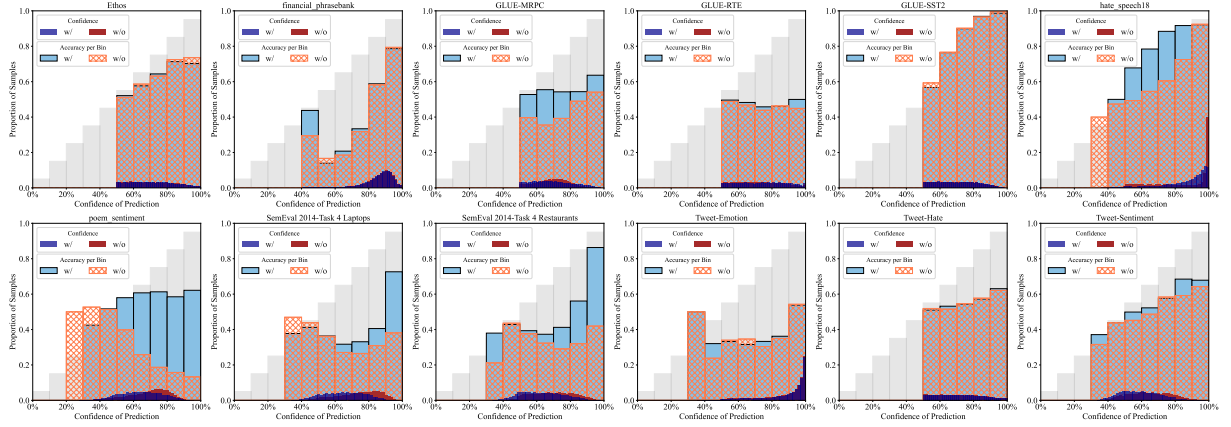


Figure 13: The reliability diagrams of GPT-J ($k = 8$).

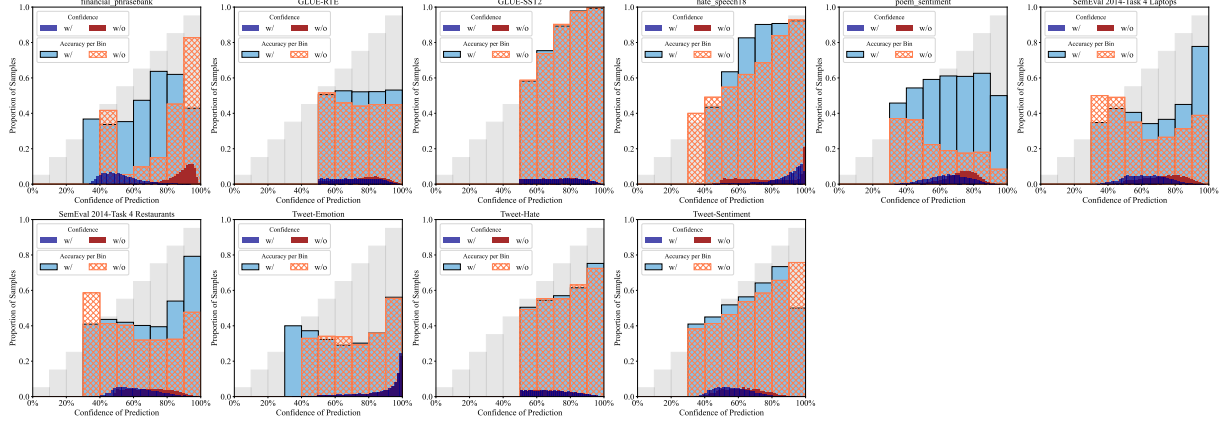


Figure 14: The reliability diagrams of GPT-J ($k = 16$).

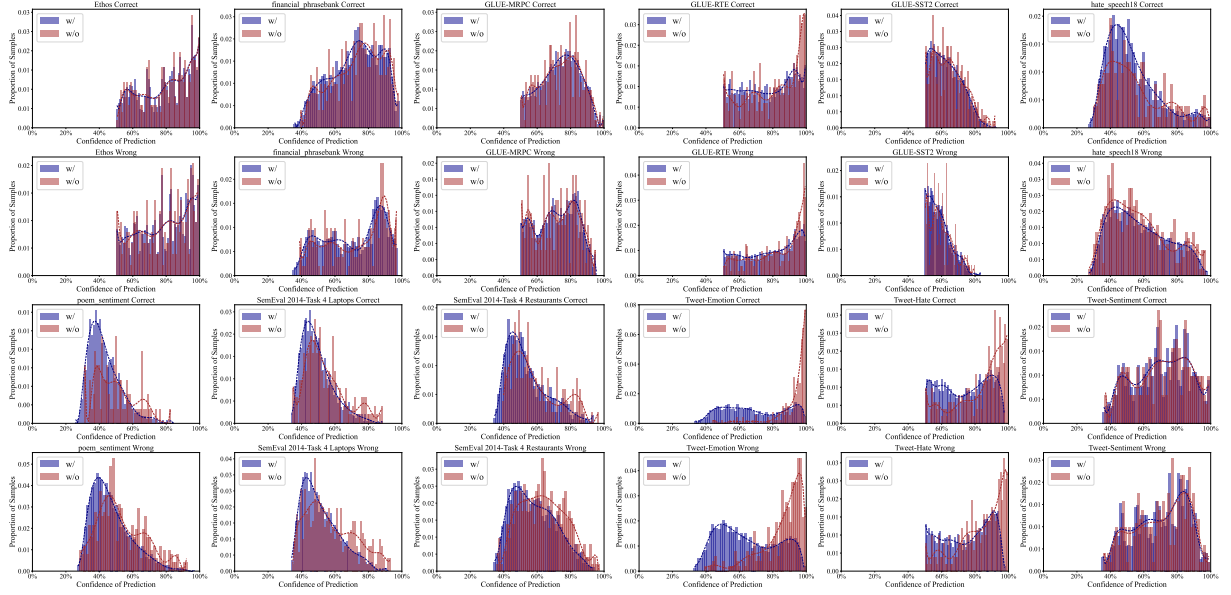


Figure 15: The confidence distribution of GPT-J ($k = 0$).

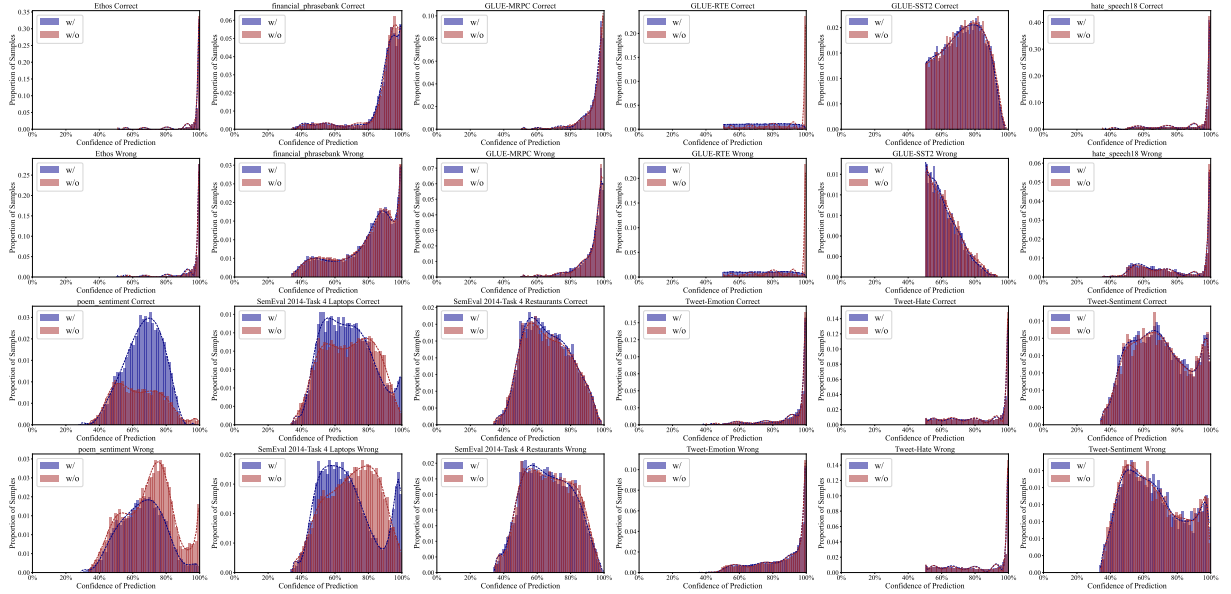


Figure 16: The confidence distribution of GPT-J ($k = 1$).

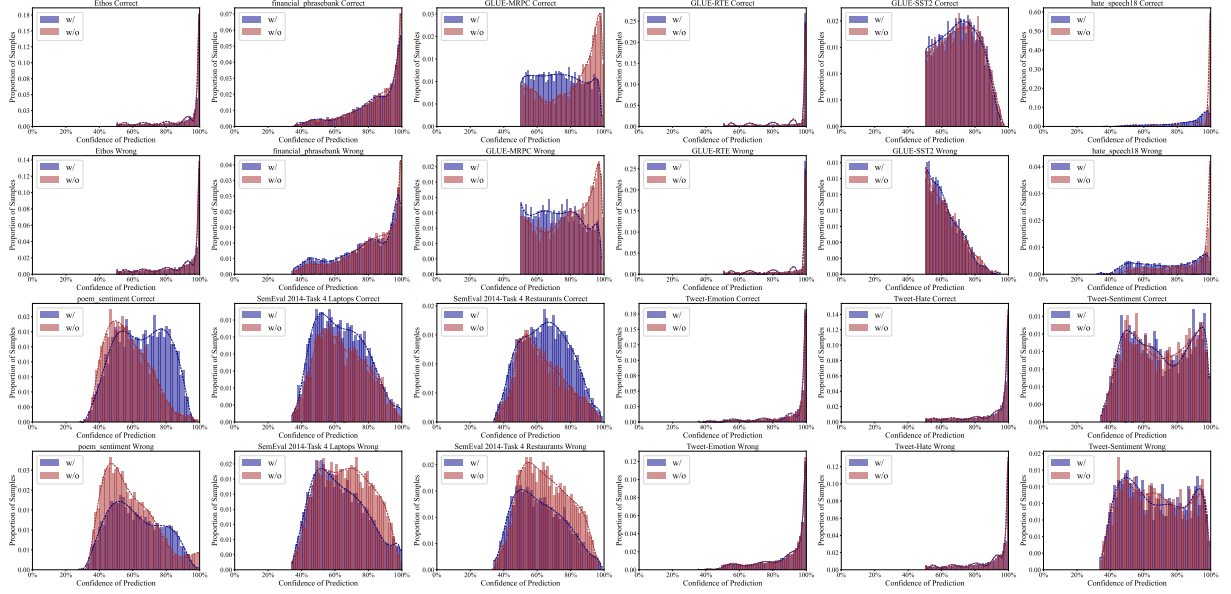


Figure 17: The confidence distribution of GPT-J ($k = 2$).

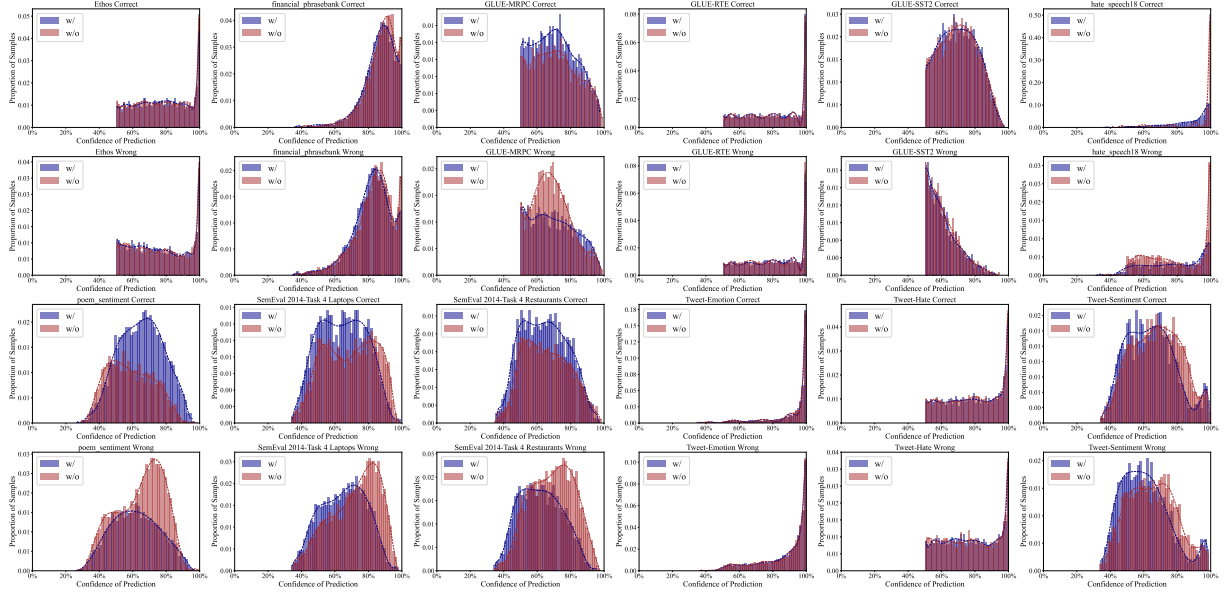


Figure 18: The confidence distribution of GPT-J ($k = 4$).

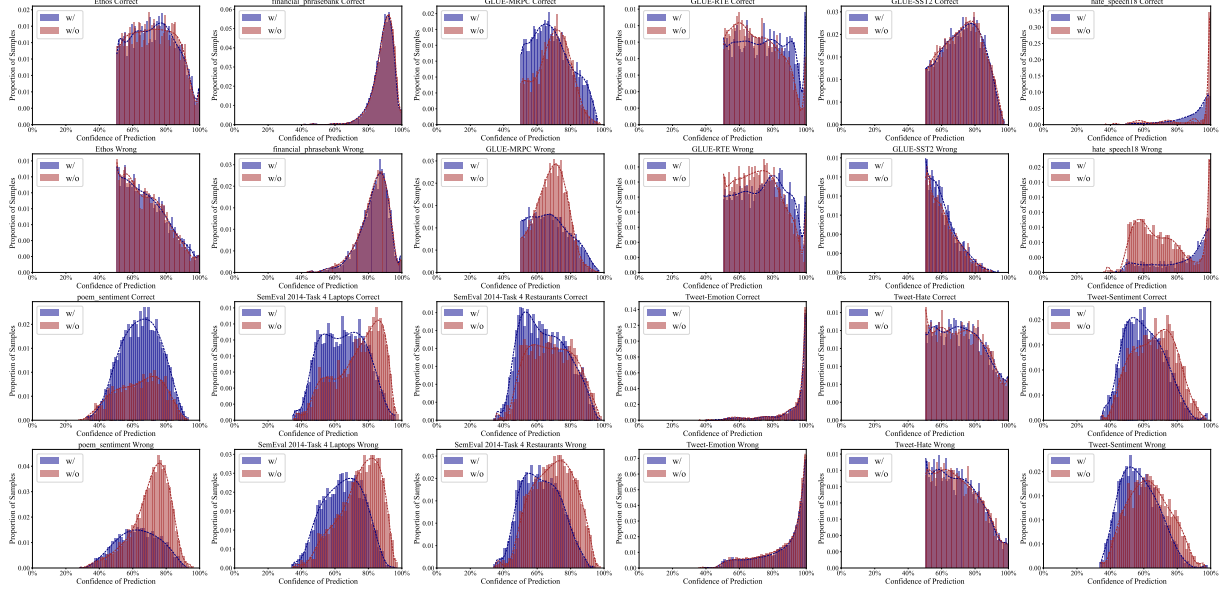


Figure 19: The confidence distribution of GPT-J ($k = 8$).

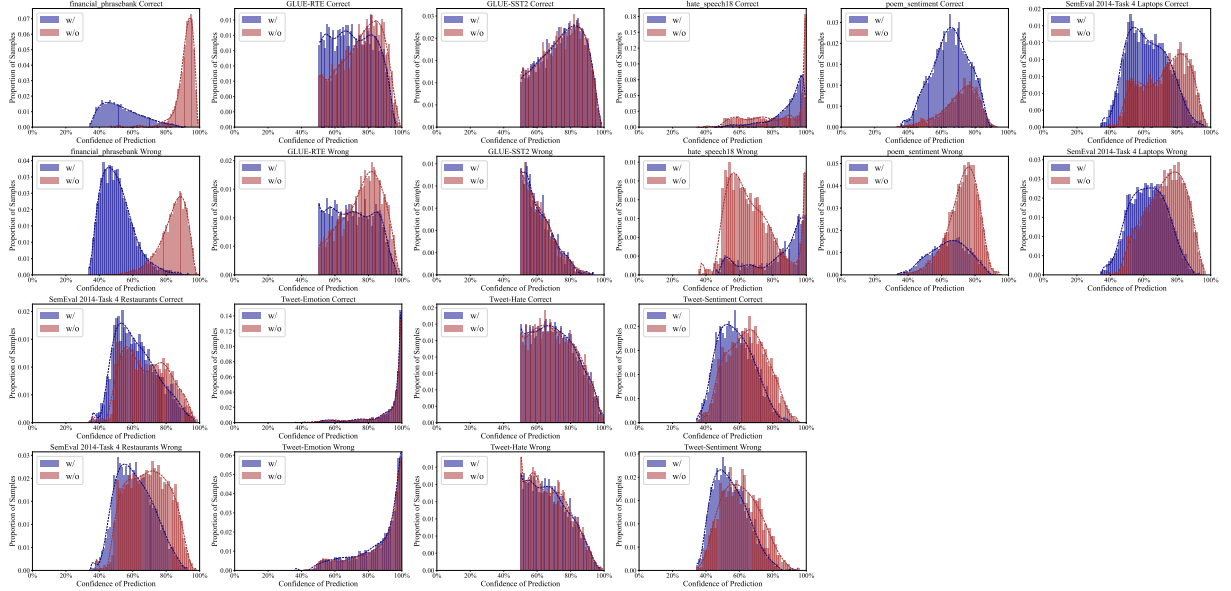


Figure 20: The confidence distribution of GPT-J ($k = 16$).

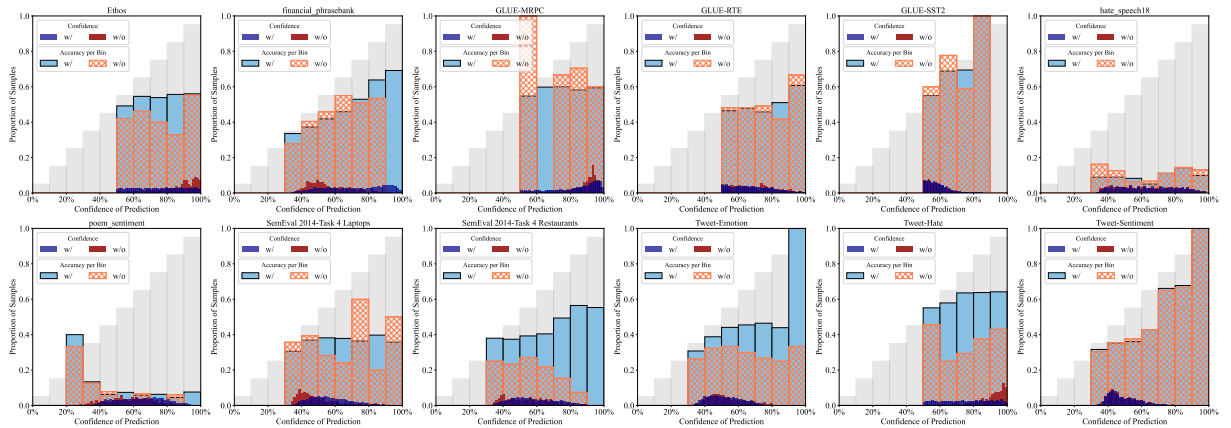


Figure 21: The reliability diagrams of GPT-2 ($k = 0$).

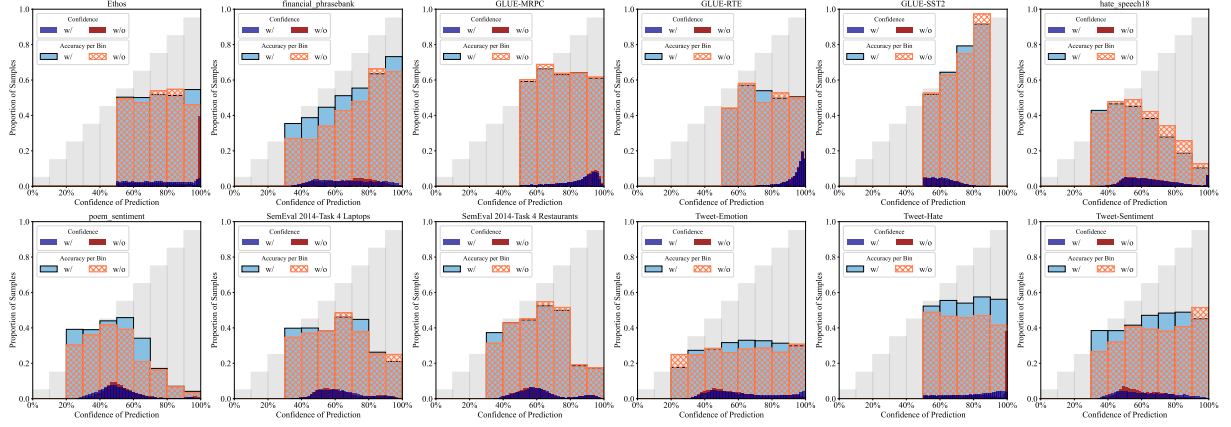


Figure 22: The reliability diagrams of GPT-2 ($k = 1$).

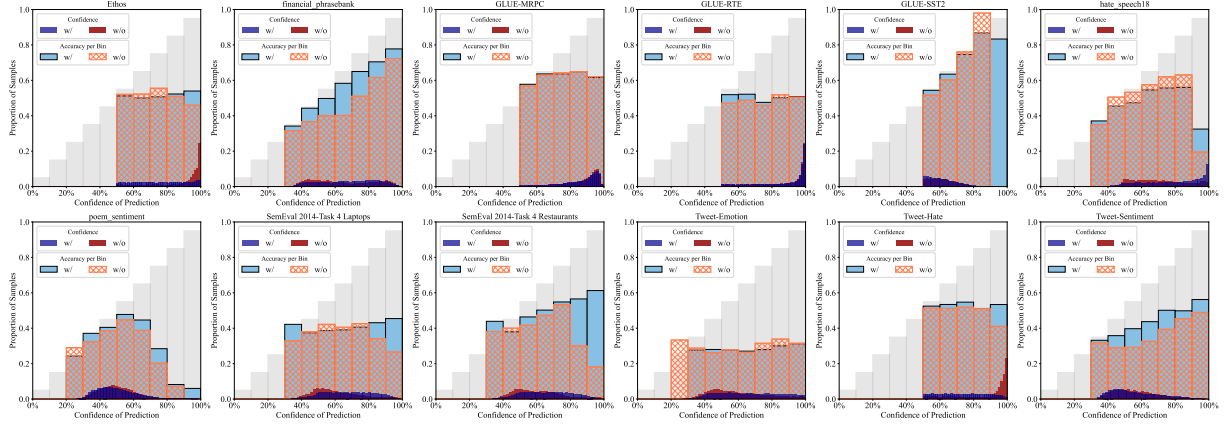


Figure 23: The reliability diagrams of GPT-2 ($k = 2$).

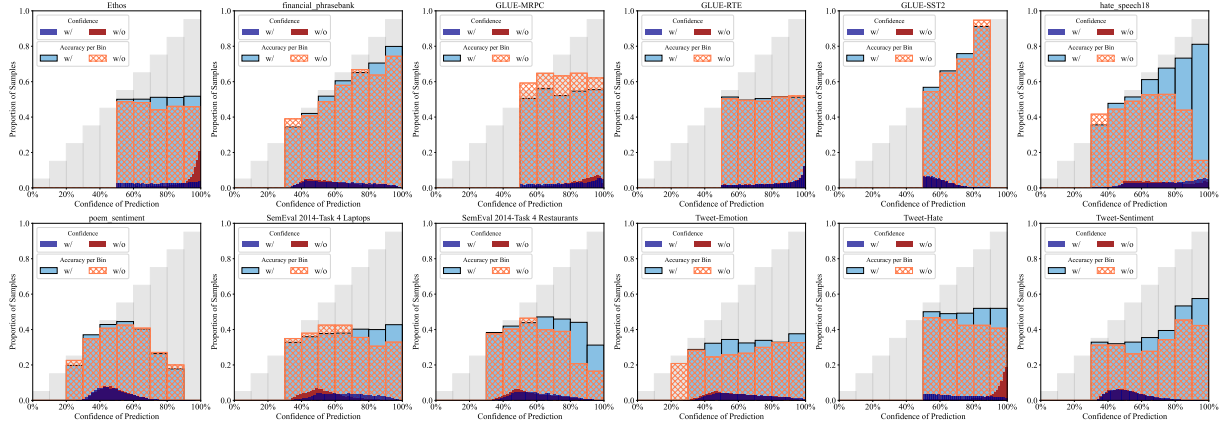


Figure 24: The reliability diagrams of GPT-2 ($k = 4$).

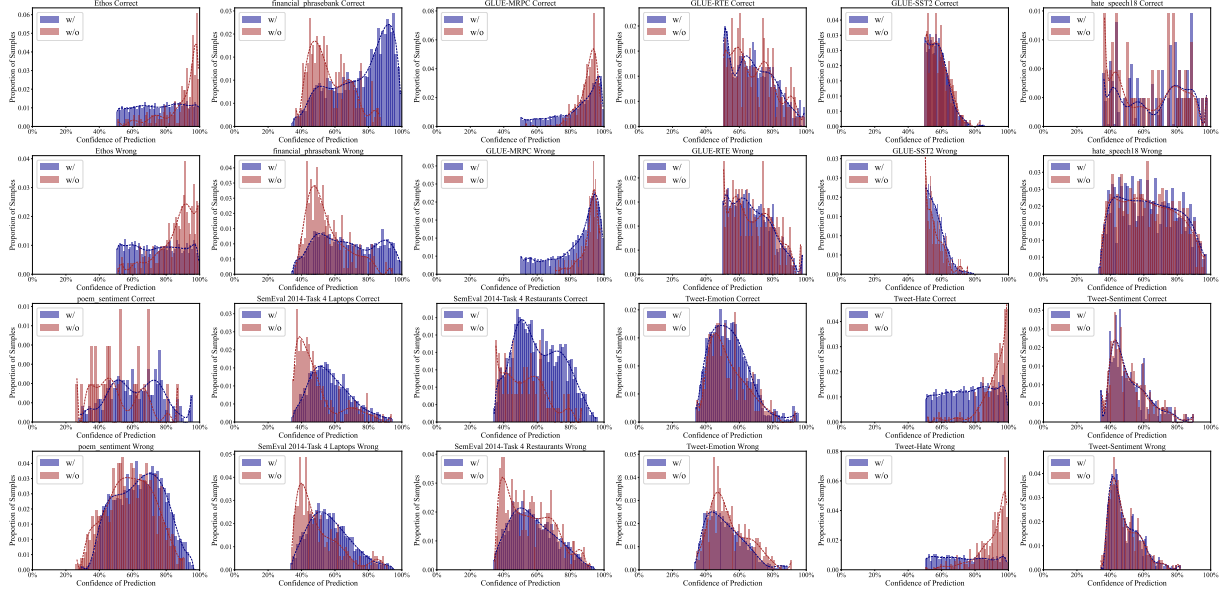


Figure 25: The confidence distribution of GPT-2 ($k = 0$).

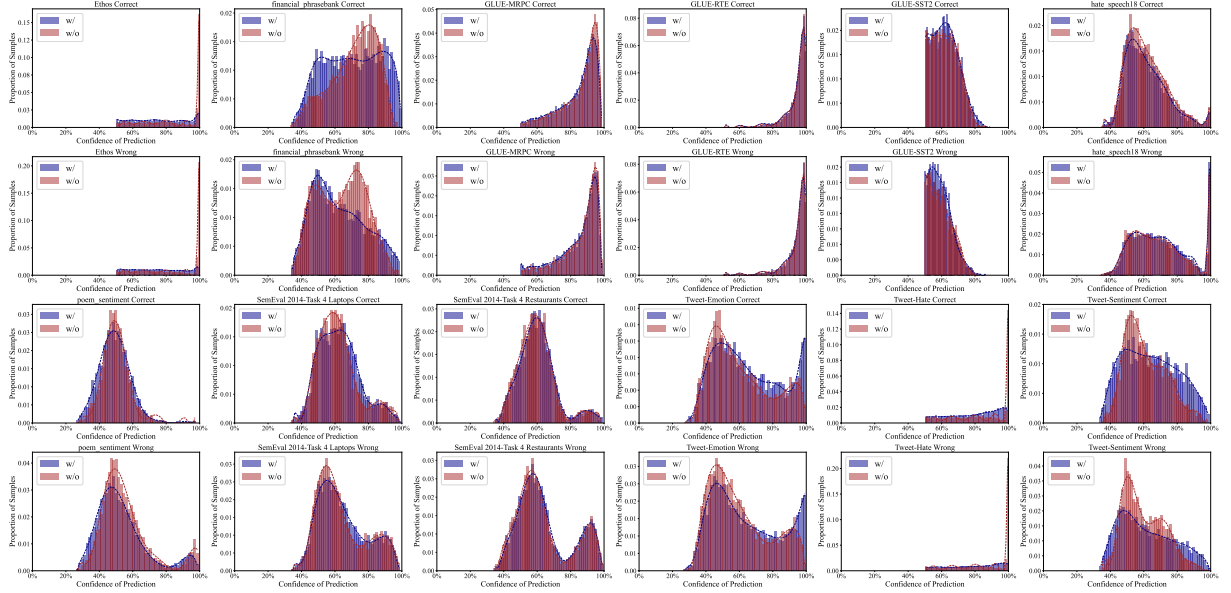


Figure 26: The confidence distribution of GPT-2 ($k = 1$).

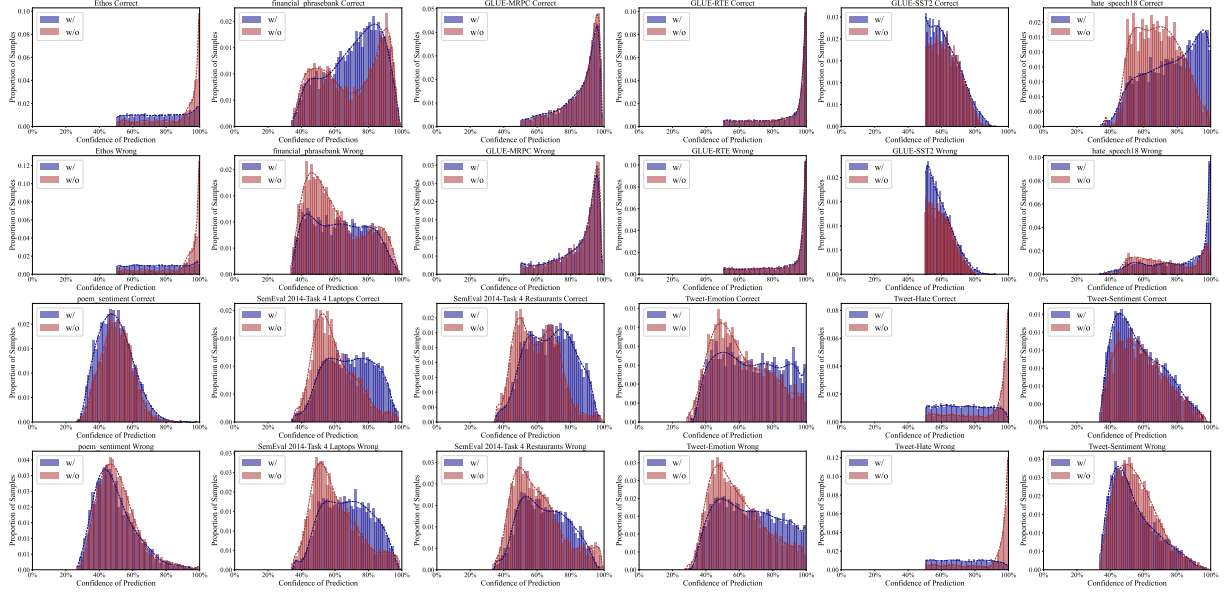


Figure 27: The confidence distribution of GPT-2 ($k = 2$).

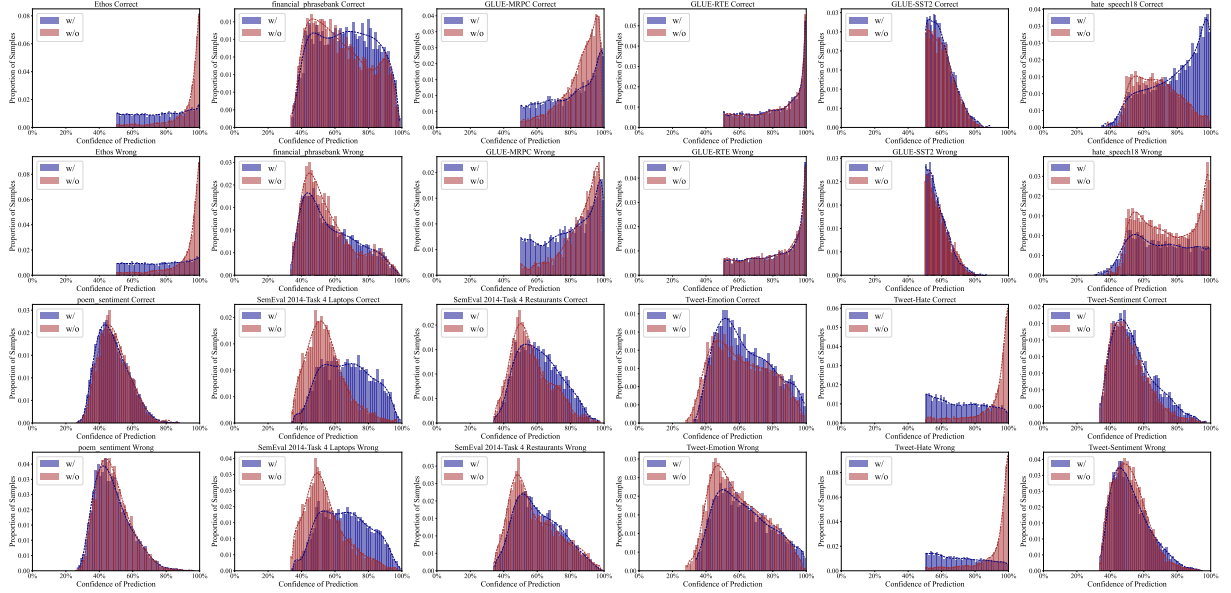


Figure 28: The confidence distribution of GPT-2 ($k = 4$).