

Labels Generated by Large Language Models Help Measure People's Empathy in Vitro

Md Rakibul Hasan[✉], *Graduate Student Member, IEEE*, Yue Yao[✉], *Member, IEEE*, Md Zakir Hossain[✉],
Aneesh Krishna[✉], *Member, IEEE*, Imre Rudas[✉], *Life Fellow, IEEE*, Shafin Rahman[✉],
and Tom Gedeon[✉], *Senior Member, IEEE*

Abstract—Large language models (LLMs) have revolutionised many fields, with LLM-as-a-service (LLMSaaS) offering accessible, general-purpose solutions without costly task-specific training. In contrast to the widely studied prompt engineering for directly solving tasks (*in vivo*), this paper explores LLMs' potential for *in-vitro* applications: using LLM-generated labels to improve supervised training of mainstream models. We examine two strategies – (1) noisy label correction and (2) training data augmentation – in empathy computing, an emerging task to predict psychology-based questionnaire outcomes from inputs like textual narratives. Crowdsourced datasets in this domain often suffer from noisy labels that misrepresent underlying empathy. We show that replacing or supplementing these crowdsourced labels with LLM-generated labels, developed using psychology-based scale-aware prompts, achieves statistically significant accuracy improvements. Notably, the RoBERTa pre-trained language model (PLM) trained with noise-reduced labels yields a state-of-the-art Pearson correlation coefficient of 0.648 on the public NewsEmp benchmarks. This paper further analyses evaluation metric selection and demographic biases to help guide the future development of more equitable empathy computing models.

Index Terms—Empathy detection, large language model, natural language processing, label noise, NewsEmp.

I. INTRODUCTION

LARGE language models (LLMs) have become a go-to approach across a variety of tasks, such as emotion recognition [1] and empathy detection [2], [3]. Due to high computational demands, coupled with environmental impact, training or fine-tuning LLMs often becomes costly. This limitation has led to increasing adoption of LLMs as a service (LLMSaaS), where users access trained LLMs via online APIs with computation on cloud [4]. LLMSaaS can be utilised *in-vivo*, i.e., prompt engineering to directly solve tasks such as named entity recognition [5], sentiment analysis [6] and empathy detection [2], [3], or *in-vitro* [7],¹ i.e., integrating LLM outputs into *other* models.

We are motivated by the following considerations. First, most current applications of LLMSaaS leverage LLM outputs *in-vivo* [2], [3], [5], [6]. We shift to their utility *in-vitro* to fine-tune smaller pre-trained language models (PLMs)² like RoBERTa [8]. In particular, we propose to utilise LLMSaaS in a *data-centric AI* approach [9] to (1) enhance the quality of training labels and to (2) increase the amount of quality training data for supervised training of PLMs.

Second, for representation learning, maintaining data quality is critical – captured succinctly by the phrase, “garbage in, garbage out” [10]. While deep learning research has mostly focused on proposing new algorithms, improvement in data-centric AI is equally important [9]. As a data-centric AI approach, we leverage LLMs to enhance data quality. The effectiveness of our proposed approach is demonstrated in an emerging field – empathy detection.

Empathy is defined as “*an affective response more appropriate to another's situation than one's own*” [11]. In psychology, various questionnaires have been developed to measure empathy. Empathy computing,³ in computer science, complements these psychology-based methods by aiming to map the questionnaire outcomes from input stimuli such as textual narratives, audiovisual interactions and physiological signals [12].

¹Like [7], we use the term “*in-vitro*” to refer to leveraging LLM outputs out of the box in a different model.

²We use “*pre-trained language models (PLMs)*” to refer specifically to smaller models like the BERT family of models, distinguishing them from LLMs, which are also pre-trained but significantly larger.

³We use the terms empathy computing, detection, prediction and measurement interchangeably.

Received 31 December 2024; revised 13 July 2025; accepted 23 February 2026. Date of publication 6 March 2026; date of current version 31 March 2026. This work was supported by Resources Provided by the Pawsey Supercomputing Research Centre through Australian Government and Government of Western Australia. The guest editor coordinating the review of this article and approving it for publication was Prof. Jiebo Luo. (*Corresponding author: Md Rakibul Hasan.*)

Md Rakibul Hasan is with the School of Electrical Engineering, Computing and Mathematical Sciences, Curtin University, Bentley, WA 6102, Australia, and also with BRAC University, Dhaka 1212, Bangladesh (e-mail: rakibul.hasan@curtin.edu.au).

Yue Yao was with the School of Electrical Engineering, Computing and Mathematical Sciences, Curtin University, Bentley, WA 6102, Australia. He is now with School of Airspace Science and Engineering, Shandong University, Weihai 264209, China (e-mail: yue.yao@anu.edu.au).

Md Zakir Hossain is with the School of Electrical Engineering, Computing and Mathematical Sciences, Curtin University, Bentley, WA 6102, Australia, and also with The Australian National University, Canberra, ACT 2600, Australia (e-mail: zakir.hossain1@curtin.edu.au).

Aneesh Krishna is with the School of Electrical Engineering, Computing and Mathematical Sciences, Curtin University, Bentley, WA 6102, Australia (e-mail: a.krishna@curtin.edu.au).

Imre Rudas is with Obuda University, 1034 Budapest, Hungary (e-mail: rudas@uni-obuda.hu).

Shafin Rahman is with North South University, Dhaka 1229, Bangladesh (e-mail: shafin.rahman@northsouth.edu).

Tom Gedeon is with the School of Electrical Engineering, Computing and Mathematical Sciences, Curtin University, Bentley, WA 6102, Australia, and also with The Australian National University, Canberra, ACT 2600, Australia, and also with Obuda University, 1034 Budapest, Hungary (e-mail: Tom.Gedeon@curtin.edu.au).

Code and LLM-generated labels are available at <https://github.com/hasan-rakibul/LLMPathy>.

Digital Object Identifier 10.1109/JSTSP.2026.3671186

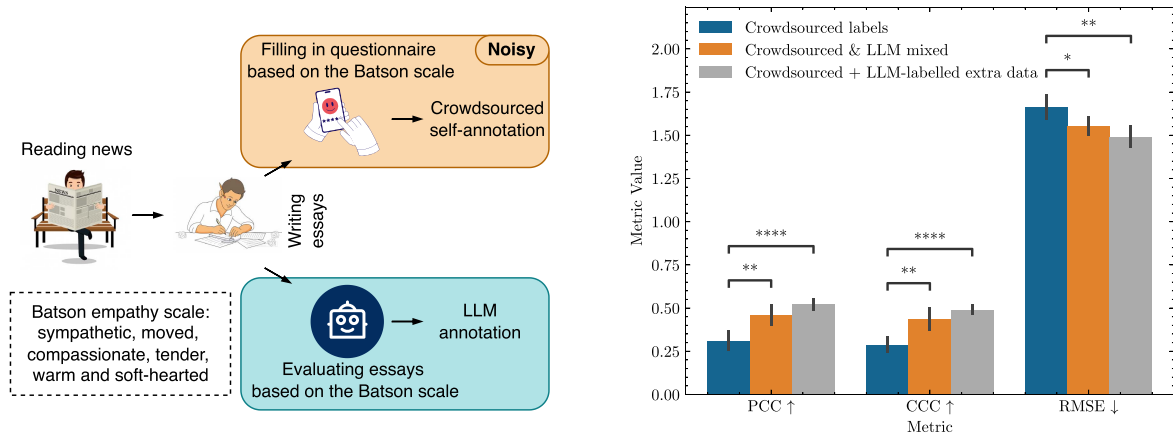


Fig. 1. Left: Overview of traditional and LLM-based methods to annotate essays for detecting empathy in the essays written in response to news articles. To be shown in our experiment, existing crowdsourced self-annotation through questionnaires is found to be incorrect in many samples. Our proposed approach involves annotating the essays using LLM, which is then used to reduce label noise and to get additional training data. Right: Impact of our LLM usage is showcased through a performance comparison between existing crowdsourced annotations, LLM-based label noise correction and the inclusion of additional data labelled by LLM. Statistical significance is calculated using Statannotations package [14], where * means $0.01 < p\text{-value} \leq 0.05$, ** means $0.001 < p\text{-value} \leq 0.01$ and **** means $p\text{-value} \leq 0.0001$ (i.e. more * means higher statistical significance).

One well-known questionnaire is the empathy measurement scale proposed by [13], which assesses empathy across six dimensions: sympathetic, moved, compassionate, tender, warm and soft-hearted.

Empathy *computing* offers the potential to improve people's empathic skills, which in turn strengthens interpersonal relationships across various human interactions [12]. In healthcare, for example, empathic writing in medical documents (e.g., patient reports) can promote understanding and trust between clinicians and patients [15]. Similarly, in education, written communication like emails and feedback on assignments has become a vital medium for expressing care and addressing students' emotional needs [16]. Journalism also demonstrates the importance of empathy in written narratives. For example, a news article on a family's recovery after a devastating event often goes beyond factual reporting and offers a compassionate perspective that engages readers emotionally and deepens their connection to the news story. Specifically, this paper measures people's empathy in essays written in response to scenarios reported in newspaper articles.

Empathy is inherently subjective, and machine learning models, including LLMs, used for its detection exhibit biases across different demographic groups [17]. We explore potential sources of such biases (Section IV-E). Additionally, while the Pearson correlation coefficient (PCC) remains the most commonly used evaluation metric in empathy computing [12], it does not account for the magnitude of the error. To address this, we advocate for adopting the concordance correlation coefficient (CCC) (Section IV-B).

Neural networks are prone to memorising training data, i.e., overfitting. This issue is exacerbated by noisy labels, where traditional regularisation techniques like dropout and weight decay often fall short [18]. A major challenge in ensuring data quality is, therefore, addressing label noise, defined as labels that deviate from their intended values. It is a significant challenge in empathy computing datasets collected through crowdsourcing. Platforms like Amazon Mechanical Turk offer quick access

to large participant pools. Accordingly, crowdsourcing with questionnaire-based self-assessment labelling is a popular way of collecting data in computational social psychology and human behaviour research, such as empathy [19] and emotion recognition [20]. However, such data often suffer from inaccuracies due to inattentiveness or multitasking among participants, compromising data reliability [21], [22], [23], [24]. This necessitates strategies to enhance data quality post-collection.

The overarching goal of this paper is to address the question: “How can LLMs enhance training of PLMs to improve empathy computing accuracy?” As illustrated in Fig. 1, our proposed *in-vitro* applications achieve statistically significant performance improvements. The first application, which automatically adjusts training labels, demonstrates consistent performance improvement across all metrics compared to the baseline PLM trained on the original dataset. The second application, leveraging additional LLM-labelled training data, further enhances model performance, yielding the highest statistically significant performance gains with a $p\text{-value} < 0.0001$ [14].

Our key contributions are summarised as follows:

- 1) We propose two *in-vitro* applications of LLMs: mitigating label noise and getting additional training data for PLMs.
- 2) We design a novel scale-aware prompt that enables LLMs to annotate data while adhering to annotation protocols grounded in theoretical frameworks.
- 3) We investigate challenges in empathy computing datasets and advocate for new evaluation metrics.
- 4) Our proposed methods achieve statistically significant performance improvements over the baseline models across multiple datasets and set a new state-of-the-art empathy computing performance.

II. RELATED WORK

A. LLM in Data Annotation

The advent of LLMs has inspired numerous studies exploring LLMs' application in data annotation, often positioning them as

a substitute for traditional human annotation. For instance, [25] examined the potential of LLMs in emotion annotation tasks and reported that LLMs can generate emotion labels closely aligned with human annotations. Similarly, [26] explored the utility of LLMs in annotating datasets for various natural language processing tasks, including sentiment analysis, question generation and topic classification. While they highlighted the cost-effectiveness of LLM-based annotations, they also noted LLMs' limitations compared to human annotators. Departing from this line of work, our approach explores LLM-generated labels to enhance the training of PLMs. Specifically, we integrate LLM-generated labels with human-generated labels rather than exclusive use of either LLM- or human-generated labels.

We examine our approach in two distinct applications: label adjustment and training data enhancement. Related to our first application, [27] also explored label noise adjustment, but that approach relies on subjects' demographic information (e.g., age, gender and race) in the prompting process. In contrast, our method deliberately avoids any use of demographic details to mitigate potential biases inherent in LLM training. Additionally, the reliance on demographic information may not always be feasible, which makes our approach more broadly applicable. Another key difference with [27]'s prompting strategy is the use of multiple input-output examples: while they rely on few-shot prompting with multiple example pairs to elicit LLM output in a consistent style, our approach does not require such examples, yet still achieves consistent outputs. Furthermore, they experimented solely on the GPT-3.5 LLM, whereas we explore both Llama 3 70B [28] and GPT-4 [29] LLMs in a recent dataset.

B. Learning With Label Noise

Noise-robust learning has been extensively studied in classification settings, especially for computer vision, leaving textual regression tasks under-explored [30]. Such learning algorithms were demonstrated either explicitly through dedicated methods [31], [32], [33], [34] or implicitly as part of broader semi-supervised learning frameworks [35], [36]. Dedicated methods such as [33], [34] address noise by a de-noising loss function based on the usual cross-entropy loss function in classification. Semi-supervised methods [35], [36] generate pseudo-labels based on class probabilities produced by `sigmoid` or `softmax` activations. Both categories of methods have shown strong performance in classification tasks; however, they are not directly applicable to regression, where the target space is continuous and lacks discrete output logits. Our *in-vitro* approach operates natively in the regression setting with the usual regression loss function and does not require any conversion to pseudo-class probabilities.

The study of label noise in textual regression remains limited. [30]'s approach iteratively identifies noisy examples and applies one of three strategies: discarding noisy data points, substituting noisy labels with pseudo-labels, or resampling clean instances to balance the dataset. While effective in identifying extreme outliers, [30] stated that their approach struggles in detecting mild disagreements. They further noted that their method performs worse in general-purpose datasets, compared

to knowledge-dense domains, such as clinical notes and academic papers. Our approach leverages LLMs as an external "teacher" to directly correct noisy labels in a *single* pass. Given the general-purpose nature of LLMs, our approach holds the potential to be effective across different domains.

C. Empathy Computing

Empathy computing is an emerging field, with significant advancements in textual empathy prediction [12]. For a detailed overview of its progress, we refer to a recent systematic literature review by [12].

In textual empathy computing, the most widely studied context is detecting people's empathy in response to newspaper articles. The Workshop on Computational Approaches to Subjectivity, Sentiment & Social Media Analysis (WASSA) Shared Tasks (2021–2024) have spurred various approaches leveraging PLMs on this task. Most approaches predominantly employed fine-tuning PLMs, with RoBERTa being the most preferred PLM [37], [38], [39], [40], [41], [42], [43], [44], [45], [46], [47], [48], [49]. Some studies have explored other BERT-based PLMs [50], [51], [52], [53] or ensemble strategies combining multiple PLMs [54], [55], [56]. The suitability of fine-tuning RoBERTa is further validated by [38], who reported that simple fine-tuning of RoBERTa outperformed more complex multi-task learning in textual empathy computing. Overall, fine-tuning PLMs has emerged as the predominant approach for this task, with RoBERTa being the leading model [12].

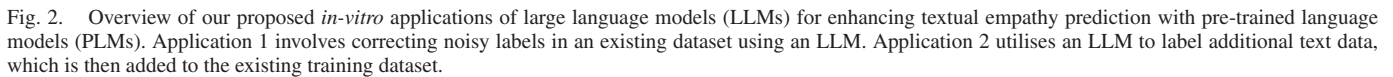
More recently, LLMs have been explored for textual empathy prediction through rephrasing text for data augmentation [27], [48], fine-tuning [2] and prompt engineering [3]. [27] adds multi-layer perception layers on top of a RoBERTa PLM to process demographic data, while [2]'s fine-tuning of LLM demands significant computational resources. Unlike these methods, our approach leverages LLM-generated labels to enhance fine-tuning of a standard RoBERTa PLM, without demographic data or high computational costs.

III. METHOD

A. Problem Formulation

Let $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$ represent a dataset, where x_i is the i -th input text sequence, and $y_i \in \mathbb{R}$ denotes corresponding continuous empathy score. Empathy, being a psychological construct, is challenging to annotate due to subjectivity. Consequently, the target variable y_i often suffers from noise, which is particularly significant in crowdsourced annotations (refer to Section IV-D1 for evidence of noise). We denote the noisy ground-truth empathy score as $\tilde{y}_i$, which serves as a proxy for the true, unobserved empathy score y_i . Thus, the dataset can be reformulated as $\mathcal{D} = \{(x_i, \tilde{y}_i)\}_{i=1}^N$. Our goal is to develop a model $\mathcal{F} : \mathcal{X} \rightarrow \mathbb{R}$, where $\mathcal{X}$ is the space of text sequences, such that $\mathcal{F}(x_i)$ accurately estimates y_i .

The dataset $\mathcal{D}$ is randomly partitioned into three non-overlapping subsets: a training set $\mathcal{D}_{\text{train}} = \{(x_i, \tilde{y}_i)\}_{i=1}^{N_{\text{train}}}$, used to train the models; a validation set $\mathcal{D}_{\text{val}} = \{(x_j, \tilde{y}_j)\}_{j=1}^{N_{\text{val}}}$, used for optimising the training and tuning hyperparameters; and a



Scheme	System prompt	User prompt
Plain	Your task is to measure the empathy of individuals based on their written essays. Human subjects wrote these essays after reading a newspaper article involving harm to individuals, groups of people, nature, etc. The essay is provided to you within triple backticks.	Essay: ```{essay}``` Now, provide empathy score between 1.0 and 7.0, where a score of 1.0 means the lowest empathy, and a score of 7.0 means the highest empathy. You must not provide any other outputs apart from the scores
Scale-aware	Your task is to measure the empathy of individuals based on their written essays. You will assess empathy using Batson's definition, which specifically measures how the subject is feeling each of the following six emotions: sympathetic, moved, compassionate, tender, warm and softhearted. Human subjects wrote these essays after reading a newspaper article involving harm to individuals, groups of people, nature, etc. The essay is provided to you within triple backticks.	Essay: ```{essay}``` Now, provide scores with respect to Batson's empathy scale. That is, provide scores between 1.0 and 7.0 for each of the following emotions: sympathetic, moved, compassionate, tender, warm and softhearted. You must provide comma-separated floating point scores, where a score of 1.0 means the individual is not feeling the emotion at all, and a score of 7.0 means the individual is extremely feeling the emotion. You must not provide any other outputs apart from the scores.

The motivation behind leveraging a scale-aware scheme includes natural alignment with the Psychology-grounded labelling protocol that was used for crowdsourced raters in the original study. For example, the NewsEmp datasets used Batson’s Empathy scale, which has six subscales, and so our scale-aware prompt instructs the LLM to provide scores on the subscales (Table I). The difference between the human- and LLM-generated labelling is minimal, with the only difference being who labels the data. Therefore, following the same protocol designed for crowdsourced raters, LLM outputs across

the subscales are averaged to calculate a single empathy score y^* . The following subsections present details about how we use these LLM labels in empathy computing.

2) *Application 1: Noise Mitigation in Labels*: The LLM-generated labels y^* are used to identify and replace potentially noisy samples in the original crowdsourced annotation $\tilde{y}$. Noisy samples are identified based on the difference between $\tilde{y}$ and y^* . Like [27], a revised label y'_i is defined as:

$$y'_i = \begin{cases} y_i^* & \text{if } |\tilde{y}_i - y_i^*| > \alpha \\ \tilde{y}_i & \text{otherwise,} \end{cases} \quad (1)$$

where α is a predefined threshold, referred to as the *annotation selection threshold*, which determines which label to use for which sample.

This selection threshold can be any real number between 0 and the range of the empathy score (i.e., $7 - 1 = 6$ for the NewsEmp dataset). A smaller α results in a more aggressive correction (replacing labels even for small differences). This could, however, lead to a larger distribution mismatch between the train and the test sets because the hold-out test set uses uncorrected crowdsourced labels. A model trained with a smaller α can, therefore, struggle to generalise on the hold-out test set.

Conversely, a larger α is a more conservative correction (replacing labels only at larger differences). Theoretically, a model trained with a larger α should generalise better on the test set because of a comparatively small distribution shift. This way, the model avoids training on crowdsourced labels that have a large deviation from LLM labels and, at the same time, enjoys the benefit of within-distribution crowdsourced labels that have slight deviations.

The revised dataset $\mathcal{D}'_{\text{train}} = \{(x_i, y'_i)\}$ is then used to train a PLM $\mathcal{F}_{y'}$. We hypothesise that the performance of $\mathcal{F}_{y'}$, trained on the mixture of $\tilde{y}$ and y^* , is better than $\mathcal{F}_{\tilde{y}}$, which is trained solely on $\tilde{y}$.

3) *Application 2: Additional Data Labelled by LLM*: Since deep learning models benefit from additional data, we propose to utilise LLM to get additional training data. While common LLM-based data augmentation techniques, such as paraphrasing [27], [45] and summarising [52], are well-documented in the literature, our approach goes a step further. Specifically, we use an LLM to label new essays following the same annotation protocol as our target domain. This method, therefore, enables the integration of any similar data points into the training process.

Mathematically, we prompt LLM to annotate new text samples u and make a new dataset $\mathcal{D}_{\text{llm}} = \{(u_i, v_i^*)\}_{i=1}^M$ with empathy scores v^* . These new data points could be any text similar to the essays x , but it may not have any prior empathy labels. We then annotate it in the same scale of y using LLM. This additional data is combined with $\mathcal{D}_{\text{train}}$ to create an extended training set:

$$\mathcal{D}_{\text{train}+} = \mathcal{D}_{\text{train}} \cup \mathcal{D}_{\text{llm}} = \{(x_i, \tilde{y}_i)\}_{i=1}^{N_{\text{train}}} \cup \{(u_i, v_i^*)\}_{i=1}^M. \quad (2)$$

Models trained on $\mathcal{D}_{\text{train}+}$ are expected to outperform models trained solely on $\mathcal{D}_{\text{train}}$ when evaluated on the hold-out test set $\mathcal{D}_{\text{test}}$. The extended dataset would enable the model to see a

more diverse set of training examples, which should improve the model's ability to generalise on unseen test data.

Similar to labelling training data required for our proposed applications, LLMs can be prompted directly to generate empathy labels for the test set $\mathcal{D}_{\text{test}}$. This zero-shot prediction leverages LLM's pre-trained knowledge without requiring further fine-tuning. Compared to the other two applications, zero-shot prediction relies heavily on the inherent capabilities of the LLM.

C. Prediction Using Pre-Trained Language Model

Fine-tuning a pre-trained language model (PLM) is a widely adopted approach in the empathy computing literature [12]. Accordingly, we utilise the dataset refined through LLM-based approaches to fine-tune a PLM. Each text sequence x_i is first encoded into a contextual representation that serves as an aggregate sequence representation:

$$h_i^{[\text{CLS}]} = \text{PLM}(x_i; \theta), \quad (3)$$

where θ are the parameters of the PLM, and $h_i^{[\text{CLS}]} \in \mathbb{R}^d$ is the [CLS] token representation. This pooled [CLS] representation is then passed through a linear regression head to predict the continuous empathy score:

$$\hat{y}_i = \mathcal{F}(h_i^{[\text{CLS}]}; \phi) = W h_i^{[\text{CLS}]} + b, \quad (4)$$

where $\phi = W \in \mathbb{R}^{1 \times d}, b \in \mathbb{R}$ denotes the learnable parameters of the linear layer. The model is trained to minimise the discrepancy between predicted scores $\hat{y}_i$ and target scores y_i^{true} across the dataset:

$$\mathcal{L} = \frac{1}{N} \sum_{i=1}^N \ell(\hat{y}_i, y_i^{\text{true}}), \quad (5)$$

where $\ell(\cdot, \cdot)$ is the mean squared error (MSE) loss function, and N is the number of training examples. The ground truth y_i^{true} refers to the mixed labels y' for Application 1, while for Application 2, it refers to crowdsourced labels y for existing dataset $\mathcal{D}_{\text{train}}$ or LLM-provided labels v^* for additional data $\mathcal{D}_{\text{llm}}$. The evaluation is always conducted on the original held-out dataset $\mathcal{D}_{\text{test}}$.

Algorithm 1 presents the overall workflow of our proposed approaches in empathy detection. After partitioning the dataset into training, validation and test subsets, one can choose between Application 1 and 2, as they are mutually exclusive. For Application 1 (label noise correction), the LLM is queried with scale-aware prompts to generate refined labels. If the difference between the original label and the LLM-generated label exceeds a threshold, the label is updated; otherwise, the original label is retained. The revised dataset is then used for PLM fine-tuning. For Application 2 (leveraging additional unlabelled data), the LLM is queried to generate labels for the unlabelled data, which is then combined with the training set to form an extended dataset. In both cases, a PLM is fine-tuned on the revised or extended dataset. The fine-tuning involves optimising the PLM to predict empathy scores based on input text embeddings, followed by evaluating its performance on the hold-out crowdsourced test set.

Algorithm 1: Leveraging LLM in Empathy Detection.

Require: Dataset $\mathcal{D} = \{(x_i, \tilde{y}_i)\}_{i=1}^N$, annotation selection threshold α , additional unlabelled data $\mathcal{U} = \{u_i\}_{i=1}^M$

Ensure: Empathy predictions $\hat{y}$

- 1: **Partition** $\mathcal{D}$ into $\mathcal{D}_{\text{train}}$, $\mathcal{D}_{\text{val}}$ and $\mathcal{D}_{\text{test}}$
- 2: **if** Application 1 **then**
- 3: **go to** 6
- 4: **else if** Application 2 **then**
- 5: **go to** 12
- 6: $\triangleright$ *Application 1: label noise correction* $\triangleleft$
- 7: **for** each i in $\mathcal{D}_{\text{train}}$ **do**
- 8: Query LLM to generate label y_i^* using scale-aware prompt
- 9: Update label: $y'_i \leftarrow \begin{cases} y_i^*, & \text{if } |\tilde{y}_i - y_i^*| > \alpha \\ \tilde{y}_i, & \text{otherwise} \end{cases}$
- 10: Form revised dataset $\mathcal{D}'_{\text{train}} = \{(x_i, y'_i)\}_{i=1}^{N_{\text{train}}}$
- 11: **go to** 18
- 12: $\triangleright$ *Application 2: additional data labelled by LLM* $\triangleleft$
- 13: **for** each u_i in $\mathcal{U}$ **do**
- 14: Query LLM to generate label v_i^* for u_i using the same prompt
- 15: Form additional dataset $\mathcal{D}_{\text{llm}} = \{(u_i, v_i^*)\}_{i=1}^M$
- 16: Combine datasets: $\mathcal{D}_{\text{train}+} = \mathcal{D}_{\text{train}} \cup \mathcal{D}_{\text{llm}}$
- 17: **go to** 18
- 18: $\triangleright$ *Prediction using pre-trained language model (PLM)* $\triangleleft$
- 19: Fine-tune PLM $\mathcal{F}_\theta$ on $\mathcal{D}'_{\text{train}}$ or $\mathcal{D}_{\text{train}+}$:
- $\hat{y}_i = \mathcal{F}_\theta(x_i) = Wh_i^{\text{[CLS]}} + b$
- 20: Evaluate $\mathcal{F}_\theta$ on $\mathcal{D}_{\text{test}}$
- 21: **return** final predictions $\hat{y}$

IV. EXPERIMENTS AND RESULTS

A. Dataset and Associated Challenges

[57] marked an important step in understanding how individuals empathise with others or nature. They designed a crowd-sourced approach where participants read newspaper articles depicting scenarios of harm to people or nature, and wrote about their emotional responses. The overarching aim was to capture individuals' reactions to adverse situations faced by others. This dataset, released in 2018, was the first of its kind, following which subsequent datasets were built. We refer to these datasets collectively as *NewsEmp* series, as their central objective is to measure empathy elicited by newspaper articles.

The second NewsEmp dataset was released in 2022, in which [19] employed 564 subjects reading 418 news articles, which led to a total of 2,655 samples distributed into training, validation and test splits. Another significant change appeared in the NewsEmp₂₃ dataset [58], which uses only the top 100 most negative articles from the pool of 418 news articles. Collectively, these datasets have become the most widely used dataset for benchmarking empathy detection approaches [12]. This popularity comes from their usage in the long-standing empathy detection challenge organised under the “Workshop on Computational Approaches to Subjectivity, Sentiment & Social

TABLE II
STATISTICS OF THE DATASETS USED IN THIS STUDY

Name	# Train	# Validation	# Test	# Total
NewsEmp ₂₂ [19]	1,860	270	525	2,655
NewsEmp ₂₃ [58]	792	208	100	1,100
NewsEmp ₂₄ [37]	1,000	63	83	1,146

Media Analysis (WASSA)” series [19], [37], [43], [59]. In particular, the WASSA 2021 [19] and WASSA 2022 [59] challenges utilised the NewsEmp₂₂ dataset, while WASSA 2023 [43] and the WASSA 2024 [37] utilised the NewsEmp₂₃ and the latest NewsEmp₂₄ datasets, respectively. Table II presents the statistics of the three datasets used in this study.

Perhaps due to the iterative nature of the datasets, there is overlap among some of these datasets. We found that the entire NewsEmp₁₈ dataset is included in the training set of NewsEmp₂₂, and the entire NewsEmp₂₃ training and validation sets appear in the NewsEmp₂₄ training set. In this study, we primarily use NewsEmp₂₂ and NewsEmp₂₄ datasets and partly NewsEmp₂₃ datasets.

Although excluding NewsEmp₂₃ would have been feasible, prior research [37] achieved state-of-the-art results on the NewsEmp₂₄ dataset by combining NewsEmp₂₂, NewsEmp₂₃ and NewsEmp₂₄ datasets to train their model. To ensure a fair comparison, we also report findings based on models trained using the combined three datasets.

While this combination may seem unusual due to the overlap, it can be beneficial for improving predictions on the NewsEmp₂₄ test split. Very likely, this test split has a similar distribution to its own training set compared to another dataset's (NewsEmp₂₂) training set, so including duplicated samples from the NewsEmp₂₄ training set allows the model to see more samples with a similar distribution.

It is worth noting that if we aim to evaluate a model on the NewsEmp₂₃ dataset, caution is necessary when combining datasets. One interesting finding on NewsEmp datasets is that – although not explicitly stated in the studies [37], [43], [58] reporting the datasets – 44 out of 100 test samples in the NewsEmp₂₃ dataset are also present in the NewsEmp₂₄ validation set. Due to this data leakage, a model trained on the NewsEmp₂₄ validation set would, therefore, inflate performance on the NewsEmp₂₃ test split. To verify this, we trained a model using NewsEmp₂₄ training and validation sets, which gives a PCC of 0.576, outperforming the state-of-the-art PCC of 0.563 in NewsEmp₂₃ test split [27]. To prevent misleading results in future research, we highlight this overlap here and recommend exercising caution when combining datasets. Throughout our experiments, we ensure that there is no data leakage between training/validation and test splits.

We compare the performance of our proposed LLM-based approaches across various dataset combinations, including NewsEmp₂₄, NewsEmp₂₃ and NewsEmp₂₂. We then benchmark our work against the evaluation metrics reported by others on the NewsEmp₂₄ dataset. This dataset was chosen because it is the most recent in this series, and it includes the NewsEmp₂₃ dataset within it. Additionally, the ground truth for

the NewsEmp₂₄ test split is publicly available, which is essential for calculating different metrics, while the ground truth for the other datasets is publicly unavailable.

B. Evaluation Metric

Pearson correlation coefficient (PCC) is the single metric used in the literature for evaluating empathy computing models across the NewsEmp datasets [12]. While it measures linear relationship between predicted and true values, it does not account for the *magnitude* of errors, meaning predictions can have a perfect correlation with true values while being consistently offset (e.g., predictions of 1, 2, and 3 corresponding to ground truths of 5, 6, and 7 yields a PCC of 1). This issue undermines its reliability for assessing model accuracy.

While PCC has been the only metric used in empathy computing literature on NewsEmp datasets, studies on other datasets sometimes use different metrics. For example, [60], detecting empathy in an audiovisual dataset, adopted the concordance correlation coefficient (CCC) as their primary metric. Previously mentioned shortcomings of PCC could be solved using CCC, as it calculates both the linear relationship and the magnitude of prediction errors. It ensures that predictions are not only aligned with the trend of true values but also close in magnitude, penalising large errors.

Root mean square error (RMSE) appears to be another choice of evaluation as it directly captures prediction error. Overall, PCC, CCC and RMSE measure three distinct qualities of performance: PCC measures linear relationship, RMSE measures the magnitude of errors, and CCC considers both linearity and error magnitude.

C. Implementation Details

We access Llama 3 (version: llama3-70b-8192) through Groq API and GPT-4 (version: gpt-4o) through OpenAI API (last accessed on 30 December 2024). The *temperature* and *top_p* parameters of the APIs control the randomness of LLM outputs during token sampling. To ensure deterministic and consistent labels from LLMs, we set *temperature* to 0 and *top_p* to 0.01.

As pre-trained language models (PLMs), we fine-tune RoBERTa (version: roberta-base) [8], which has 125.7 M trainable parameters, and DeBERTa (version: deberta-v3-base) [61], which has 184 M trainable parameters, from Huggingface [62]. We apply a delayed-start early-stopping strategy that starts monitoring validation CCC after five epochs and stops training if the score does not improve for two successive epochs. Early stopping based on PCC was ineffective in our experience due to higher fluctuations of the PCC score, whereas CCC performed better due to its smoother behaviour. Since most experiments converged within 20 epochs with early stopping, we set the maximum number of training epochs to 20. Deterministic behaviour was enforced using the PyTorch Lightning framework to ensure reproducibility. We save the model checkpoint corresponding to the last epoch of training.

We adopt reported hyperparameters from the original work [8] reporting RoBERTa. Following their reported approach to

TABLE III
HYPERPARAMETERS FOR MODEL TRAINING

Parameter	Value	Parameter	Value
Optimiser	AdamW	Learning rate scheduler	Linear
Learning rate	3e-5	Warmup ratio	0.06
AdamW (β_1, β_2)	(0.9, 0.98)	Batch size	16
AdamW ϵ	1e-6	Maximum epochs	20
Weight decay	0.1	Max sequence length	512

fine-tuning RoBERTa for downstream tasks, we only tuned the learning rate and batch size for our task. Values of the hyperparameters are reported in Table III. While experimenting with the DeBERTa PLM, we use the same hyperparameters as used for RoBERTa.

Following [8], we report median statistics over five different random initialisations (seeds: 0, 42, 100, 999 and 1234). Since prior works on empathy computing on these datasets reported a single peak score of their model, we also report the peak score from these five runs.⁴ All experiments are conducted in Python 3 running on a single AMD Instinct MI250X GPU (64 GB).

D. Main Results

This section presents the quantitative results of our proposed applications of LLM as a service (LLMaaS) in empathy prediction. Unless otherwise stated, results are reported primarily using the RoBERTa PLM, with DeBERTa results explicitly specified.

1) *Noise Mitigation*: We first show evidence of noise in the NewsEmp₂₄ dataset. Table IV illustrates a comparison between human participants' and LLMs' assessments of empathy on a scale of 1 (lowest empathy) to 7 (highest empathy) in two example essays. It demonstrates interesting disparities between crowdsourced and LLM evaluations – for instance, in one essay expressing deep emotional concern for affected people and children, the human rater assigned a relatively low score of 1.0, while the Llama and GPT LLMs rated it much higher at 6.4 and 6.08, respectively. Conversely, a less empathic account of a mining disaster received the maximum possible empathy score (7.0) from human raters but a much lower 1.83 and 1.67 from the Llama and GPT LLMs, respectively. Interestingly, both LLMs, despite differences in size and provider, produce highly consistent annotations, which further underscores the potential inaccuracies of the crowdsourced annotation.

We evaluate our proposed LLMaaS application for noise mitigation in three dataset configurations: NewsEmp₂₄ alone, NewsEmp₂₄₊₂₂, and NewsEmp₂₄₊₂₃₊₂₂. While combining the datasets, we combine their training and validation splits of the additional dataset with the training split of the base dataset for training the model. For example, the NewsEmp₂₄₊₂₂ experimental setup uses the training split of NewsEmp₂₄ and training and validation splits of the NewsEmp₂₂ dataset to train the model. In all cases, the model training is optimised for the

⁴We define *peak* score as the best score (maximum PCC, maximum CCC or minimum RMSE) across five random runs within a *single* experimental setup. Another related terminology used throughout this paper is the *best* score, which refers to the best scores across *different* experimental setups.

TABLE IV

EXAMPLES OF MISLABELLED CROWDSOURCED ANNOTATION, DEVIATING FROM BATSON'S DEFINITION OF EMPATHY. THE FIRST EXAMPLE SHOWS AN ESSAY WITH EMPATHIC ELEMENTS, BUT THE PARTICIPANT'S ANNOTATION INDICATES THE LOWEST EMPATHY. THE SECOND EXAMPLE HAS THE HIGHEST EMPATHY SCORE DESPITE THE ESSAY LACKING EMPATHIC CONTENT. LLM LABELS APPEAR ACCURATE AND CONSISTENT BETWEEN LLAMA AND GPT.

Essay	Crowd	Llama	GPT
"After reading the article, my heart just breaks for the people that are affected by this. Not only are innocent people being killed daily but also little children as well as babies. These children do not deserve this and it's sad because they have their whole lives ahead of them. I really hope that war will end one day although it is looking unlikely."	1.0	6.4	6.08
"I read the article on the China mining disaster. There were 33 miners trapped in the mine. Only two of them survived. Officials stated whoever was responsible would be punished. Smaller mines were shut down immediately until further notice. China has always been known for the deadliest mining."	7.0	1.83	1.67

Empathic expressions are highlighted in blue.

Labels are in a continuous range from 1 to 7, where 1 and 7 refer to the lowest and highest empathy, respectively.

TABLE V
RESULTS OF OUR LLM-BASED NOISE MITIGATION APPROACH, EVALUATED ON THE NEWSEMP₂₄ TEST SET

Labels	α	PCC $\uparrow$	CCC $\uparrow$	RMSE $\downarrow$
Data: NewsEmp₂₄, PLM: RoBERTa				
CS	–	0.331(0.378)	0.307(0.329)	1.656(0.066)
CS & Llama	3.5	<u>0.453</u> (0.462)	0.435 (0.455)	1.604(0.087)
	4.0	0.384(0.464)	0.378(0.454)	<u>1.647</u> (0.075)
	4.5	0.421(0.482)	0.392(0.463)	<u>1.566</u> (0.098)
CS & GPT	3.5	0.473 (0.509)	<u>0.431</u> (0.496)	1.558 (1.499)
	4.0	0.415(0.519)	<u>0.398</u> (<u>0.482</u>)	1.601(1.461)
	4.5	0.370(0.422)	0.325(0.400)	1.646(1.532)
Data: NewsEmp₂₄, PLM: DeBERTa				
CS	–	0.447(0.481)	0.398(0.420)	1.481(1.441)
CS & Llama	3.5	0.502(0.516)	0.450(0.511)	1.531(1.445)
	4.0	<u>0.536</u> (<u>0.576</u>)	<u>0.500</u> (0.518)	<u>1.413</u> (1.369)
	4.5	0.476(0.525)	0.433(0.479)	<u>1.470</u> (1.398)
CS & GPT	3.5	0.529(0.568)	0.489(<u>0.536</u>)	1.419(1.393)
	4.0	0.558 (0.596)	0.533 (0.571)	1.408 (1.326)
	4.5	0.526(0.554)	0.484(0.521)	1.425(<u>1.341</u>)
Data: NewsEmp₂₄₊₂₂, PLM: RoBERTa				
CS	–	0.536(0.597)	0.461(0.505)	1.356(0.042)
CS & Llama	0.0	0.483(0.504)	0.445(0.480)	1.655(1.628)
	0.5	0.488(0.503)	0.445(0.481)	1.646(1.617)
	1.0	0.479(0.491)	0.432(0.451)	1.756(1.664)
	1.5	0.474(0.498)	0.434(0.448)	1.694(1.583)
	2.0	0.455(0.519)	0.421(0.453)	1.661(1.575)
	2.5	0.502(0.547)	0.447(0.456)	1.622(1.585)
	3.0	0.543(0.571)	<u>0.507</u> (0.512)	1.461(1.435)
	3.5	<u>0.558</u> (0.612)	0.496(0.559)	1.389(0.084)
	4.0	0.589 (0.627)	0.516 (<u>0.563</u>)	1.338 (<u>0.054</u>)
	4.5	0.551(<u>0.620</u>)	0.478(0.575)	1.378(0.100)
	5.0	0.551(0.605)	0.478(0.520)	<u>1.348</u> (1.265)
CS & GPT	5.5	0.542(0.604)	0.464(0.516)	1.352(1.265)
	6.0	0.536(0.597)	0.461(0.505)	1.356(1.274)
Data: NewsEmp₂₄₊₂₃₊₂₂, PLM: RoBERTa				
CS	–	0.528(0.551)	0.469(0.498)	1.380(0.086)
CS & Llama	3.5	<u>0.556</u> (0.573)	<u>0.511</u> (<u>0.552</u>)	1.381(<u>0.029</u>)
	4.0	0.574 (<u>0.582</u>)	0.529 (0.548)	1.333 (0.021)
	4.5	0.548(0.648)	0.479(0.597)	<u>1.346</u> (0.092)

CS – crowdsourced labels; α – annotation selection threshold.

Reported metrics are in median (peak) format, calculated from five random initialisations.

Boldface and underline texts indicate the best and the second-best scores, respectively.

validation split of the NewsEmp₂₄ dataset and finally evaluated on the hold-out NewsEmp₂₄ test split.

As presented in Table V, our LLM-based noise mitigation approach demonstrates consistent performance improvements

across all configurations. For the NewsEmp₂₄ dataset configuration, we report results with both Llama- and GPT-generated labels. While we choose RoBERTa as the primary PLM, we also report results with DeBERTa PLM in this setup, which shows even better performance improvement. Specifically, GPT labels at $\alpha = 4.0$ achieve the best median PCC (0.558), CCC (0.533) and RMSE (1.408) scores. Overall, the performance improvement between Llama and GPT labels is comparable, with each achieving the best results in certain metrics. Since Llama is an open-weight LLM and freely available, we proceed with Llama for the remaining experiments in this application scenario.

Including the NewsEmp₂₂ dataset enhances performance further, with $\alpha = 4.0$ yielding the best median PCC (0.589) and median CCC (0.516). When combined with NewsEmp₂₃, the baseline metrics remain comparable, but $\alpha = 4.0$ again delivers the highest median PCC (0.574) and median CCC (0.529), with RMSE achieving its lowest value of 1.333. Considering peak scores instead of median statistics across five runs, our approach also outperforms the baseline model by achieving the peak PCC of 0.648, CCC of 0.597 and RMSE of 0.021. As illustrated earlier in Fig. 1, the performance improvements are *statistically* significant.

The value of the threshold α controls the proportion of LLM and crowdsourced labels. As demonstrated earlier, a smaller value of α means having a higher amount of LLM labels, which may hurt the model's performance in the test set. This hypothesis is verified in Table V's results on varying α on the NewsEmp₂₄₊₂₂ scenario, which shows that $\alpha = 3.5 \sim 4.5$ provides the best performance across the three dataset configurations.

Our noise mitigation approach outperforms the baseline in terms of median PCCs, CCCs and RMSEs, as well as peak PCCs and CCCs across all four configurations. Out of 24 test cases,⁵ only two cases (RoBERTa on NewsEmp₂₄ and NewsEmp₂₄₊₂₂), show a better peak RMSE achieved by the baseline. This discrepancy can be primarily attributed to how training was controlled: we applied early stopping based on CCC to mitigate overfitting. Since CCC and RMSE are not strongly correlated [63], [64], early stopping by CCC does not necessarily optimise RMSE.

Recent works in affective computing, including emotion recognition [65] and empathy computing [12], preferred correlation-based metrics over error-based metrics. Our findings

⁵2 types (median, peak) $\times$ 3 metrics $\times$ 4 configurations in Table V.

TABLE VI
EFFECT OF ADDITIONAL LABELLED DATA

Training data	PCC $\uparrow$	CCC $\uparrow$	RMSE $\downarrow$
PLM: RoBERTa			
NewsEmp ₂₄	0.331(0.378)	0.307(0.329)	1.656(0.066)
+ Crowd-labelled NewsEmp ₂₂	0.485(0.594)	0.439(0.480)	1.417 (0.093)
+ Llama-labelled NewsEmp ₂₂	0.513 (0.571)	0.490 (0.523)	<u>1.484</u> (0.059)
+ GPT-labelled NewsEmp ₂₂	<u>0.495</u> (0.549)	<u>0.446</u> (0.455)	1.581(1.514)
PLM: DeBERTa			
NewsEmp ₂₄	0.447(0.481)	0.398(0.420)	1.481(1.441)
+ Crowd-labelled NewsEmp ₂₂	<u>0.564</u> (0.638)	0.478(0.566)	1.329 (1.233)
+ Llama-labelled NewsEmp ₂₂	0.581 (0.601)	0.535 (0.584)	<u>1.445</u> (1.366)
+ GPT-labelled NewsEmp ₂₂	0.554(0.596)	<u>0.493</u> (0.534)	1.468(1.391)
PLM: RoBERTa			
NewsEmp ₂₂	0.459(0.477)	0.363(0.411)	1.776(0.046)
+ Crowd-labelled NewsEmp ₂₄	<u>0.467</u> (0.478)	<u>0.392</u> (0.435)	1.756(0.062)
+ Llama-labelled NewsEmp ₂₄	0.496 (0.519)	0.429 (0.434)	1.729 (0.034)

All evaluations use test splits, except for NewsEm22, where CCC and RMSE are computed on the validation split due to unavailable test labels.

Reported metrics are in median (peak) format, calculated across five random initialisations.

Boldface and underline texts indicate the best and the second-best scores, respectively.

align with this trend: although the baseline approach incidentally achieves a better peak RMSE in two isolated runs, our method demonstrates more consistent and robust performance across all metrics, including median RMSE across five runs, which reflects typical behaviour rather than outliers. In terms of the choice of the primary metric, PCC only captures linear correlation but not error magnitude, while RMSE only captures error magnitude but not correlation; therefore, our recommendation is to consider CCC as the primary metric, as it captures the best of both worlds – both linear correlation and error magnitude.

2) *Additional Data Labelled by LLM*: The first application described above demonstrates that additional training data helps achieve better performance. However, we may not always have the flexibility of having extra data labelled by human participants. This application, therefore, explores whether additional data labelled by LLM could help.

We evaluate this application in two configurations: either NewsEmp₂₄ or NewsEmp₂₂ as the base dataset. While evaluating the model on the NewsEmp₂₄ dataset, we consider the NewsEmp₂₂ dataset as additional data and vice versa. Since we have both crowdsourced and LLM-generated labels for the additional data, we compare whether the additional data is labelled by (1) human participants or (2) LLM.

Table VI reports the performance in both settings. When trained a RoBERTa model with the NewsEmp₂₄ dataset alone, it achieved a median PCC of 0.331 and 0.307 CCC. Additional NewsEmp₂₂ dataset labelled by human participants boosted performance to 0.485 PCC, 0.439 CCC and 1.417 RMSE. The same NewsEmp₂₂ dataset – labelled by Llama LLM – makes the highest median PCC of 0.513 and CCC of 0.490 while maintaining a competitive RMSE of 1.484. For this NewsEmp₂₄ setup, we further report results with the DeBERTa PLM, which shows even better performance across different metrics. Like our earlier experiment (Application 1), the use of either Llama or GPT LLMs yields similar performance in this setup, and we proceed with Llama LLMs for the rest of the experiments.

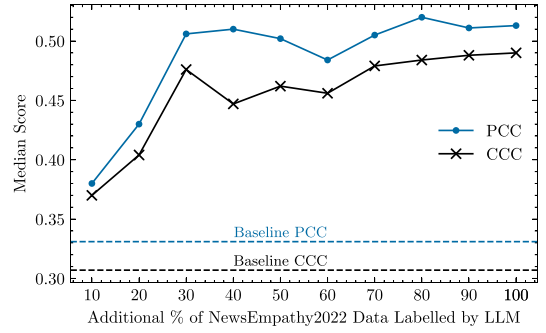


Fig. 3. Median performance in the NewsEmp₂₄ test set with gradual increase of additional LLM-labelled data. *Baseline* scores refer to the scores achieved using only NewsEmp₂₄ data.

Experiments using NewsEmp₂₂ as the base dataset exhibited a similar trend to those with NewsEmp₂₄ as the base dataset. The model trained with LLM-generated labels achieves the best PCC (0.496) and CCC (0.429), as well as the lowest RMSE (1.729). Notably, the human-labelled data shows improvement over the base dataset but falls short of the performance achieved with LLM-labelled data. Overall, LLM labels are as good as crowdsourced labels, and additional data, either crowdsourced or LLM-labelled, boosts the performance.

Compared to our first application scenario (mixed labels to reduce label noise), performance improvement from baseline in this application is statistically more significant across all three evaluation metrics (Fig. 1). In particular, this application scenario demonstrates the highest level of statistically significant improvements in terms of PCC and CCC. This is likely because, in this scenario, the labels of the base training data remain unchanged, which is presumably of a similar distribution to the hold-out test set. The additional data provides extra supervision, which helps achieve a better score. Results using additional data labelled by LLM are better than using human labelling in most cases (12 out of 18 test cases), likely because of the higher quality of labels from LLM.

To understand how the amount of additional data affects the performance, we gradually increase the amount of additional data and report model performance (Fig. 3). In each case, we randomly sample a percentage of the additional data ranging from 10% to 100%. Surprisingly, the performance increases most rapidly from 10 to 30%, after which the improvement slows down.

Fig. 4 presents 3D t-SNE visualisations of embeddings derived from different labelling schemes, alongside their clustering performance measured by the mean Silhouette score [66] on the embeddings from PLM. Given the continuous nature of empathy labels in the datasets, we discretise them into six bins (1–2, 2–3, ..., 6–7) to calculate the metric. This score ranges from –1 to +1, with –1 being “misclassified” and +1 being “well-clustered” [66].

Embeddings from crowdsourced labels exhibit a slightly dispersed distribution visually and the lowest Silhouette score (–0.005), which suggests many samples are mislabelled (i.e., label noise). In contrast, embeddings based on LLM labels display smoother distributions (especially the rightmost plot) and higher

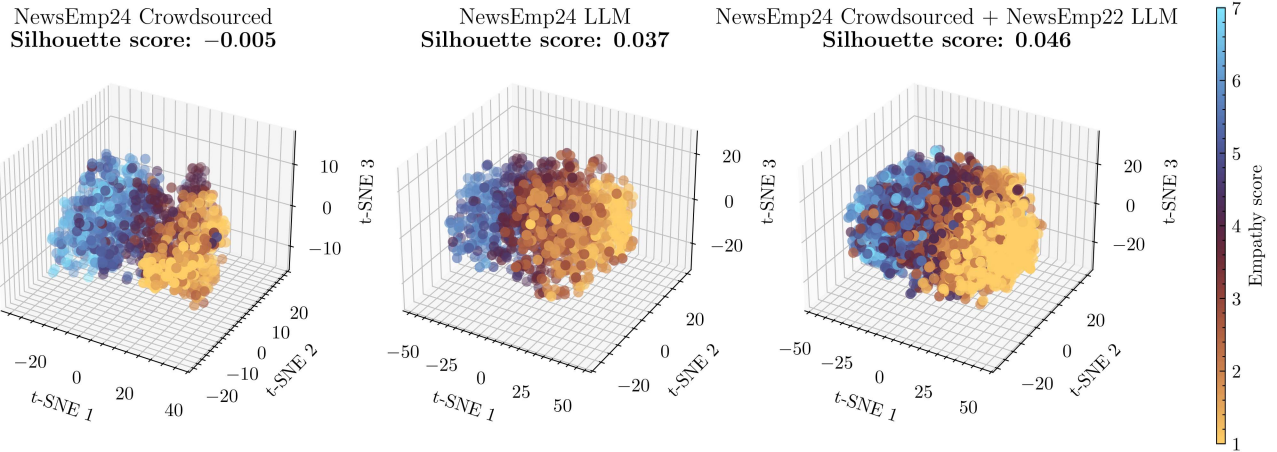


Fig. 4. 3D t-SNE visualisation and Silhouette scores on CLS embeddings from pre-trained language models fine-tuned using crowdsourced labels (**left**), LLM-generated labels (**middle**) and crowdsourced + additional LLM-labelled data (**right**). Continuous empathy labels are discretised into six bins to calculate Silhouette scores (higher is better), which suggests better clustering after integrating LLM-generated labels.

TABLE VII
ZERO-SHOT EMPATHY PREDICTION USING LLMs

Dataset	LLM	Split	PCC $\uparrow$	CCC $\uparrow$	RMSE $\downarrow$
NewsEmp ₂₄	Llama	Test	0.441	0.436	1.731
		Validation	0.502	0.457	1.952
	Llama _{plain}	Test	0.405	0.358	1.859
	GPT	Test	0.581	0.489	1.715
		Validation	0.480	0.375	2.038
NewsEmp ₂₃	Llama	Test	0.380	—	—
		Validation	0.108	0.108	2.18
NewsEmp ₂₂	Llama	Test	0.517	—	—
		Validation	0.579	0.573	1.728

Llama_{plain}: Llama with the plain prompt. All other entries use the scale-aware prompt.

Ground truth of the NewsEm22 and NewsEmp₂₃ test splits are unavailable to calculate CCC and RMSE.

Silhouette scores (0.037 and 0.046),⁶ which suggests many of the mislabelled samples are now correctly labelled (i.e., reduction of label noise). The smoothest distribution and the highest Silhouette score on the NewsEmp₂₄ Crowdsourced + NewsEmp₂₂ LLM configuration can be attributed to extra supervision from the LLM-generated labels.

E. Zero-Shot Prediction & Demographic Biases

The most direct application of LLM is to predict empathy in a zero-shot manner, i.e., without any training or fine-tuning of the LLM. Table VII reports the performances of zero-shot prediction across all validation and test splits.

On the NewsEmp₂₄ test split, we compare plain prompting with our proposed scale-aware prompting scheme. The improved performance with the scale-aware prompt (PCC:

0.405 $\rightarrow$ 0.441; CCC: 0.358 $\rightarrow$ 0.436; RMSE: 1.859 $\rightarrow$ 1.731) demonstrates its effectiveness. Accordingly, we adopt the scale-aware prompt as the default in all our experiments.

We further examine how the agreement between crowdsourced annotation and LLM annotation varies across different demographic groups. For this analysis, we combine training, validation and test splits of the NewsEmp₂₄ dataset and compare crowdsourced and Llama-generated labels. As demonstrated in Fig. 5, both have similar levels of CCC in gender and education demographics. However, CCC changed wildly across race, age and income groups. In particular, it went to negatives in two race groups: “Hispanic/Latino” and “Other” categories, noting that there are only four samples in the “Other” category of race.

In terms of the number of samples across different demographics, we see that certain demographic groups (e.g., “4-year bachelor’s degree” education, “White” race, and “31-40” age groups) are highly represented compared to their counterparts. Some demographics, for example, “Less than high school” education and “Native American / American Indian (Indigenous People)” are not represented in the dataset at all. Such imbalances can presumably introduce bias in the empathy detection model built from such biased datasets.

LLMs designed for empathy detection may exhibit biases across different demographic groups [17]. A proper assessment of such bias would require accurate and unbiased ground truth labels to compare with. Note that Fig. 5 compares LLM outputs against potentially noisy crowdsourced labels and so may not offer an objective assessment of bias. Instead, it reflects potential *sources* of bias that may influence an empathy detection model, insofar as biased training data can propagate bias into the model’s predictions.

Our prompt to the LLM is *demographic-unaware*, meaning it does not include any explicit demographic information. Theoretically, LLM outputs through such an unaware prompt should be less biased (because the LLM does not have access to demographic information) compared to a demographic-aware prompt [17]. [17] argued that a

⁶Note that although Fig. 4 shows relative differences in Silhouette scores across embeddings, their absolute values remain low, presumably due to the discretisation of continuous labels (the clustered labels are not well-separated groups).

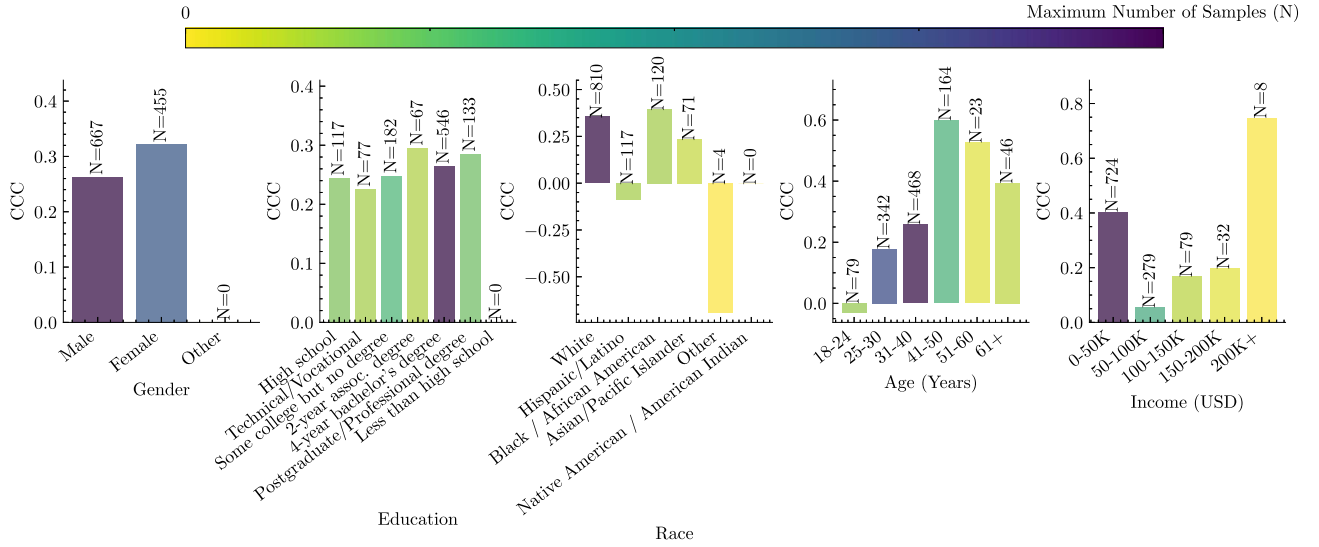


Fig. 5. Number of samples and zero-shot (Llama) prediction performance across different demographic groups in the NewsEmp₂₄ dataset. CCC varies rapidly across different racial groups.

TABLE VIII

COMPARISON OF OUR PROPOSED MODEL WITH THE LITERATURE ON THE NEWSEMP₂₄ TEST DATASET

Approach	Base Model (Ref.)	PCC ↑	CCC ↑	RMSE ↓
Training	BERT [53]	0.290	—	—
	MLP [67]	0.345	—	—
	RoBERTa [49]	0.375	—	—
	Not mentioned [68]	0.390	—	—
	Llama 3 8B [2]	0.474	—	—
	RoBERTa [37]	0.629	—	—
	RoBERTa [37] ^a	0.607	0.498	0.075
	RoBERTa (Ours)	0.648	0.597	0.092
Zero-shot	GPT 3.5 [3]	0.523	—	—
	GPT 4 (Ours)	0.581	0.489	1.715

^a Our implementation of the earlier SOTA work [37].

demographic-aware prompt may also result in less biased assessment, since the LLMs are being alerted to potential bias. They found mixed outcomes on mitigating bias through demographic-aware and unaware prompts across different LLMs [17]. Mitigating bias is critical for real-world deployment, and so future exploration towards a universally applicable prompting strategy is warranted to mitigate demographic biases.

F. Comparison With the Literature

A quantitative comparison between our proposed framework and other empathy detection works in the literature on the NewsEmp₂₄ dataset is presented on Table VIII. The best performance in the literature is 0.629 PCC [37], while our best performance is 0.648 PCC. Both [37] and we use the same amount of training data – combined NewsEmp₂₄₊₂₃₊₂₂.

The reported results in the literature are in terms of a single evaluation metric, which suggests the peak performance of their model. To compare in terms of other evaluation metrics in a similar setting of ours (five random initialisations), we implemented the state-of-the-art work [37]. The mismatch between our implementation and [37]’s reported result (0.607 vs 0.629)

TABLE IX

CONSISTENCY AND INTER-RATER RELIABILITY AMONG LLAMA, GPT AND CROWDSOURCED ANNOTATIONS ON THE NEWSEMP₂₄ TRAINING SET. THE LOW RELIABILITY BETWEEN LLMS AND CROWDSOURCED ANNOTATIONS, CONTRASTED WITH THE HIGH RELIABILITY BETWEEN TWO DIFFERENT LLMS, MAY SUGGEST THAT THE CROWDSOURCED ANNOTATIONS ARE NOISY

Annotator 1	Annotator 2	Krippendorff’s Alpha	MAE ± SD
Llama	Llama	0.99	0.10 ± 0.21
Llama	GPT	0.80	0.78 ± 0.70
Llama	Crowd	0.27	1.72 ± 1.34
GPT	Crowd	0.19	1.81 ± 1.27

MAE – mean absolute error; SD – standard deviation of absolute error

is likely due to hyperparameter choice. Having no public implementation of [37], we chose default hyperparameters, apart from the minimal amount of hyperparameter details reported in their work. Our approach outperforms [37]’s results in terms of both PCC and CCC (Table VIII).

G. Consistency and Inter-Rater Reliability

LLMs are known to produce varying outputs across different API calls [69]. This variability could raise concerns about using LLM to label data as well as evaluating on the test set. We calculate two types of consistency: *intra-LLM* consistency, which evaluates whether annotations generated by the same LLM model remain consistent across multiple API calls, and *inter-LLM* consistency, which assesses whether annotations are consistent between two different LLMs.

To assess intra-LLM consistency, we label the NewsEmp₂₄ training set (1,000 samples) twice in the Llama LLM, using separate independent API calls. The results (Table IX) demonstrate “almost perfect” agreement between the two annotation rounds, with a Krippendorff’s Alpha (K-Alpha) [70] score of 0.99. Between GPT and LLM annotations (inter-LLM), a K-Alpha of 0.80 is achieved, which lies on the boundary between “substantial” and “almost perfect” reliability [70]. Such a high level of consistency, including inter-LLM consistency, suggests the

effectiveness of our prompting strategy, which clearly specifies the expectations from the LLM.

As presented in Table IX, the inter-rater reliability between LLMs and crowdsourced annotations is notably lower than the reliability observed between LLMs. It potentially supports our hypothesis that crowdsourced annotations are inherently noisy.

H. Human Preference Study: Crowdsourced vs LLM Labels

To evaluate the credibility of LLM-generated empathy scores relative to crowdsourced annotations, we conducted a small-scale human assessment study involving three co-authors (all PhD holders, including two mid-career and one senior academic) as independent assessors. The corresponding author designed the experiment, while the assessors were blind to the origin of each label (LLM or crowdsourced). We selected 20 samples that exhibited the largest disagreement between LLM and crowdsourced annotations to better test discernibility. Each assessor was shown an essay along with two empathy scores (randomised in order) and asked to choose the score that best reflected the empathy expressed in the essay, following Batson's definition of empathy.

Across the 20 samples, 8 were unanimously rated in favour of the LLM-generated labels by all three assessors. In 10 other samples, two out of three assessors preferred the LLM labels. Only in 2 cases did the majority (2 of 3) select the crowdsourced label. Based on majority voting, LLM-generated labels were preferred in 18 out of 20 instances, suggesting that, at least in these high-discrepancy cases, LLM annotations aligned more closely with human expert judgment than the original crowdsourced labels.

I. Limitations and Future Work

As we note the inherent biases in LLM, it is crucial to exercise caution when using LLM-generated labels if such a system is to be deployed in real life. Zero-shot predictions are likely to exhibit greater bias, as LLMs may inherit biases from their training data. Therefore, downstream models should be trained on diverse and representative datasets that reflect the demographics in which empathy would be detected.

As this paper investigates how LLM-generated labels can help measure empathy by smaller PLMs, we restrict the LLMs to generating labels only. However, incorporating the reasoning behind these labels could enhance explainability and potentially further improve label quality. Future work may explicitly prompt LLMs to demand justification, employ chain-of-thought prompting to elicit reasoning [71], or leverage LLMs having implicit reasoning capability (e.g., OpenAI o3). Lastly, as all experiments are conducted on English datasets (NewsEmp series), exploring the multilingual adaptability of our methods remains an important direction for future work.

V. CONCLUSION

This work demonstrates the potential of large language models (LLMs) in addressing challenges in empathy computing through two *in-vitro* applications: label noise reduction and

training data expansion. Both applications resulted in statistically significant performance gains over baseline methods. The proposed framework outperformed state-of-the-art methods, achieving new benchmarks on a public empathy dataset with a Pearson correlation coefficient (PCC) of 0.648, among other metrics. Beyond the empirical results, this paper contributes a critical rethinking of evaluation practices in empathy computing, advocating for the adoption of the concordance correlation coefficient (CCC). The novel scale-aware prompting technique introduced here ensures alignment between LLM annotations and theoretical annotation protocols. We further highlight biases in the dataset across different demographic groups. Similar to the empathy detection dataset addressed in this paper, many other tasks, such as detecting depression, anxiety and mental health conditions, rely on questionnaire-based self-annotations. The proposed approach, therefore, opens exciting avenues for leveraging LLMs as complementary tools to enhance model training across different domains.

REFERENCES

- [1] Z. Ma et al., "Leveraging speech PTM, text LLM, and emotional TTS for speech emotion recognition," in *Proc. IEEE Int. Conf. Acoust. Speech Signal Process.*, 2024, pp. 11146–11150.
- [2] T. Li, N. Rusnachenko, and H. Liang, "Chinchunmei at WASSA 2024 empathy and personality shared task: Boosting LLM's prediction with role-play augmentation and contrastive reasoning calibration," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. De Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, Aug. 2024, pp. 385–392. [Online]. Available: <https://aclanthology.org/2024.wassa-1.32>
- [3] H. Kong and S. Moon, "RU at WASSA 2024 shared task: Task-aligned prompt for predicting empathy and distress," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. De Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, Aug. 2024, pp. 380–384. [Online]. Available: <https://aclanthology.org/2024.wassa-1.31>
- [4] T. Sun, Y. Shao, H. Qian, X. Huang, and X. Qiu, "Black-box tuning for language-model-as-a-service," in *Proc. Int. Conf. Mach. Learn.*, 2022, pp. 20841–20855.
- [5] Y. Hu et al., "Improving large language models for clinical named entity recognition via prompt engineering," *J. Amer. Med. Informat. Assoc.*, vol. 31, no. 9, pp. 1812–1820, Jan. 2024.
- [6] H. Fei, B. Li, Q. Liu, L. Bing, F. Li, and T.-S. Chua, "Reasoning implicit sentiment with chain-of-thought prompting," in *Proc. 61st Annu. Meeting Assoc. Comput. Linguistics*, vol. 2, A. Rogers, J. Boyd-Graber, and N. Okazaki, Eds., Toronto, ON, Canada, Jul. 2023, pp. 1171–1182.
- [7] Z. Zheng, L. Zheng, and Y. Yang, "Unlabeled samples generated by GAN improve the person re-identification baseline in vitro," in *Proc. IEEE Int. Conf. Comput. Vis.*, 2017, pp. 3774–3782.
- [8] Y. Liu et al., "RoBERTa: A robustly optimized BERT pretraining approach," 2019, *arXiv:1907.11692*.
- [9] M. Mazumder et al., "DataPerf: Benchmarks for data-centric AI development," in *Proc. Adv. Neural Inf. Process. Syst.*, A. Oh, T. Naumann, A. Globerson, K. Saenko, M. Hardt, and S. Levine, Eds., 2023, vol. 36, pp. 5320–5347.
- [10] R. S. Geiger et al., "Garbage in, garbage out? Do machine learning application papers in social computing report where human-labeled training data comes from?," in *Proc. Conf. Fairness, Accountability, Transparency*, 2020, pp. 325–336.
- [11] M. L. Hoffman, *Empathy and Moral Development: Implications for Caring and Justice*. Cambridge, U.K.: Cambridge Univ. Press, 2000.
- [12] M. R. Hasan, M. Z. Hossain, S. Ghosh, A. Krishna, and T. Gedeon, "Empathy detection from text, audiovisual, audio or physiological signals: A systematic review of task formulations and machine learning methods," *IEEE Trans. Affect. Comput.*, vol. 16, no. 4, pp. 2507–2525, Oct.–Dec. 2025.
- [13] C. D. Batson, J. Fultz, and P. A. Schoenrade, "Distress and empathy: Two qualitatively distinct vicarious emotions with different motivational consequences," *J. Pers.*, vol. 55, no. 1, pp. 19–39, 1987.

- [14] F. Charlier et al., *Statannotations*, version 0.6., Zenodo, Oct. 2023, doi: [10.5281/zenodo.8396665](https://doi.org/10.5281/zenodo.8396665).
- [15] B. D. Jani, D. N. Blane, and S. W. Mercer, "The role of empathy in therapy and the physician-patient relationship," *Complement. Med. Res.*, vol. 19, no. 5, pp. 252–257, 2012.
- [16] K. Aldrup, B. Carstensen, and U. Klusmann, "Is empathy the key to effective teaching? A systematic review of its association with teacher-student interactions and student outcomes," *Educ. Psychol. Rev.*, vol. 34, no. 3, pp. 1177–1216, 2022.
- [17] S. Gabriel, I. Puri, X. Xu, M. Malgaroli, and M. Ghassemi, "Can AI relate: Testing large language model response for mental health support," in *Proc. Findings Assoc. Comput. Linguistics*, Y. Al-Onaizan, M. Bansal, and Y.-N. Chen, Eds., Miami, FL, USA, Nov. 2024, pp. 2206–2221.
- [18] C. Zhang, S. Bengio, M. Hardt, B. Recht, and O. Vinyals, "Understanding deep learning (still) requires rethinking generalization," *Commun. ACM*, vol. 64, no. 3, pp. 107–115, 2021.
- [19] S. Tafreshi, O. De Clercq, V. Barriere, S. Buechel, J. Sedoc, and A. Balahur, "WASSA 2021 shared task: Predicting empathy and emotion in reaction to news stories," in *Proc. 11th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, Apr. 2021, pp. 92–104. [Online]. Available: <https://aclanthology.org/2021.wassa-1.10>
- [20] S. Mohammad and P. Turney, "Emotions evoked by common words and phrases: Using mechanical turk to create an emotion lexicon," in *Proc. NAACL HLT 2010 Workshop Comput. Approaches Anal. Gener. Emotion Text*, 2010, pp. 26–34.
- [21] K. B. Sheehan, "Crowdsourcing research: Data collection with Amazon's mechanical turk," *Commun. Monographs*, vol. 85, no. 1, pp. 140–156, 2018.
- [22] R. Jia, Z. R. Steelman, and B. H. Reich, "Using mechanical turk data in is research: Risks, rewards, and recommendations," *Commun. Assoc. Inf. Syst.*, vol. 41, no. 1, 2017, Art. no. 14.
- [23] J. L. Huang, P. G. Curran, J. Keeney, E. M. Poposki, and R. P. DeShon, "Detecting and deterring insufficient effort responding to surveys," *J. Bus. Psychol.*, vol. 27, pp. 99–114, 2012.
- [24] W. O'Brochta and S. Parikh, "Anomalous responses on amazon mechanical turk: An indian perspective," *Res. Politics*, vol. 8, no. 2, 2021, Art. no. 20531680211016971.
- [25] M. Niu, M. Jaiswal, and E. M. Provost, "From text to emotion: Unveiling the emotion annotation capabilities of LLMs," in *Proc. Interspeech*, 2024, pp. 2650–2654.
- [26] S. Wang, Y. Liu, Y. Xu, C. Zhu, and M. Zeng, "Want to reduce labeling cost? GPT-3 can help," in *Proc. Findings Assoc. Comput. Linguistics*, Punta Cana, Dominican Republic, Nov. 2021, pp. 4195–4205.
- [27] M. R. Hasan, M. Z. Hossain, T. Gedeon, and S. Rahman, "LLM-Gen: Large language model-guided prediction of people's empathy levels towards newspaper article," in *Proc. Findings Assoc. Comput. Linguistics*, Y. Graham and M. Purver, Eds., St. Julian's, Malta, Mar. 2024, pp. 2215–2231. [Online]. Available: <https://aclanthology.org/2024.findings-eacl.147>
- [28] A. Grattafiori et al., "The Llama 3 herd of models," 2024, *arXiv:2407.21783*.
- [29] OpenAI et al., "GPT-4 technical report," 2024, *arXiv:2303.08774*.
- [30] Y. Wang and T. Baldwin, "Noisy label regularisation for textual regression," in *Proc. 29th Int. Conf. Comput. Linguistics*, N. Calzolari et al., Eds., Gyeongju, Republic of Korea, Oct. 2022, pp. 4228–4240.
- [31] E. Engleson and H. Azizpour, "Robust classification via regression for learning with noisy labels," in *Proc. 12th Int. Conf. Learn. Representations*, 2024, pp. 1–21. [Online]. Available: <https://openreview.net/forum?id=wfgZc3IMqo>
- [32] N. Natarajan, I. S. Dhillon, P. K. Ravikumar, and A. Tewari, "Learning with noisy labels," in *Proc. Adv. Neural Inf. Process. Syst.*, C. Burges, L. Bottou, M. Welling, Z. Ghahramani, and K. Weinberger, Eds., vol. 26, 2013, pp. 1–9. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2013/file/3871bd64012152bfb53fd04b401193f-Paper.pdf
- [33] S. Garg, G. Ramakrishnan, and V. Thumbe, "Towards robustness to label noise in text classification via noise modeling," in *Proc. 30th ACM Int. Conf. Inf. Knowl. Manage.*, 2021, pp. 3024–3028.
- [34] P. Li et al., "An improved categorical cross entropy for remote sensing image classification based on noisy labels," *Expert Syst. Appl.*, vol. 205, 2022, Art. no. 117296.
- [35] B. Zhang et al., "Flexmatch: Boosting semi-supervised learning with curriculum pseudo labeling," in *Proc. Adv. Neural Inf. Process. Syst.*, M. Ranzato, A. Beygelzimer, Y. Dauphin, P. Liang, and J. W. Vaughan, Eds., 2021, vol. 34, pp. 18408–18419. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2021/file/995693c15f439e3d189b06e89d145dd5-Paper.pdf
- [36] K. Sohn et al., "Fixmatch: Simplifying semi-supervised learning with consistency and confidence," in *Proc. Adv. Neural Inf. Process. Syst.*, H. Larochelle, M. Ranzato, R. Hadsell, M. Balcan, and H. Lin, Eds., 2020, vol. 33, pp. 596–608. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2020/file/06964dce9addb1c5cb5d6e3d9838f733-Paper.pdf
- [37] S. Giorgi, J. Sedoc, V. Barriere, and S. Tafreshi, "Findings of WASSA 2024 shared task on empathy and personality detection in interactions," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. De Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, Aug. 2024, pp. 369–379. [Online]. Available: <https://aclanthology.org/2024.wassa-1.30>
- [38] S. Qian, C. Orašan, D. Kanojia, H. Saadany, and F. Do Carmo, "SURREY-CTS-NLP at WASSA2022: An experiment of discourse and sentiment analysis for the prediction of empathy, distress and emotion," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, 2022, pp. 271–275.
- [39] A. Lahnala, C. Welch, and L. Flek, "CAISA at WASSA 2022: Adapter-tuning for empathy prediction," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, 2022, pp. 280–285.
- [40] Y. Chen, Y. Ju, and S. Kübler, "IUCL at WASSA 2022 shared task: A text-only approach to empathy and emotion detection," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, 2022, pp. 228–232.
- [41] F. M. Plaza-del Arco, J. Collado-Montañez, L. A. Ureña, and M.-T. Martiñ-Valdivia, "Empathy and distress prediction using transformer multi-output regression and emotion analysis with an ensemble of supervised and zero-shot learning models," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, Dublin, Ireland, May 2022, pp. 239–244.
- [42] Y. Wang, J. Wang, and X. Zhang, "YNU-HPCC at WASSA-2023 shared task 1: Large-scale language model with LoRA fine-tuning for empathy detection and emotion classification," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, Canada, Jul. 2023, pp. 526–530. [Online]. Available: <https://aclanthology.org/2023.wassa-1.45>
- [43] V. Barriere, J. Sedoc, S. Tafreshi, and S. Giorgi, "Findings of WASSA 2023 shared task on empathy, emotion and personality detection in conversation and reactions to news articles," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, ON, Canada, Jul. 2023, pp. 511–525. [Online]. Available: <https://aclanthology.org/2023.wassa-1.44>
- [44] F. Gruschka, A. Lahnala, C. Welch, and L. Flek, "CAISA at WASSA 2023 shared task: Domain transfer for empathy, distress, and personality prediction," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, Canada, Jul. 2023, pp. 553–557. [Online]. Available: <https://aclanthology.org/2023.wassa-1.50>
- [45] H. Vasava, P. Uikey, G. Wasnik, and R. Sharma, "Transformer-based architecture for empathy prediction and emotion classification," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, 2022, pp. 261–264.
- [46] A. Kulkarni, S. Somwase, S. Rajput, and M. Marathe, "PVG at WASSA 2021: A multi-input, multi-task, transformer-based architecture for empathy and distress prediction," in *Proc. 11th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, Apr. 2021, pp. 105–111. [Online]. Available: <https://aclanthology.org/2021.wassa-1.11>
- [47] A. S. Srinivas, N. Barua, and S. Pal, "Team_Hawk at WASSA 2023 empathy, emotion, and personality shared task: Multi-tasking multi-encoder based transformers for empathy and emotion prediction in conversations," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, ON, Canada, Jul. 2023, pp. 542–547. [Online]. Available: <https://aclanthology.org/2023.wassa-1.48>
- [48] X. Lu, Z. Li, Y. Tong, Y. Zhao, and B. Qin, "HIT-SCIR at WASSA 2023: Empathy and emotion analysis at the utterance-level and the essay-level," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, Canada, Jul. 2023, pp. 574–580. [Online]. Available: <https://aclanthology.org/2023.wassa-1.54>
- [49] R. Frick and M. Steinebach, "Fraunhofer SIT at WASSA 2024 empathy and personality shared task: Use of sentiment transformers and data augmentation with fuzzy labels to predict emotional reactions in conversations and essays," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. De Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, Aug. 2024, pp. 435–440. [Online]. Available: <https://aclanthology.org/2024.wassa-1.40>

- [50] S. Ghosh, D. Maurya, A. Ekbal, and P. Bhattacharyya, "Team IITP-AINLPML at WASSA 2022: Empathy detection, emotion classification and personality detection," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, 2022, pp. 255–260.
- [51] Y. Butala, K. Singh, A. Kumar, and S. Shrivastava, "Team phoenix at WASSA 2021: Emotion analysis on news stories with pre-trained language models," in *Proc. 11th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, Apr. 2021, pp. 274–280. [Online]. Available: <https://aclanthology.org/2021.wassa-1.30>
- [52] M. R. Hasan, M. Z. Hossain, T. Gedeon, S. Soon, and S. Rahman, "Curtin OCAI at WASSA 2023 empathy, emotion and personality shared task: Demographic-aware prediction using multiple transformers," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, ON, Canada, Jul. 2023, pp. 536–541.
- [53] A. Numanoğlu, S. Ateş, N. Cicekli, and D. Küçük, "Empathify at WASSA 2024 empathy and personality shared task: Contextualizing empathy with a BERT-based context-aware approach for empathy detection," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. D. Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, 2024, vol. 8, pp. 393–398. [Online]. Available: <https://aclanthology.org/2024.wassa-1.33>
- [54] J. Mundra, R. Gupta, and S. Mukherjee, "WASSA, IITK at WASSA 2021: Multi-task learning and transformer finetuning for emotion classification and empathy prediction," in *Proc. 11th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, Apr. 2021, pp. 112–116. [Online]. Available: <https://aclanthology.org/2021.wassa-1.12>
- [55] T.-M. Lin, J.-Y. Chang, and L.-H. Lee, "NCUEE-NLP at WASSA 2023 shared task 1: Empathy and emotion prediction using sentiment-enhanced RoBERTa transformers," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, Canada, Jul. 2023, pp. 548–552. [Online]. Available: <https://aclanthology.org/2023.wassa-1.49>
- [56] T. Chavan, K. Deshpande, and S. Sonawane, "PICT-CLRL at WASSA 2023 empathy, emotion and personality shared task: Empathy and distress detection using ensembles of transformer models," in *Proc. 13th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, Toronto, ON, Canada, Jul. 2023, pp. 564–568. [Online]. Available: <https://aclanthology.org/2023.wassa-1.52>
- [57] S. Buechel, A. Buffone, B. Slaff, L. Ungar, and J. Sedoc, "Modeling empathy and distress in reaction to news stories," in *Proc. Conf. Empirical Methods Natural Lang. Process.*, Brussels, Belgium, Oct. 2018, pp. 4758–4765.
- [58] D. Omiaomu et al., "Empathic conversations: A multi-level dataset of contextualized conversations," 2022, *arXiv:2205.12698*.
- [59] V. Barriere, S. Tafreshi, J. Sedoc, and S. Alqahtani, "WASSA 2022 shared task: Predicting empathy, emotion and personality in reaction to news stories," in *Proc. 12th Workshop Comput. Approaches Subjectivity, Sentiment Social Media Anal.*, Dublin, Ireland, May 2022, pp. 214–227.
- [60] P. Barros, N. Churamani, A. Lim, and S. Wermter, "The OMG-empathy dataset: Evaluating the impact of affective behavior in storytelling," in *Proc. 8th Int. Conf. Affect. Comput. Intell. Interact.*, 2019, pp. 1–7.
- [61] P. He, J. Gao, and W. Chen, "DeBERTaV3: Improving deberta using electra-style pre-training with gradient-disentangled embedding sharing," in *Proc. 11th Int. Conf. Learn. Representations*, 2023, pp. 1–16. [Online]. Available: <https://openreview.net/forum?id=sE7-XhLxHA>
- [62] T. Wolf et al., "Huggingface's transformers: State-of-the-art natural language processing," 2020, *arXiv:1910.03771*.
- [63] S. Khorram, M. G. McInnis, and E. M. Provost, "Jointly aligning and predicting continuous emotion annotations," *IEEE Trans. Affect. Comput.*, vol. 12, no. 4, pp. 1069–1083, Oct.–Dec. 2021.
- [64] J. Han, Z. Zhang, M. Schmitt, M. Pantic, and B. Schuller, "From hard to soft: Towards more human-like emotion recognition by modelling the perception uncertainty," in *Proc. 25th ACM Int. Conf. Multimedia*, 2017, pp. 890–897.
- [65] B. T. Atmaja and M. Akagi, "Evaluation of error-and correlation-based loss functions for multitask learning dimensional speech emotion recognition," *J. Phys., Conf. Ser.*, vol. 1896, no. 1, 2021, Art. no. 012004.
- [66] P. J. Rousseeuw, "Silhouettes: A graphical aid to the interpretation and validation of cluster analysis," *J. Comput. Appl. Math.*, vol. 20, pp. 53–65, 1987.
- [67] R. Chevi and A. Aji, "Daisy at WASSA 2024 empathy and personality shared task: A quick exploration on emotional pattern of empathy and distress," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. De Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, Aug. 2024, pp. 420–424. [Online]. Available: <https://aclanthology.org/2024.wassa-1.37>
- [68] P. Pereira, H. Moniz, and J. P. Carvalho, "ConText at WASSA 2024 empathy and personality shared task: History-dependent embedding utterance representations for empathy and emotion prediction in conversations," in *Proc. 14th Workshop Comput. Approaches Subjectivity, Sentiment, Social Media Anal.*, O. De Clercq, V. Barriere, J. Barnes, R. Klinger, J. Sedoc, and S. Tafreshi, Eds., Bangkok, Thailand, 2024, vol. 8, pp. 448–453. [Online]. Available: <https://aclanthology.org/2024.wassa-1.42>
- [69] S. Ouyang, J. M. Zhang, M. Harman, and M. Wang, "An empirical study of the non-determinism of ChatGPT in code generation," *ACM Trans. Softw. Eng. Methodol.*, vol. 34, pp. 1–28, 2025, doi: [10.1145/3697010](https://doi.org/10.1145/3697010).
- [70] K. Krippendorff, *Reliability*. Thousand Oaks, CA, USA: Sage Publications, 2019.
- [71] J. Wei et al., "Chain-of-thought prompting elicits reasoning in large language models," in *Proc. Adv. Neural Inf. Process. Syst.*, S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, Eds., 2022, vol. 35, pp. 24824–24837. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2022/file/9d5609613524ecf4f15af0f7b31abca4-Paper-Conference.pdf