

DepSy: A Dataset and Benchmark for *Depression* Symptoms Detection with Hierarchical Transformers and Fine-Tuned LLMs

Anonymous ACL submission

Abstract

Early detection of symptoms of depression can help minimise its impact on people suffering from depression. Social media, where users often share emotions and life experiences, offers a valuable resource for NLP-driven mental health research. We posit that mining social media posts enables researchers to identify clinically significant depressive symptoms. This paper introduces: a) the DepSy dataset, a novel resource annotated by psychologists for depressive symptoms, containing over 40k posts; and b) the DepSy model, a fine-tuned model trained to identify and extract depressive symptoms. We conducted comparative experiments between BERT-based models and large language models (LLMs) for symptom extraction. Our results show that both BERT-based models and LLMs demonstrated comparable performance, with BERT achieving the highest overall f-1 score of 0.522.

a rich resource for NLP-driven mental health research. However, identifying high-quality datasets for training models to monitor depressive symptoms remains a significant challenge (Gadzama et al., 2024).

According to recent surveys (Montejo-Ráez et al., 2024; Garg, 2023), most existing studies utilizing NLP on social media data have centred around detecting the presence or absence of depression and identifying shifts from depression to suicidal ideation (De Choudhury et al., 2016; Gong et al., 2019; Sawhney et al., 2020; Kour and Gupta, 2022; Baghdadi et al., 2022; Khafaga et al., 2023; Adarsh et al., 2023). However, to date, little attention has been given to examining the occurrence of depressive symptoms. This study aims to fill this identified gap by enhancing early detection and intervention strategies through improved datasets and model development. Our key contributions can be summarized as follows:

1 Introduction

Mental health issues are rising globally, with WHO estimating that one in four people will experience a condition in their lifetime (World Health Organization, 2001). These conditions significantly impact the quality of life and contribute to disability and high suicide rates (DSM5, 2013). Understanding symptoms is essential for advancing mental health research and developing effective models to support both individuals with mental health conditions and those with similar experiences.

Most NLP research on mental health monitoring has focused on electronic health records and diagnostic assessments (Kim et al., 2020; Pradier et al., 2021; Mesbah et al., 2021; de Oliveira et al., 2021; Ignashina et al., 2025). While these provide valuable insights, social media offers an alternative by capturing large-scale, real-time expressions of emotions and mental states. With billions of users sharing experiences online, social media data presents

- DepSy Dataset: The first English dataset of depression symptoms in textual posts of users who self-reported being diagnosed with depression that is **fully annotated by psychologists**¹.
- Hier-DepSy: a BERT-based hierarchical model architecture for depression symptom classification from social media data.
- DepSyLlama: a fine-tuned LLM for identifying depression symptoms
- Empirical work comparing multiple predictive models (based on BERT, RoBERTa, MentalBERT, GPT, Llama 2, Llama 3, MentalLlama) built using our dataset for the task of classifying/extracting depression symptoms from posts.

¹The DepSy dataset, DepSy model, and code will be made available upon paper acceptance

2 Related Work

2.1 Datasets for Depression Monitoring on Social Media

Several studies have developed social media datasets for depression analysis, emphasising the need for labelled data. [Kabir et al. \(2023\)](#) introduced a dataset categorising tweets as "non-depressed" or "depressed," with severity levels. However, reliance on symptom-related keywords may exclude relevant posts, and crowdworker annotations can lack contextual depth. The PRIMATE dataset ([Gupta et al., 2022](#)), based on PHQ-9 responses from Reddit, raises similar concerns. The lack of expertise among crowdworkers may lead to inaccuracies. [Milintsevich et al. \(2024\)](#) later re-annotated PRIMATE, finding errors and false positives, highlighting the need for more rigorous annotation processes to ensure high-quality mental health datasets. To address these limitations, our work involves constructing a large dataset fully annotated by expert psychologists, with the goal of achieving a high level of annotation agreement.

2.2 Monitoring Depression Symptoms through Social Media

Several studies have focused on monitoring depression through social media platforms such as X (Twitter) and Reddit ([Aragon et al., 2023](#); [Shah et al., 2020](#); [Tavast et al., 2022](#); [Zogan et al., 2022](#); [Chiong et al., 2021a,b](#); [Wang et al., 2022](#)). These studies typically annotate posts using matching terms, which can result in data loss or inaccurate labelling, as the context in which these terms are used may often be sarcastic ([Ezerceli and Dehkharghani, 2024](#); [Pavlova and Berkers, 2022](#)). This limitation highlights the challenges inherent in relying solely on keyword matching for sentiment analysis in sensitive mental health contexts. Several studies have adopted the PHQ-9 framework to detect depressive symptoms in social media. Early work by [Mowery et al. \(2016\)](#) identified three PHQ-9 symptoms using a two-stage classifier, laying groundwork for symptom-level annotation. [Yazdavar et al. \(2017\)](#) used a labelled LDA model on tweets from keyword-identified users to detect nine symptoms, but relied on non-expert annotations and showed performance disparities across symptoms. More recent work by [Yadav et al. \(2020\)](#) proposed a BERT-based multi-task model incorporating figurative language, though keyword-based data collection limited generalisability. [Yadav et al. \(2023\)](#)

introduced RESTORE, a multimodal dataset of annotated memes, showing improvements with orthogonal constraints across modalities. To improve cross-domain robustness, [Nguyen et al. \(2022\)](#) grounded predictions in PHQ-9 descriptions, enhancing interpretability, though performance may be limited for contextually subtle symptoms. In a large-scale study, [Liu et al. \(2023\)](#) used Reddit data and RoBERTa to detect 13 expert-validated symptoms, achieving strong results on an external benchmark, though subreddit-based labelling may limit clinical validity. Other studies employed the Beck Depression Inventory (BDI) for symptom estimation, including DepressMind ([Fernández-Iglesias et al., 2024](#)), which uses sentence similarity, and a prompting-based method by [Aragón et al. \(2024\)](#) that maps user posts to BDI questions via ChatGPT. While promising, both rely on surface-level or model-generated inferences rather than clinical annotation. Accurate symptom detection requires modelling symptoms and involving domain experts in annotation. However, few existing works offer clinically annotated, large-scale, multi-symptom datasets with strong modelling baselines, leaving a gap that this study aims to address.

3 DepSy Dataset

Recognizing the critical importance of precise and expert-driven annotations in the study of mental health through natural language processing (NLP), we undertook a rigorous annotation process guided by a robust annotation scheme. We annotated a dataset originally collected by [Alhamed et al. \(2024b\)](#) of users who self-reported being clinically diagnosed with depression on X (formerly Twitter). We annotated the posts of users in the class "after" being diagnosed with depression. The annotation process followed the annotation scheme in ([Alhamed et al., 2024a](#)) that was designed for annotating posts for depressive symptom and severity, and where high annotation agreement was reported.

The depressive symptoms in this scheme were synthesised from well-established and validated depression assessment tools, including the PHQ-9 ([Kroenke et al., 2001](#)), BDI ([Beck et al., 1961](#)), and CES-D ([Radloff, 1977](#)) questionnaires; the final symptoms list is shown in Table 1. Experienced psychologists meticulously annotated each post in the dataset for the mentioned list of depression symptoms. This annotation process has resulted in the creation of the first and largest dataset fully an-

Depression Symptoms
Poor appetite or eating disturbances
Feeling down and depressed
Crying
Concentration problems
Feeling tired or having little energy
Feelings of failure
Sleep disturbances
Loss of interest
Self-blame and shame
Loneliness
Suicidal thoughts

Table 1: List of depression symptoms

notated by professional psychologists, comprising over 40,000 posts annotated for depression symptoms. Each post is annotated to indicate whether it contains no symptoms, one symptom, or multiple symptoms, with the corresponding symptom name(s) included in the annotation. This resource is intended to support further research and model development in the detection and analysis of depression within NLP applications².

3.1 Data Annotation

The annotation task was carried out by five psychologists (co-authors), each with at least three years of specialised experience in diagnosing depression and/or anxiety disorders. Their involvement in this task was entirely voluntary and driven by a shared commitment to advancing mental health research. The annotators completed consent forms, which are securely stored on the college’s OneDrive servers, and were provided with information sheets and detailed annotation guidelines. They were asked to select all (if any) depressive symptoms that existed in a post from an existing list of symptoms. We used Labelstudio³ as a labelling interface for all experiments in this work. Within Label Studio, we designed a custom labelling interface to meet the specific needs of our task, as none of the available templates offered a suitable match. To evaluate the consistency of our annotations, we utilised Cohen’s kappa (κ), a widely recognised statistic for assessing inter-rater reliability on nominal data (Cohen, 1960). Cohen’s kappa accounts for agreement occurring by chance, thus offering a more rigorous measure of concordance than simple percentage agreement. We obtained a pairwise kappa score of 0.67 across 10% of annotated posts. According to the interpretation scale proposed by Landis

²The dataset will be made available upon paper acceptance

³<https://labelstud.io/>

and Koch (1977), this score falls within the “substantial” agreement range. The 10% subset was selected to ensure robust and reliable results, exceeding the smaller subsets (20–100 posts) commonly used in similar studies (e.g., Chancellor et al. (2021); Harrigan and Dredze (2022)).

4 Data Analysis

This section presents an analysis of depression symptom prevalence within the corpus of social media posts. By examining the frequency distribution of symptoms, we aim to identify the most commonly occurring depressive symptoms that users are willing to expose online. This analysis enables us to gain deeper insights into the dynamic nature of depressive symptoms and potential trajectories of the condition.

4.1 Symptoms Co-morbidity

Number of Symptoms	Number of Posts
0 (No Symptoms)	37,335
1	2,080
2	407
3	78
4	9
5	1

Table 2: Distribution of symptom co-occurrence in DepSy dataset.

4.2 Depressive Symptoms Frequency

When we analysed the frequency distribution of depressive symptoms within the users’ posts, feeling down or depressed emerged as the most prevalent symptom, followed by feeling tired or having little energy, and crying. Symptoms such as self-blame and suicidal ideation were reported to be the least frequent. Figure 1 illustrates the frequency distribution of symptoms in the posts.

We further analysed the distribution of depressive symptoms in the dataset to understand their prevalence and co-occurrence. As shown in Table 2, the dataset is highly imbalanced, with the majority of posts (37,335) containing no symptoms and only 2,575 posts labelled with one or more symptoms. Most of the symptom-labelled posts contain a single symptom, while posts with multiple symptoms are increasingly rare. This imbalance is expected to pose a challenge for model training.

Figure 2 shows the co-morbidity matrix of depressive symptoms, capturing how often symptom

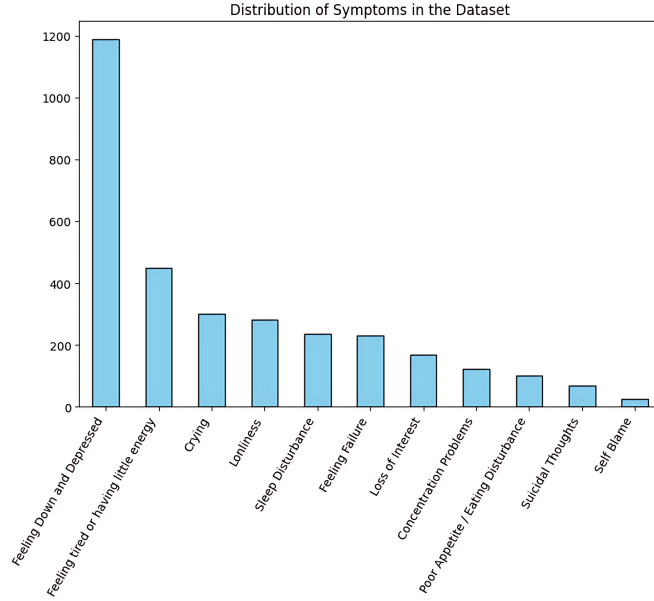


Figure 1: depressive symptom frequency in posts

pairs co-occur. Diagonal values reflect individual symptom frequency, while off-diagonal values indicate joint occurrences. “Feeling Down and Depressed” is the most frequent symptom and co-occurs frequently with others such as “Feeling Failure” and “Feeling tired or having little energy,” suggesting a core cluster. In contrast, symptoms like “Self Blame” and “Suicidal Thoughts” appear less often and with weaker co-morbidity. These patterns highlight inter-symptom dependencies relevant for symptom-level classification.

5 Task Definition

Given a user-generated post P_i , the objective is to predict a binary label for each depressive symptom from the predefined set of 11 symptoms. The task is multi-label multi-class classification and is formulated as:

$$Y_i = f(P_i; \theta), \quad (1)$$

$$Y_i = \{y_{i,1}, y_{i,2}, \dots, y_{i,11}\}, \quad y_{i,j} \in \{0, 1\}$$

Where, P_i represents an individual user-generated post, and $f(\cdot)$ denotes the classification model parameterized by θ . The output Y_i corresponds to the predicted set of binary labels for the 11 depressive symptoms. Each element $y_{i,j}$ in Y_i indicates the binary label for symptom j in post P_i .

6 Hier-DepSy: A Hierarchical Model for Classifying Unbalanced Data

We introduce *Heir-DepSy*, a hierarchical model architecture for automated depression symptom clas-

sification from social media data. When examined more thoroughly, the task of symptom classification can be naturally divided into two sequential components. The first component involves determining whether a post expresses any depressive symptom. The second, conditional on the first, involves identifying which specific symptoms are present. Heir-DepSy explicitly models this structure using two successive classification stages describe in Section 6.2. In addition to aligning with the task’s inherent structure, this approach also addresses the significant class imbalance in the dataset—where a large proportion of posts are non-symptomatic and several symptoms are underrepresented.

6.1 Model Selection and Training

Multiple pre-trained transformer models, including BERT, RoBERTa, and MentalBERT, were evaluated independently for each classification stage. For each task, models were trained and evaluated separately, and the best-performing model for each stage was selected to construct the final Heir-DepSy architecture. The models performance on each stage can be found in Table 3

Based on the results, **BERT** was selected for the binary classification task and **RoBERTa** was selected for multi-label classification, as they achieved the best micro-averaged F1-score and handled symptoms more effectively.

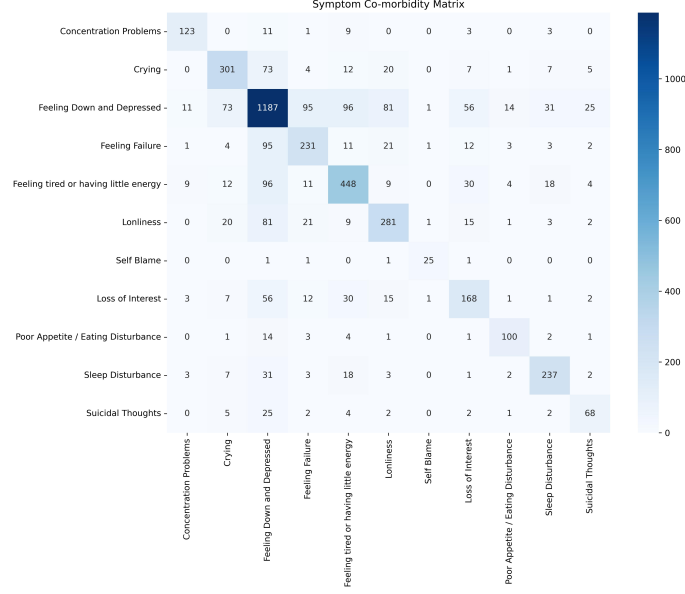


Figure 2: depressive symptom frequency in posts

	Model	Accuracy	Precision	Recall	F-1
Binary Classification	BERT	0.83	0.80	0.80	0.80
	RoBERTa	0.83	0.79	0.80	0.79
	MentalBERT	0.82	0.78	0.77	0.78
Multi-Label Symptom Classification	BERT	0.47	0.71	0.59	0.64
	RoBERTa	0.49	0.67	0.65	0.66
	MentalBERT	0.47	0.70	0.59	0.64

Table 3: Results for models on classifying depressive symptoms, preparing for Hier-DepSy model. Precision, recall, and F1 are micro-average scores.

6.2 Hier-DepSy Model Architecture

Hier-DepSy is implemented as a two-stage hierarchical classification model composed of two independently trained transformer-based classifiers. As we mentioned earlier, the first stage detects whether a post expresses any depressive symptom (binary classification), and the second stage, triggered only for positive cases, identifies the specific symptoms present (multi-label classification).

7 DepSyLlama Model

7.1 DepSy Instruction Dataset

The DepSy Instruction dataset is constructed using all posts from the raw dataset introduced in Section 3, along with the selected prompts. The symptoms listed in the "symptoms" columns are directly utilized as responses to the corresponding questions. To create the training portion of the DepSy Instruction dataset, we merge the prompt, post, and response into a single text. For optimal model selection, a test set of 10% of the DepSy dataset developed employing the same methodol-

Algorithm 1: Heir-DepSy Model Architecture

Input: Input post x
Output: Predicted symptom vector $\hat{y}_{\text{symptom}} \in \{0, 1\}^{11}$

```

// Stage 1: Binary Classification
 $x_{\text{tok}} \leftarrow \text{BERT\_Tokenizer}(x, \text{max\_len}=512)$ ;
 $[CLS]_{\text{bert}} \leftarrow \text{BERT\_Encoder}(x_{\text{tok}})$ ;
 $z_{\text{binary}} \leftarrow \text{LinearLayer\_Binary}([CLS]_{\text{bert}})$ ;
 $b \leftarrow \arg \max(z_{\text{binary}})$ ; // Binary prediction
if  $b = 0$  then
     $\hat{y}_{\text{symptom}} \leftarrow [0, 0, \dots, 0]$ ; // No symptoms
else
    // Stage 2: Multi-label Classification
     $x_{\text{tok}} \leftarrow \text{RoBERTa\_Tokenizer}(x, \text{max\_len}=512)$ ;
     $[CLS]_{\text{roberta}} \leftarrow \text{RoBERTa\_Encoder}(x_{\text{tok}})$ ;
     $z_{\text{multi}} \leftarrow \text{LinearLayer\_MultiLabel}([CLS]_{\text{roberta}})$ ;
     $p \leftarrow \sigma(z_{\text{multi}})$ ; // Apply sigmoid
     $\hat{y}_{\text{symptom}}[i] \leftarrow 1$  if  $p[i] \geq 0.5$  else 0,  $\forall i \in \{1, \dots, n\}$ 
return  $\hat{y}$ 

```

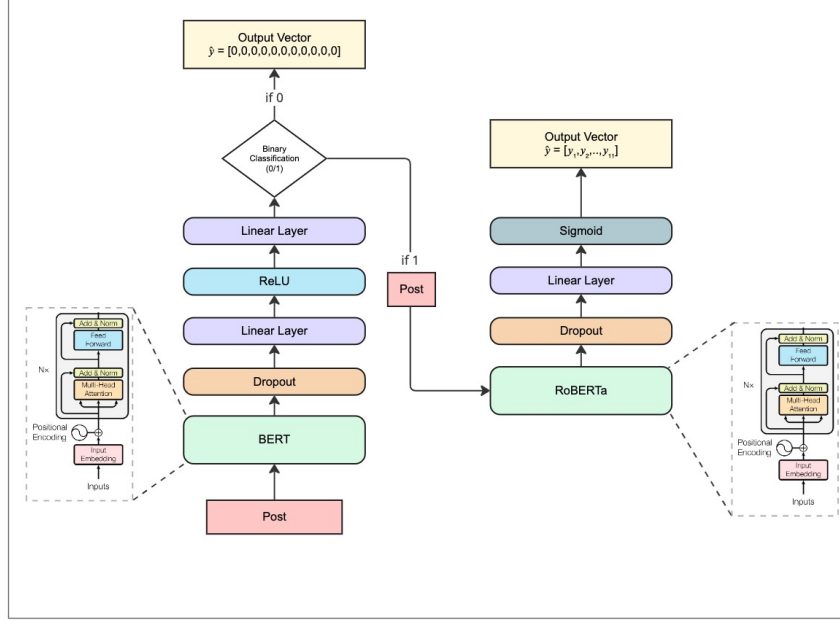


Figure 3: Architecture of the Hier-DepSy model for depression symptoms classification. The model operates in two independent stages without joint training.

ogy was used to test the model.

7.2 DepSyLlama Model Training

We fine-tune the LLaMA models on the DepSy Instruction dataset to develop our DepSyLlama model. Specifically, we constructed 2 versions of DepSy by training LLaMA2-7B and Llama3-8B on the DepSy training set for 16 epochs, selecting the optimal model based on performance on the DepSy validation set. The training process employs a batch size of 8 and is optimized using the AdamW optimizer, with a maximum learning rate of $1e-5$. The maximum input length for the model is set to 4096 tokens. To expedite the training process, we utilize QLoRA, configuring the LoRA rank to 8 and the alpha parameter to 16. All models are trained on Nvidia Tesla A100 GPUs, each with 40GB of memory.⁴ For model inference we used Depsy with the settings (temperature = 0.2, max_new_tokens = 30)

8 Experiments and Results

We evaluate a range of models for our classification tasks, including BERT, RoBERTa, MentalBERT, and large language models (LLMs). Full model details are provided in the appendix. We explore the use of LLMs in zero-shot, zero-shot with labels, and few-shot settings using various prompt formats.

⁴The model will be made publicly available on Hugging Face upon paper acceptance

The specific prompts used in our experiments are detailed in Appendix A, Table 6.

8.1 Evaluation Metrics

For BERT-based models, we used 5-fold cross-validation to evaluate performance based on accuracy, micro-averaged precision, recall, and F1 scores. Our implementation utilizes Scikit-learn (Pedregosa et al., 2011). To evaluate the performance of large language models (LLMs), we assessed the presence of each symptom label within the generated responses by identifying the corresponding text spans. For the "no symptoms" label, we verified the presence of the phrase "does not contain symptoms" as this was introduced in the prompt when no symptoms were indicated in the post. After this step, we got 11 predicted columns (11 symptoms) and 11 true labels. For evaluating the model's output, we used Scikit-learn's multilabel confusion matrix and extracted micro-averaged precision, recall, F1-score, and accuracy. Additionally, we performed 1,000 bootstrap resamples to compute 95% confidence intervals (CI) for these metrics. DepSyLlama was evaluated using 3-fold cross-validation.⁵

8.2 Results

Table 4 presents the performance comparison of various models on the task, evaluating accuracy,

⁵The code will be made available upon paper acceptance

	Model	Accuracy	Precision	Recall	F-1	CI 95% F-1
	BERT	0.711	0.569	0.482	0.522	[0.514, 0.529]
	RoBERTa	0.686	0.505	0.501	0.501	[0.488, 0.512]
	MentalBERT	0.701	0.546	0.483	0.512	[0.500, 0.524]
	Hier-DepSy	0.845	0.668	0.679	0.673	[0.614, 0.702]
Zero shot	GPT-4o	0.527	0.714	0.033	0.064	[0.038, 0.091]
	GPT-3.5	0.527	0.704	0.032	0.061	[0.037, 0.086]
	MentalLlama_7b	0.526	0.418	0.068	0.118	[0.086, 0.149]
	DepSy Llama2	0.526	0.419	0.66	0.486	[0.47, 0.502]
	DepSy Llama3	0.313	0.370	0.494	0.423	[0.411, 0.435]
Zero shot with labels	GPT-4o	0.564	0.634	0.185	0.287	[0.245, 0.325]
	GPT-3.5	0.53	0.342	0.432	0.382	[0.352, 0.411]
	Llama2	0.479	0.116	0.099	0.107	[0.082, 0.132]
	Llama3	0.521	0.305	0.066	0.108	[0.079, 0.141]
	MentalLlama	0.446	0.099	0.15	0.12	[0.099, 0.142]
	Depsy Llama 2	0.196	0.14	0.62	0.226	[0.208, 0.244]
	Depsy Llama 3	0.425	0.138	0.543	0.198	[0.177, 0.218]
Few shots	GPT-4o	0.612	0.527	0.492	0.509	[0.475, 0.543]
	GPT-3.5	0.476	0.309	0.574	0.402	[0.375, 0.429]
	Llama2	0.397	0.119	0.135	0.126	[0.103, 0.15]
	Llama3	0.422	0.13	0.197	0.157	[0.132, 0.181]
	MentalLlama	0.168	0.121	0.457	0.191	[0.173, 0.21]
	DepSy Llama2	0.360	0.259	0.379	0.308	[0.29, 0.325]
	DepSy Llama3	0.168	0.121	0.457	0.191	[0.173, 0.21]

Table 4: Results for models on classifying depressive symptoms using DepSy dataset. Precision, recall, and F1 are micro-average scores.

Symptom	BERT			RoBERTa			MentalBERT			Hier-DepSy		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
Feeling Down and Depressed	0.57	0.54	0.55	0.53	0.60	0.56	0.55	0.53	0.54	0.64	0.76	0.70
Feeling tired or having little energy	0.49	0.46	0.47	0.64	0.50	0.56	0.56	0.43	0.48	0.60	0.60	0.60
Crying	0.67	0.75	0.71	0.62	0.72	0.67	0.64	0.89	0.74	0.74	0.91	0.82
Loneliness	0.61	0.23	0.33	0.66	0.37	0.47	0.46	0.27	0.34	0.65	0.71	0.68
Sleep Disturbance	0.69	0.72	0.71	0.67	0.72	0.70	0.67	0.70	0.68	0.78	0.90	0.84
Feeling Failure	0.29	0.11	0.16	0.18	0.20	0.19	0.38	0.13	0.20	0.48	0.52	0.50
Loss of Interest	0.27	0.23	0.25	0.14	0.16	0.15	0.17	0.16	0.17	0.36	0.33	0.35
Concentration Problems	0.32	0.26	0.29	0.15	0.13	0.14	0.28	0.22	0.24	0.50	0.55	0.52
Poor Appetite / Eating Disturbance	0.56	0.65	0.60	0.46	0.52	0.49	0.58	0.61	0.60	0.70	0.78	0.74
Suicidal Thoughts	1	0.33	0.50	1	0.50	0.67	0.62	0.42	0.50	0.62	0.80	0.70
Self Blame	0	0	0	0	0	0	0	0	0	0	0	0

Table 5: Per-symptom precision, recall, and F1-score across four models. Symptoms are sorted by their frequency in the dataset.

precision, recall, and F1-score with 95% confidence intervals. Among all models, our Hier-Depsy model achieved the highest F1-score (0.673) and accuracy (0.845), outperforming BERT, RoBERTa, and MentalBERT and LLMs. In the LLMs zero-shot setting, GPT-4o showed the highest precision (0.714), but with extremely low recall (0.033), resulting in a low F1-score (0.064). When provided with labels, GPT-4o’s performance improved substantially, achieving an F1-score of 0.385, suggesting the benefit of added contextual guidance for general-purpose LLMs. Among few-shot models, GPT-4o again outperformed others with an F1-score of 0.475 and the highest accuracy (0.612), approaching the performance of fine-tuned models.

Overall, LLaMA models showed limited effectiveness across all settings, as reflected by low precision, recall, and F1-scores. This trend supports prior findings (Ignashina et al., 2025), indicating their difficulty in recognising nuanced mental health symptoms. Given the importance of recall in depression screening—where missing symptoms can have serious implications—models with higher recall are more appropriate for this task (Ren et al., 2015). Depsy Llama models consistently demonstrated strong recall across settings. In the zero-shot setting, Depsy Llama 2 achieved the highest recall (0.66), while in the few-shot setting, Depsy Llama 3 reached a recall of 0.494. Although these models had low precision, their ability to detect a

wider range of depressive symptoms makes them useful for early-stage mental health assessment.

Table 5 reports per-symptom precision, recall, and F1-scores across four models, with symptoms sorted by frequency. Notably, higher frequency does not always correspond to better performance. Symptoms like *Crying* and *Sleep Disturbance* achieved the highest F1-scores (up to 0.82 and 0.84 with Hier-DepSy), likely due to their more explicit linguistic expression. In contrast, abstract symptoms such as *Loss of Interest* or *Feeling Down and Depressed* showed lower performance, despite being more frequent. Rare symptoms like *Self-Blame* remained difficult to detect across all models. Hier-DepSy consistently outperformed baselines, particularly on less frequent symptoms, highlighting its effectiveness in handling imbalanced and subtle symptom classes.

9 Discussion

The Hier-DepSy model outperformed all other approaches in our symptom classification task, demonstrating the value of a hierarchical structure for handling task complexity and class imbalance. By decoupling binary detection from multi-label classification, the model was better able to capture symptom-specific patterns. Similar layered strategies have proven effective in related domains, reinforcing the broader applicability of hierarchical approaches in psychological NLP tasks.

Large language models (LLMs), in contrast, underperformed compared to fine-tuned BERT-based models. This aligns with previous findings in health-related NLP and may be partly due to embedded safety alignment, which can lead to conservative outputs when handling sensitive content. While we did not observe explicit refusals or safety prompts, implicit alignment may still have influenced prediction behaviour. However, per-symptom confusion matrix analysis (Appendix C) did not reveal systematic avoidance of sensitive symptoms such as suicidal ideation.

DepSyLLaMA, fine-tuned on our dataset using QLoRA, did not outperform general LLMs. This may be due to limitations in parameter-efficient fine-tuning, which updates only a subset of parameters, reducing the model’s ability to adapt fully to the task. Additionally, the limited size of posts with symptoms and class imbalance likely affected generalisation. Interestingly, DepSyLLaMA performed best in the pure zero-shot setting, sug-

gesting that fine-tuning may reduce reliance on in-context prompts and thus reduce the size of inference.

Finally, domain-specific models like Mental-LLaMA and MentalBERT did not outperform general-purpose models. This suggests that pre-training on mental health data alone is insufficient for accurate symptom classification. These models may have been tuned for empathetic dialogue rather than multi-label detection. Our findings underscore the need for task-specific fine-tuning strategies that directly optimise for distinguishing nuanced symptom categories, beyond domain relevance alone.

To further investigate model limitations, we conducted detailed error analysis across symptoms and models. Table 5 shows that symptoms such as Crying, Sleep Disturbance, and Feeling Down and Depressed achieved consistently high F1-scores, likely due to their prevalence and clearer linguistic cues. In contrast, low-frequency symptoms like Loss of Interest and Self Blame were rarely detected, with some models failing to predict them at all. This highlights the effect of data imbalance and the need for symptom-specific learning strategies. Confusion matrix analysis of LLMs (Appendix C) showed no strong evidence of safety-driven suppression of sensitive symptoms, though underprediction was observed for abstract or less explicitly stated symptoms such as Loneliness and Feeling Failure. These findings reinforce the strengths of hierarchical models in managing imbalance and the limitations of LLMs in capturing subtle symptom expressions without direct optimisation.

10 Conclusions

This paper addressed the task of detecting depressive symptoms in social media posts as a multi-label classification problem. We evaluated transformer-based models, large language models (LLMs), and a hierarchical architecture. The proposed Hier-DepSy model improved detection of low-frequency symptoms by mitigating class imbalance through a two-stage structure. We also introduced DepSyLLaMA, a domain-adapted LLM, which outperformed general-purpose open-source LLMs in zero-shot settings and showed competitive results compared to proprietary models. Despite these advances, symptom detection remains a challenging task, with the best model achieving an F1-score of 0.673, highlighting the need for further research in this area.

11 Limitations

Despite the novel contributions of this study, several limitations should be acknowledged. First, while our model was developed using the DepSy dataset, which is substantial in size (over 40,000 entries) compared to existing datasets and believed to be robust, we aimed to further evaluate its generalizability by testing it on additional datasets. However, our request for access to the SAD depressive symptoms dataset (Mowery et al., 2017), which would have enabled broader validation, did not receive a response from its owners. This limited our ability to assess the model’s generalizability across diverse data sources. Secondly, while the DepSy dataset is carefully annotated by expert psychologists, it relies on publicly available social media posts, which may not fully capture the diversity of individuals experiencing depression. Social media content often reflects self-presentation biases, potentially affecting symptom reporting and model generalizability.

Ethical Consideration

This study has received ethics approval from XXXXXX⁶ (Reference: 21IC7222). The dataset contains only publicly available posts from X, and we are committed to following ethical practices to protect the privacy and anonymity of the users. To ensure this, the author’s usernames, which could contain sensitive information related to the names or locations of the user, are not saved or used. Instead, the information was pre-processed and replaced with user IDs. Social media data is often sensitive, particularly when it is related to mental health, and we take great care to ensure that our dataset is handled responsibly. Since the dataset is related to mental disorders, it might trigger some people, thus, annotators were advised to take breaks during annotation and were given plenty of time.

References

V Adarsh, P Arun Kumar, V Lavanya, and GR Gadharan. 2023. Fair and explainable depression detection in social media. *Information Processing & Management*, 60(1):103168.

Falwah Alhamed, Rebecca Bendayan, Julia Ive, and Lucia Specia. 2024a. Monitoring depression severity

and symptoms in user-generated content: An annotation scheme and guidelines. In *Proceedings of the 14th Workshop on Computational Approaches to Subjectivity, Sentiment, & Social Media Analysis*, pages 227–233.

Falwah Alhamed, Julia Ive, and Lucia Specia. 2024b. Classifying social media users before and after depression diagnosis via their language usage: A dataset and study. In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 3250–3260.

Mario Aragon, Adrián Pastor López Monroy, Luis Gonzalez, David E Losada, and Manuel Montes. 2023. Disorbert: A double domain adaptation model for detecting signs of mental disorders in social media. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15305–15318.

Mario Ezra Aragón, Javier Parapar, and David E Losada. 2024. Delving into the depths: evaluating depression severity through bdi-biased summaries. In *9th Workshop on Computational Linguistics and Clinical Psychology (CLPsych 2024)*, pages 12–22. Association for Computational Linguistics.

Nadiah A Baghdadi, Amer Malki, Hossam Magdy Balaha, Yousry AbdulAzeem, Mahmoud Badawy, and Mostafa Elhosseini. 2022. An optimized deep learning approach for suicide detection through arabic tweets. *PeerJ Computer Science*, 8:e1070.

Aaron T Beck, Charles H Ward, Morris Mendelsohn, John Mock, and James Erbaugh. 1961. [An inventory for measuring depression](#). *Archives of General Psychiatry*, 4(6):561–571.

Stevie Chancellor, Steven A Sumner, Corinne David-Ferdon, Tahirah Ahmad, and Munmun De Choudhury. 2021. Suicide risk and protective factors in online support forum posts: annotation scheme development and validation study. *JMIR mental health*, 8(11):e24471.

Raymond Chiong, Gregorius Satia Budhi, and Sandeep Dhakal. 2021a. Combining sentiment lexicons and content-based features for depression detection. *IEEE Intelligent Systems*, 36(6):99–105.

Raymond Chiong, Gregorius Satia Budhi, Sandeep Dhakal, and Fabian Chiong. 2021b. A textual-based featuring approach for depression detection using machine learning classifiers and social media texts. *Computers in Biology and Medicine*, 135:104499.

Jacob Cohen. 1960. A coefficient of agreement for nominal scales. *Educational and Psychological Measurement*, 20(1):37–46.

Munmun De Choudhury, Emre Kiciman, Mark Dredze, Glen Coppersmith, and Mrinal Kumar. 2016. [Discovering shifts to suicidal ideation from mental health content in social media](#). *Conference on Human*

⁶masked for anonymity

611	<i>Factors in Computing Systems - Proceedings</i> , pages	Mohsinul Kabir, Tasnim Ahmed, Md Bakhtiar Hasan,	664
612	2098–2110.	Md Tahmid Rahman Laskar, Tarun Kumar Joarder,	665
		Hasan Mahmud, and Kamrul Hasan. 2023. Deptweet:	666
613	Joseigla Pinto de Oliveira, Karen Jansen, Taiane	A typology for social media texts to detect depres-	667
614	de Azevedo Cardoso, Thaíse Campos Mondin, Lu-	sion severities. <i>Computers in Human Behavior</i> ,	668
615	ciano Dias de Mattos Souza, Ricardo Azevedo	139:107503.	669
616	da Silva, and Fernanda Pedrotti Moreira. 2021. Pre-		
617	dictors of conversion from major depressive disorder	D Sami Khafaga, Maheshwari Auvdaiappan, K Deepa,	670
618	to bipolar disorder . <i>Psychiatry Research</i> , 297(Jan-	Mohamed Abouhawwash, and F Khalid Karim. 2023.	671
619	uary):113740.	Deep learning for depression detection using twitter	672
		data. <i>Intelligent Automation & Soft Computing</i> ,	673
620	DSM5. 2013. <i>Diagnostic and statistical manual of</i>	36(2):1301–1313.	674
621	<i>mental disorders : DSM-5</i> , 5th ed. edition. American		
622	Psychiatric Association Arlington, VA.	Kiana Kheiri and Hamid Karimi. 2023. Sentimentgpt:	675
		Exploiting gpt for advanced sentiment analysis and	676
623	Özay Ezerceci and Rahim Dehkharghani. 2024. Mental	its departure from current machine learning . <i>Preprint</i> ,	677
624	disorder and suicidal ideation detection from social	arXiv:2307.10234.	678
625	media using deep neural networks. <i>Journal of Com-</i>		
626	<i>putational Social Science</i> , 7(3):2277–2307.	Hyewon Kim, Yuwon Kim, Ji Hyun Baek, Maurizio	679
		Fava, David Mischoulon, Andrew A. Nierenberg,	680
627	Roque Fernández-Iglesias, Marcos Fernández-Pichel,	Kwan Woo Choi, Eun Jin Na, Myung Hee Shin, and	681
628	Mario Aragon, and David E Losada. 2024. Depress-	Hong Jin Jeon. 2020. Predictive factors of diagnostic	682
629	mind: A depression surveillance system for social	conversion from major depressive disorder to bipolar	683
630	media analysis. In <i>Proceedings of the 18th Confer-</i>	disorder in young adults ages 19–34: A nationwide	684
631	<i>ence of the European Chapter of the Association for</i>	population study in South Korea . <i>Journal of Affective</i>	685
632	<i>Computational Linguistics: System Demonstrations</i> ,	<i>Disorders</i> , 265(December 2019):52–58.	686
633	pages 35–43.		
634	Wadzani Aduwamai Gadzama, Danlami Gabi,	Harnain Kour and Manoj K Gupta. 2022. An hy-	687
635	Musa Sule Argungu, and Hassan Umar Suru. 2024.	brid deep learning approach for depression predic-	688
636	The use of machine learning and deep learning	tion from user tweets using feature-rich cnn and bi-	689
637	models in detecting depression on social media: A	directional lstm. <i>Multimedia Tools and Applications</i> ,	690
638	systematic literature review. <i>Personalized Medicine</i>	81(17):23649–23685.	691
639	<i>in Psychiatry</i> , 45:100125.		
640	Muskan Garg. 2023. Mental health analysis in social	Kurt Kroenke, Robert L Spitzer, and Janet B Williams.	692
641	media posts: a survey. <i>Archives of Computational</i>	2001. The phq-9: validity of a brief depression sever-	693
642	<i>Methods in Engineering</i> , 30(3):1819–1842.	ity measure. <i>Journal of General Internal Medicine</i> ,	694
		16(9):606–613.	695
643	Jue Gong, Gregory E. Simon, and Shan Liu. 2019. Ma-	J Richard Landis and Gary G Koch. 1977. The mea-	696
644	chine learning discovery of longitudinal patterns of	surement of observer agreement for categorical data.	697
645	depression and suicidal ideation . <i>PLoS ONE</i> , 14(9):1–	<i>Biometrics</i> , 33(1):159–174.	698
646	15.		
647	Shrey Gupta, Anmol Agarwal, Manas Gaur, Kaushik	Tingting Liu, Devansh Jain, Shivani R Rapole, Brenda	699
648	Roy, Vignesh Narayanan, Ponnurangam Kumaraguru,	Curtis, Johannes C Eichstaedt, Lyle H Ungar, and	700
649	and Amit Sheth. 2022. Learning to automate follow-	Sharath Chandra Guntuku. 2023. Detecting symp-	701
650	up question generation using process knowledge for	toms of depression on reddit. In <i>Proceedings of the</i>	702
651	depression triage on Reddit posts . In <i>Proceedings of</i>	<i>15th ACM web science conference 2023</i> , pages 174–	703
652	<i>the Eighth Workshop on Computational Linguistics</i>	183.	704
653	<i>and Clinical Psychology</i> , pages 137–147, Seattle,		
654	USA. Association for Computational Linguistics.	Rahele Mesbah, Nienke de Bles, Nathaly Rius-	705
		Ottenheim, A. J. Willem van der Does, Brenda W.J.H.	706
655	Keith Harrigan and Mark Dredze. 2022. Then and	Penninx, Albert M. van Hemert, Max de Leeuw,	707
656	now: Quantifying the longitudinal validity of self-	Erik J. Giltay, and Manja Koenders. 2021. Anger	708
657	disclosed depression diagnoses . In <i>Proceedings of</i>	and cluster B personality traits and the conversion	709
658	<i>the Eighth Workshop on Computational Linguistics</i>	from unipolar depression to bipolar disorder . <i>Depres-</i>	710
659	<i>and Clinical Psychology</i> , pages 59–75, Seattle, USA.	<i>sion and Anxiety</i> , (August 2020):1–11.	711
660	Association for Computational Linguistics.		
		Kirill Milintsevich, Kairit Sirts, and Gaël Dias. 2024.	712
661	Mariia Ignashina, Paulina Bondaronek, Dan Santel,	Your model is not predicting depression well and	713
662	John Pestian, and Julia Ive. 2025. Llm assistance for	that is why: A case study of primate dataset. <i>arXiv</i>	714
663	pediatric depression . <i>Preprint</i> , arXiv:2501.17510.	<i>preprint arXiv:2403.00438</i> .	715
		Arturo Montejo-Ráez, M Dolores Molina-González,	716
		Salud María Jiménez-Zafra, Miguel Ángel García-	717
		Cumbreras, and Luis Joaquín García-López. 2024.	718
		A survey on detecting mental disorders with natural	719

720	language processing: Literature review, trends and	Mikke Tavast, Anton Kunnari, and Perttu Hämäläinen.	777
721	challenges. <i>Computer Science Review</i> , 53:100654.	2022. Language models can generate human-like	778
722		self-reports of emotion . In <i>27th International Confer-</i>	779
723	Danielle Mowery, Hilary Smith, Tyler Cheney, Greg	<i>ence on Intelligent User Interfaces</i> , IUI '22 Compan-	780
724	Stoddard, Glen Coppersmith, Craig Bryan, Mike	ion, page 69–72, New York, NY, USA. Association	781
725	Conway, et al. 2017. Understanding depressive symp-	for Computing Machinery.	782
726	toms and psychosocial stressors on twitter: a corpus-		
727	based study. <i>Journal of medical Internet research</i> ,		
728	19(2):e6895.	Wei-Yao Wang, Yu-Chien Tang, Wei-Wei Du, and Wen-	783
729		Chih Peng. 2022. Nycu_twd@ It-edu-acl2022: En-	784
730	Danielle L. Mowery, Albert Park, Craig Bryan, and	semble models with vader and contrastive learning	785
731	Mike Conway. 2016. Towards automatically clas-	for detecting signs of depression from social me-	786
732	sifying depressive symptoms from Twitter data for	dia. In <i>Proceedings of the second workshop on lan-</i>	787
733	population health . In <i>Proceedings of the Workshop</i>	<i>guage technology for equality, diversity and inclu-</i>	788
734	<i>on Computational Modeling of People's Opinions,</i>	<i>sion</i> , pages 136–139.	789
735	<i>Personality, and Emotions in Social Media (PEO-</i>		
736	<i>PLES)</i> , pages 182–191, Osaka, Japan. The COLING	World Health Organization. 2001. The world health re-	790
737	2016 Organizing Committee.	port 2001: Mental disorders affect one in four people.	791
738		Accessed: 2024-10-05.	792
739	Thong Nguyen, Andrew Yates, Ayah Zirikly, Bart		
740	Desmet, and Arman Cohan. 2022. Improving	Shweta Yadav, Cornelia Caragea, Chenye Zhao, Naincy	793
741	the generalizability of depression detection by	Kumari, Marvin Solberg, and Tanmay Sharma. 2023.	794
742	leveraging clinical questionnaires. <i>arXiv preprint</i>	Towards identifying fine-grained depression symp-	795
743	<i>arXiv:2204.10432</i> .	toms from memes . In <i>Proceedings of the 61st Annual</i>	796
744		<i>Meeting of the Association for Computational Lin-</i>	797
745	Alina Pavlova and Pauwke Berkers. 2022. “mental	<i>guistics (Volume 1: Long Papers)</i> , pages 8890–8905.	798
746	health” as defined by twitter: Frames, emotions,		
747	stigma. <i>Health communication</i> , 37(5):637–647.	Shweta Yadav, Jainish Chauhan, Joy Prakash Sain, Kr-	799
748		ishnaprasad Thirunarayan, Amit Sheth, and Jeremiah	800
749	Fabian Pedregosa, Gaël Varoquaux, Alexandre Gram-	Schumm. 2020. Identifying depressive symp-	801
750	fort, Vincent Michel, Bertrand Thirion, Olivier Grisel,	toms from tweets: Figurative language enabled	802
751	Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vin-	multitask learning framework. <i>arXiv preprint</i>	803
752	cent Dubourg, et al. 2011. Scikit-learn: Machine	<i>arXiv:2011.06149</i> .	804
753	learning in python. <i>the Journal of machine Learning</i>		
754	<i>research</i> , 12:2825–2830.	Kailai Yang, Tianlin Zhang, Ziyang Kuang, Qianqian Xie,	805
755		Jimin Huang, and Sophia Ananiadou. 2024. Mental-	806
756	Melanie F. Pradier, Michael C. Hughes, Thomas H. Mc-	lama: Interpretable mental health analysis on social	807
757	Coy, Sergio A. Barroilhet, Finale Doshi-Velez, and	media with large language models . In <i>Proceedings</i>	808
758	Roy H. Perlis. 2021. Predicting change in diagnosis	<i>of the ACM on Web Conference 2024</i> , volume 35 of	809
759	from major depression to bipolar disorder after an-	WWW '24, page 4489–4500. ACM.	810
760	tidepressant initiation . <i>Neuropsychopharmacology</i> ,		
761	46(2):455–461.	Amir Hossein Yazdavar, Hussein S Al-Olimat, Monireh	811
762		Ebrahimi, Goonmeet Bajaj, Tanvi Banerjee, Krish-	812
763	Lenore S Radloff. 1977. The ces-d scale: A self-report	naprasad Thirunarayan, Jyotishman Pathak, and Amit	813
764	depression scale for research in the general popula-	Sheth. 2017. Semi-Supervised Approach to Moni-	814
765	tion. <i>Applied Psychological Measurement</i> , 1(3):385–	toring Clinical Depressive Symptoms in Social Me-	815
766	401.	dia . <i>Proceedings of the ... IEEE/ACM International</i>	816
767	Yanping Ren, Hui Yang, Colette Browning, Shane	<i>Conference on Advances in Social Network Analysis</i>	817
768	Thomas, and Meiyang Liu. 2015. Performance of	<i>and Mining. IEEE/ACM International Conference on</i>	818
769	screening tools in detecting major depressive disorder	<i>Advances in Social Network Analysis and Mining</i> ,	819
770	among patients with coronary heart disease: a	2017:1191–1198.	820
771	systematic review. <i>Medical science monitor: interna-</i>		
772	<i>tional medical journal of experimental and clinical</i>	Hamad Zogan, Imran Razzak, Xianzhi Wang, Shoaib	821
773	<i>research</i> , 21:646.	Jameel, and Guandong Xu. 2022. Explainable de-	822
774		pression detection with multi-aspect features using a	823
775	Ramit Sawhney, Harshit Joshi, Saumya Gandhi, and	hybrid deep learning model on social media. <i>World</i>	824
776	Rajiv Ratn Shah. 2020. A Time-Aware Transformer	<i>Wide Web</i> , 25(1):281–304.	825
777	Based Model for Suicide Ideation Detection on So-		
778	cial Media . pages 7685–7697.	A Prompts Used for LLMs	826
779	Faisal Muhammad Shah, Farzad Ahmed, Sajib Ku-		
780	mar Saha Joy, Sifat Ahmed, Samir Sadek, Rimom	Table6 presents the prompts used for LLMs to ex-	827
781	Shil, and Md Hasanul Kabir. 2020. Early depres-	tract depression symptoms from posts. The se-	828
782	sion detection from social network using deep learn-	lected prompts are highlighted in bold, indicating	829
783	ing techniques. In <i>2020 IEEE region 10 symposium</i>	those chosen for the final analysis	830
784	<i>(TENSYP)</i> , pages 823–826. IEEE.		

Zero-shot	Identify depression symptoms in this post
	Can you extract depression symptoms from this text
	Extract depression symptoms from this text, return a list of symptoms
	Extract depression symptoms from this text, return a list of symptoms, return "no symptoms" if none existed
	Extract depression symptoms from the post below, if there is no depression symptom, return "this post does not contain depression symptoms"
Zero-shot with labels	Extract depression symptoms from the post below , the symptoms could be one or more from this list [poor appetite or eating disturbances, feeling down and depressed, crying, concentration problems, feeling tired or having little energy, feelings of failure, sleep disturbances, loss of interest, self-blame and shame, loneliness, suicidal thoughts] if there is no depression symptoms, return "this post does not contain depression symptoms."
	Extract depression symptoms from the post below , the symptoms could be one or more from this list only [poor appetite or eating disturbances, feeling down and depressed, crying, concentration problems, feeling tired or having little energy, feelings of failure, sleep disturbances, loss of interest, self-blame and shame, loneliness, suicidal thoughts] if there is no depression symptoms, return only "this post does not contain depression symptoms."
Few shots	Extract depression symptoms from the post below, the symptoms could be one or more from this list only [poor appetite or eating disturbances, feeling down and depressed, crying, concentration problems, feeling tired or having little energy, feelings of failure, sleep disturbances, loss of interest, self-blame and shame, loneliness, suicidal thoughts] if there is no depression symptoms, return only "this post does not contain depression symptoms." Examples: - post: "I Received an unexpected surprise it's been an emotional afternoon. I'm feeling sentimental about someone. This season has been challenging, but this moment has lifted my spirits. I'm feeling nostalgic" symptoms: feeling down and depressed - post: "others have commented on my appearance, saying I seem more toned, and that makes me thrilled. Personally, I don't notice the difference, but it's obvious that my new eating habits are paying off." Symptoms: this post does not contain depression symptoms - post: "i've noticed a discrepancy between my online presence and the response I get when I share my thoughts on a critical issue. It seems that despite having a large following, my words often fall on deaf ears" symptoms: feeling failure, loneliness POST: {post} symptoms:

Table 6: Prompts used for LLMs in extracting depression symptoms from a post. The chosen prompts are bold-faced

B Models Hyper-parameters for Extracting Depressive Symptoms

We performed hyperparameter tuning for BERT-based models to optimize performance. The best results, obtained after extensive experimentation, are presented below.

BERT

Model_card: "bert-base-uncased"
Epochs: 64
Batch_size: 8
Learning_rate: 5e-5
Hidden_size: 128
Optimizer: Adam
Loss: BCEWithLogitsLoss

RoBERTa

Model_card: "roberta-base"
Epochs: 128
Batch_size: 32
Learning_rate: 1e-05
Hidden_size: 128
Optimizer: Adam
Loss: BCEWithLogitsLoss

MentalBERT

Model_card: "mental-bert-base-uncased"

Epochs: 64

Batch_size: 32

Learning_rate: 1e-05

Hidden_size: 128

Optimizer: Adam

Loss: BCEWithLogitsLoss

GPT. We used the GPT-3.5 "gpt-3.5-turbo" and GPT-4 "gpt-4-turbo" versions, as these have shown a strong ability to understand human-like emotional context (Tavast et al., 2022), and in sentiment analysis (Kheiri and Karimi, 2023). We used the Official OpenAI Python library⁷ to collect responses with the settings (temperature = 0.2, max tokens = 30).

Llama 2. The Hugging Face library is used for MentalBERT tokenization and fine-tuning, namely the 'meta-llama/Llama-2-7b-hf' model card.

Llama 3. The Hugging Face library is used for MentalBERT tokenization and fine-tuning, namely the meta-llama/Meta-Llama-3-8B' model card.

MentalLama. MentalLlama is a specialized large language model built on Llama designed

⁷<https://platform.openai.com/docs/libraries/python-library>

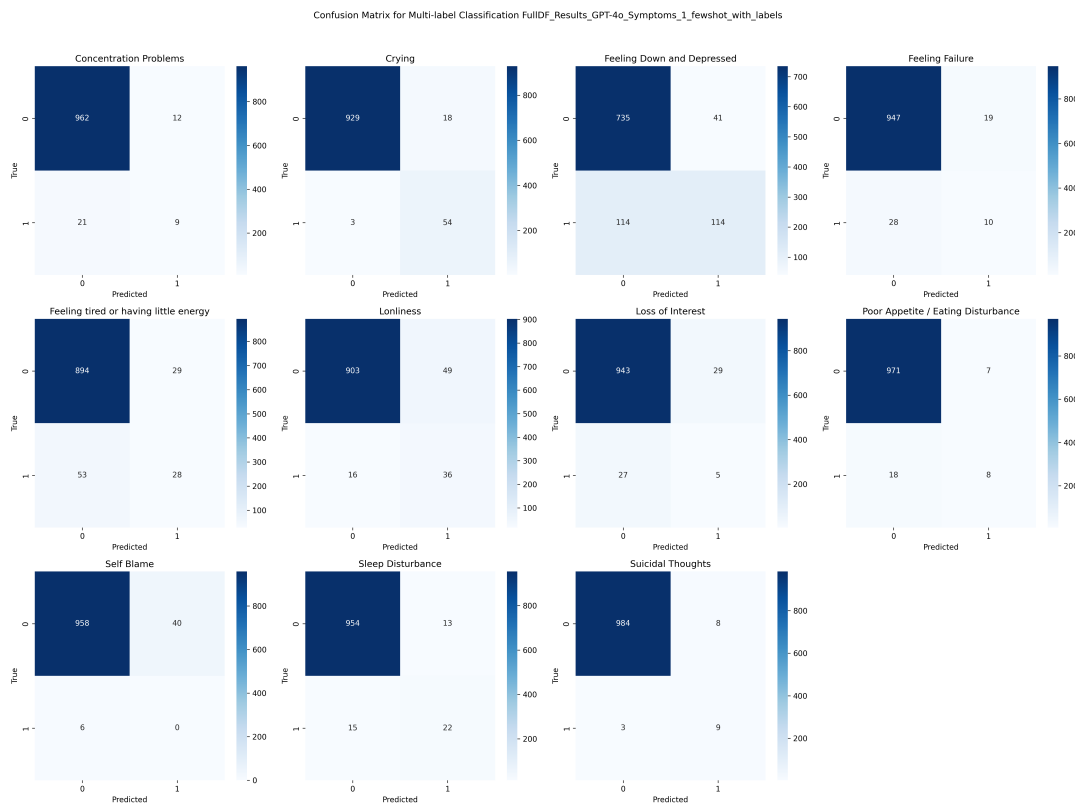


Figure 4: Confusion matrices for GPT-4o in the few-shot setting for symptom-level classification.

to excel in the domain of mental health (Yang et al., 2024). MentaLLama incorporates domain-specific training data and techniques to enhance its ability to understand, process, interpret, and generate text related to mental health. ‘klyang/MentaLLaMA-chat-7B’

C Appendix: Per-Symptom Confusion Matrices for LLMs

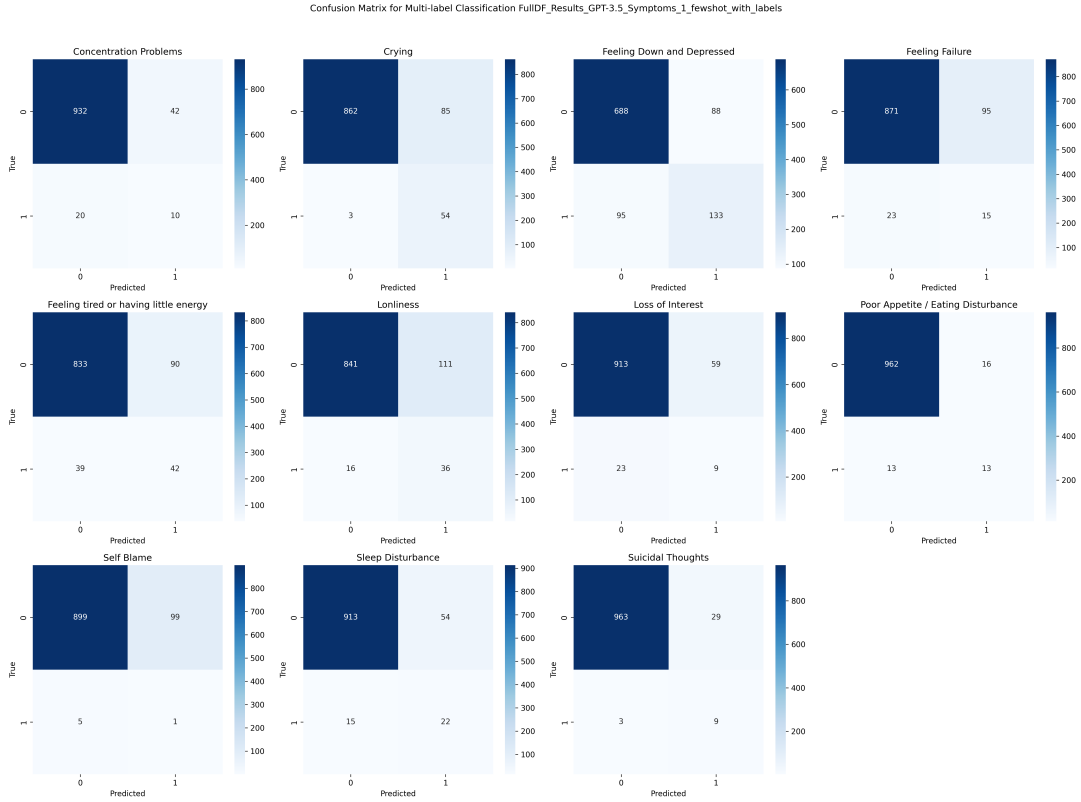


Figure 5: Confusion matrices for GPT-3.5 in the few-shot setting for symptom-level classification.

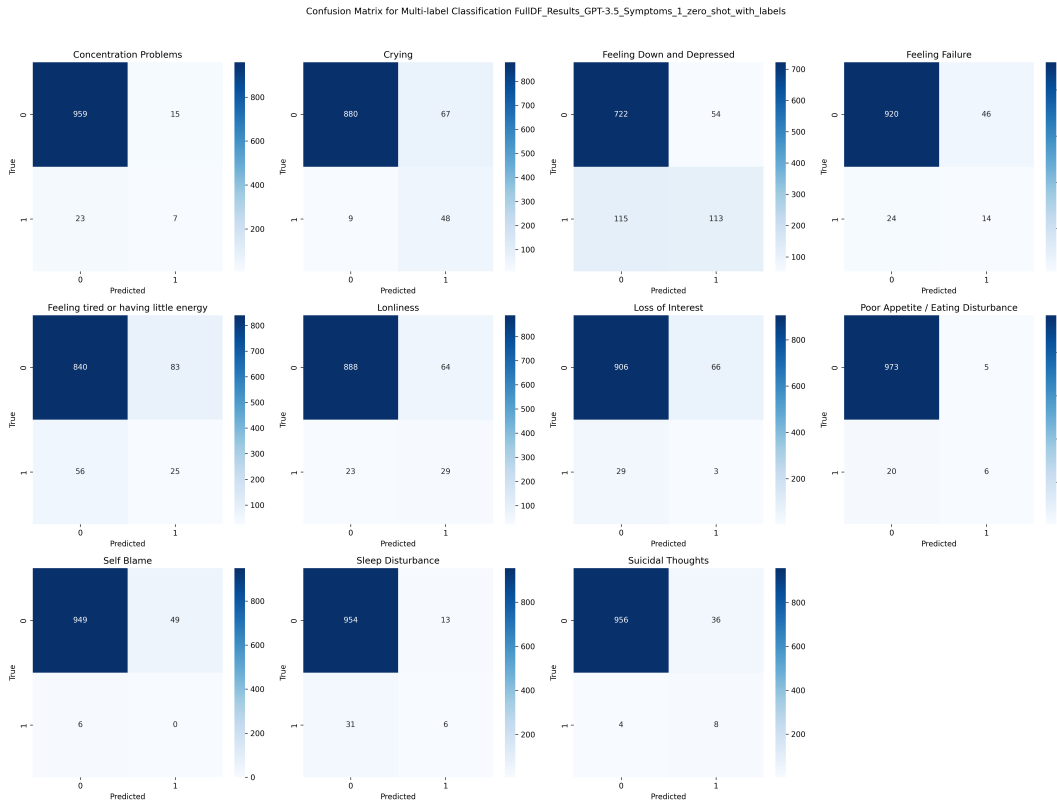


Figure 6: Confusion matrices for GPT-3.5 in the zero-shot with labels setting.

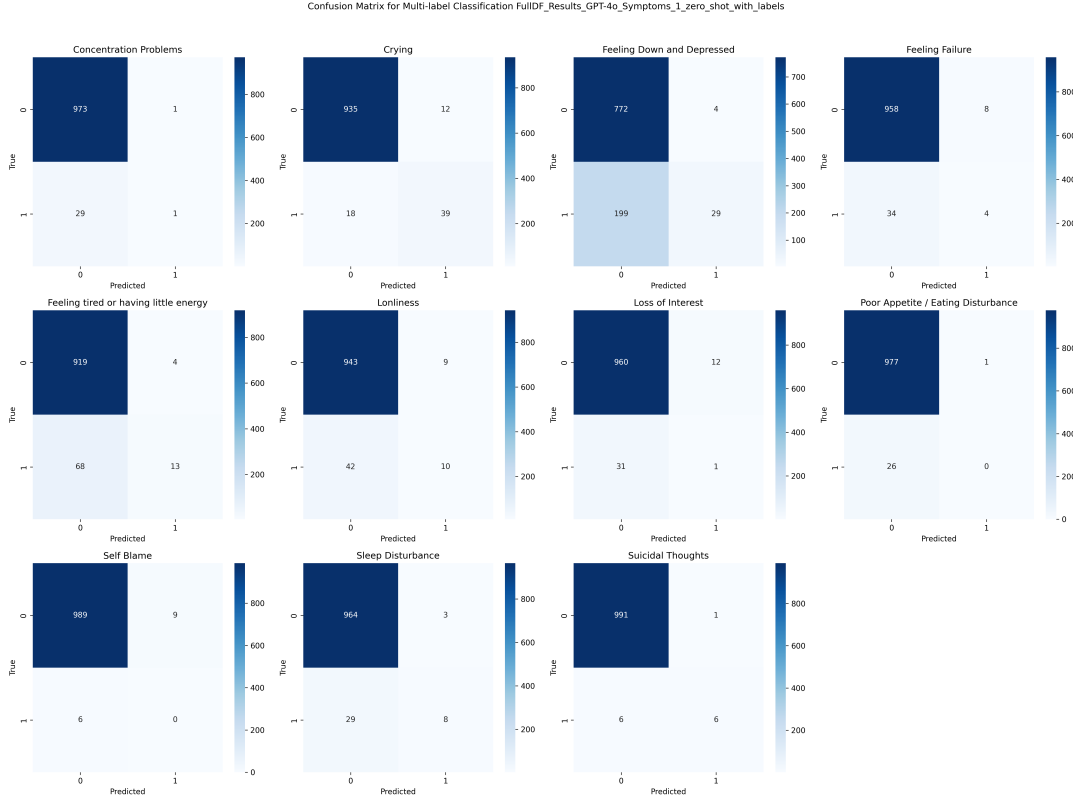


Figure 7: Confusion matrices for GPT-4o in the zero-shot with labels setting.

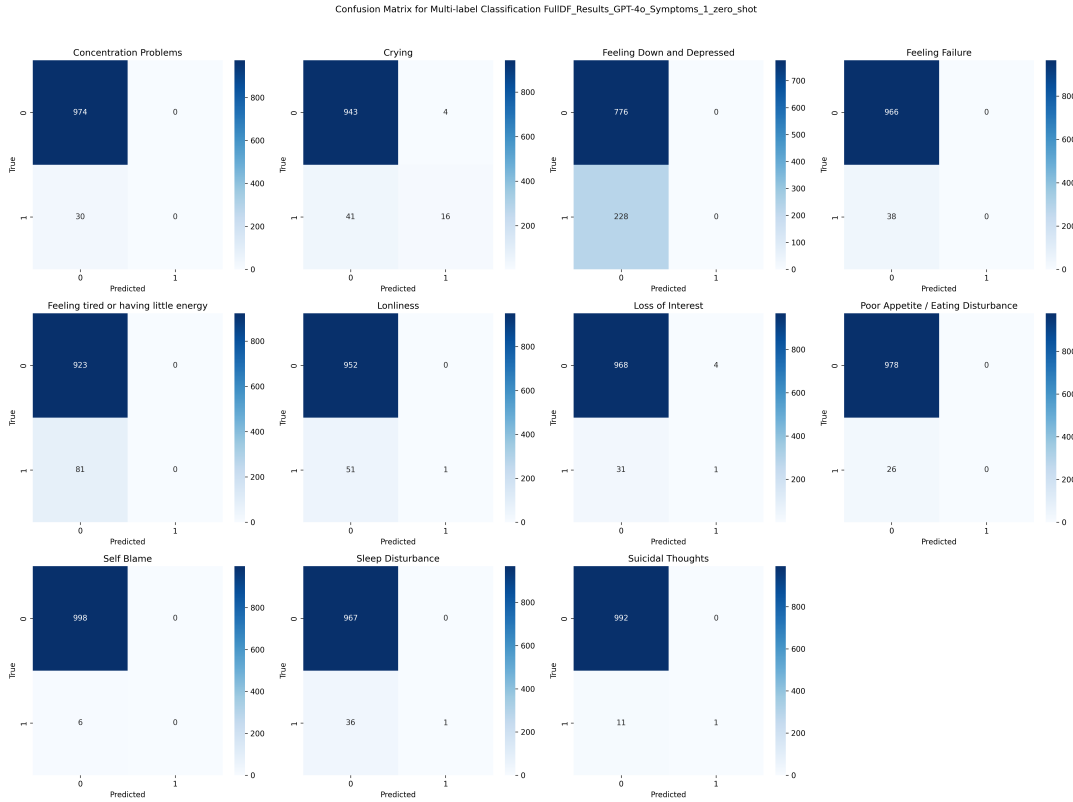


Figure 8: Confusion matrices for GPT-4o in the pure zero-shot setting (no labels or examples provided).

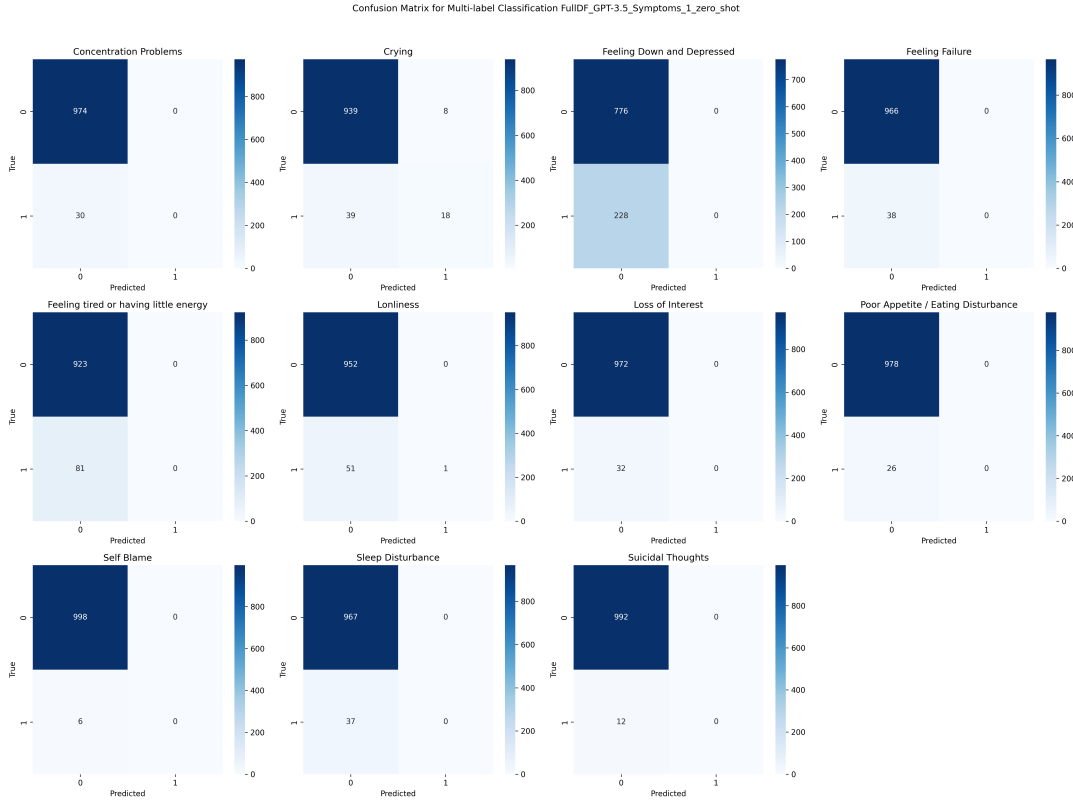


Figure 9: Confusion matrices for GPT-3.5 in the pure zero-shot setting (no labels or examples provided).

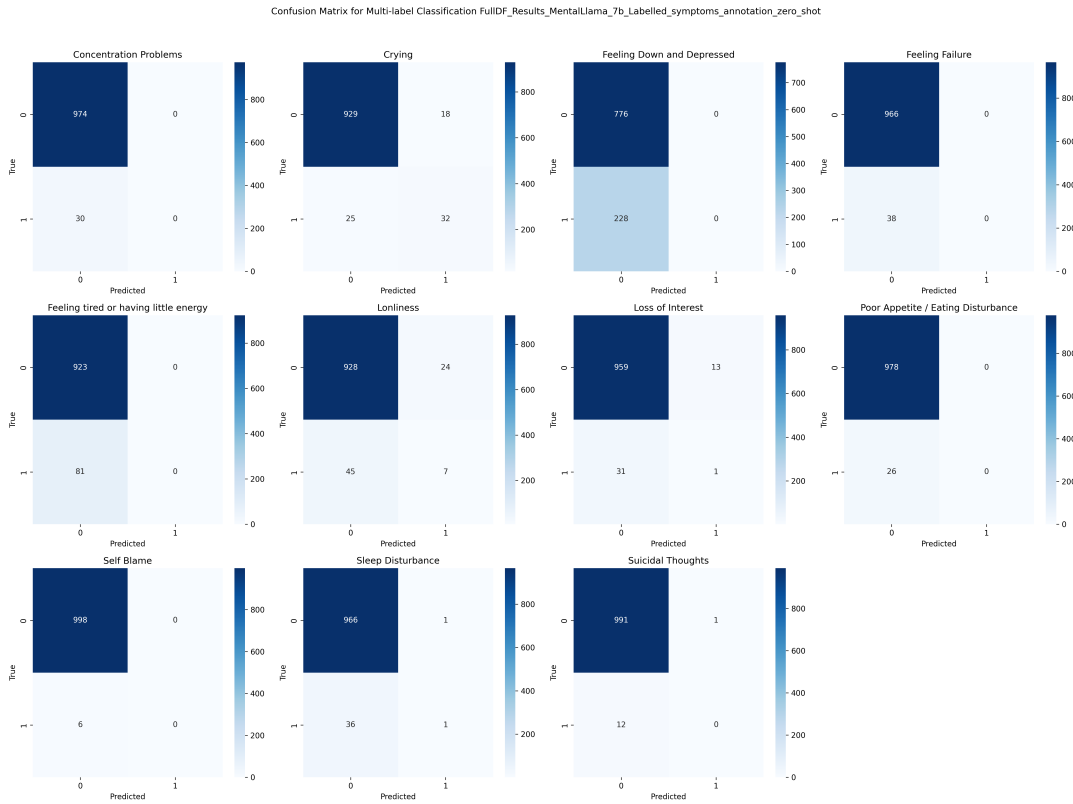


Figure 10: Confusion matrices for MentalLLaMA-7B in the zero-shot setting.

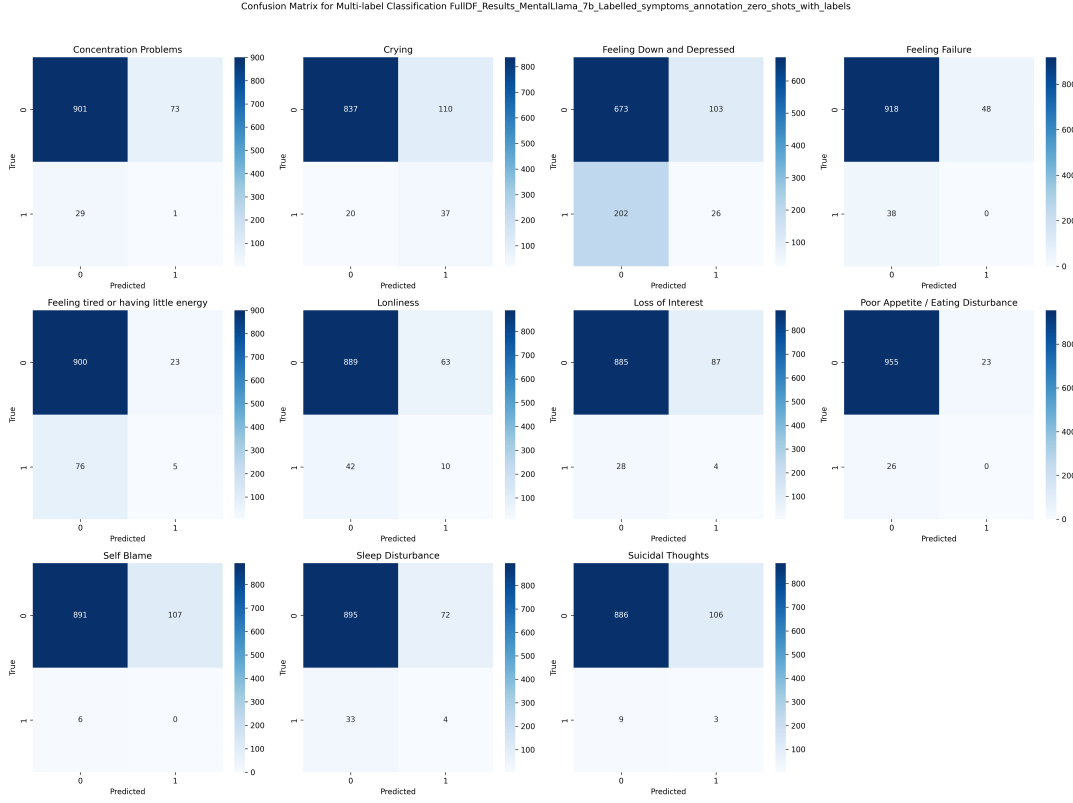


Figure 11: Confusion matrices for MentalLLaMA-7B in the zero-shot with labels setting.

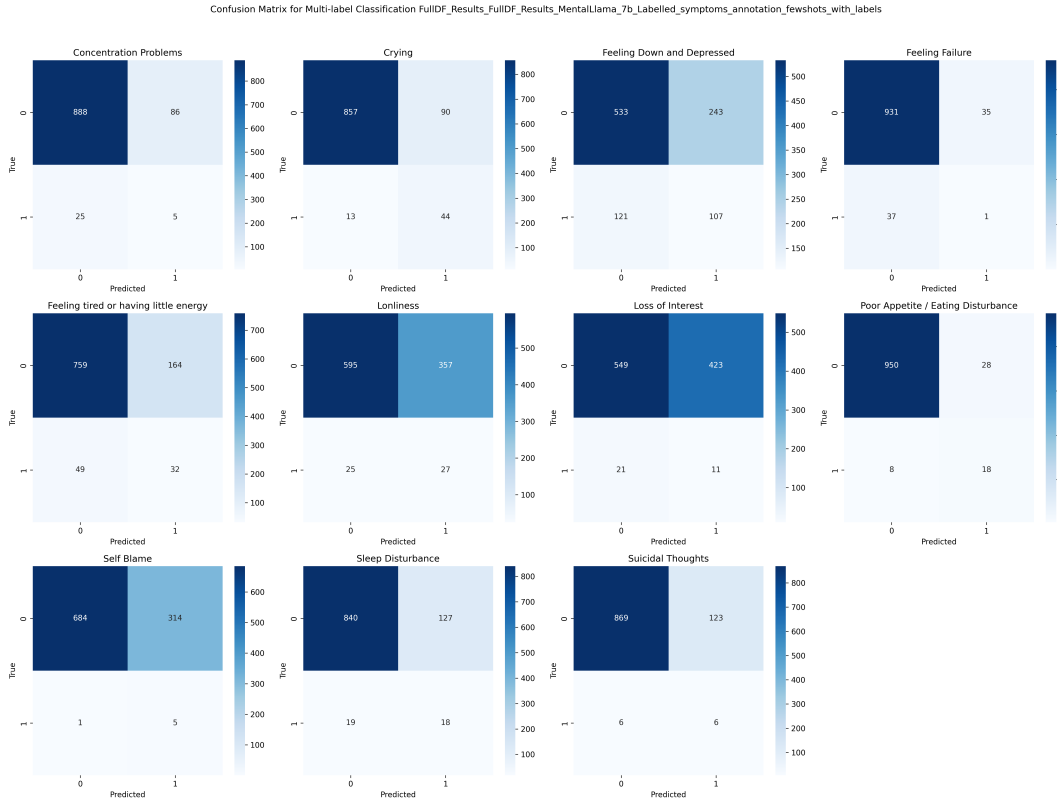


Figure 12: Confusion matrices for MentalLLaMA-7B in the few-shot setting.

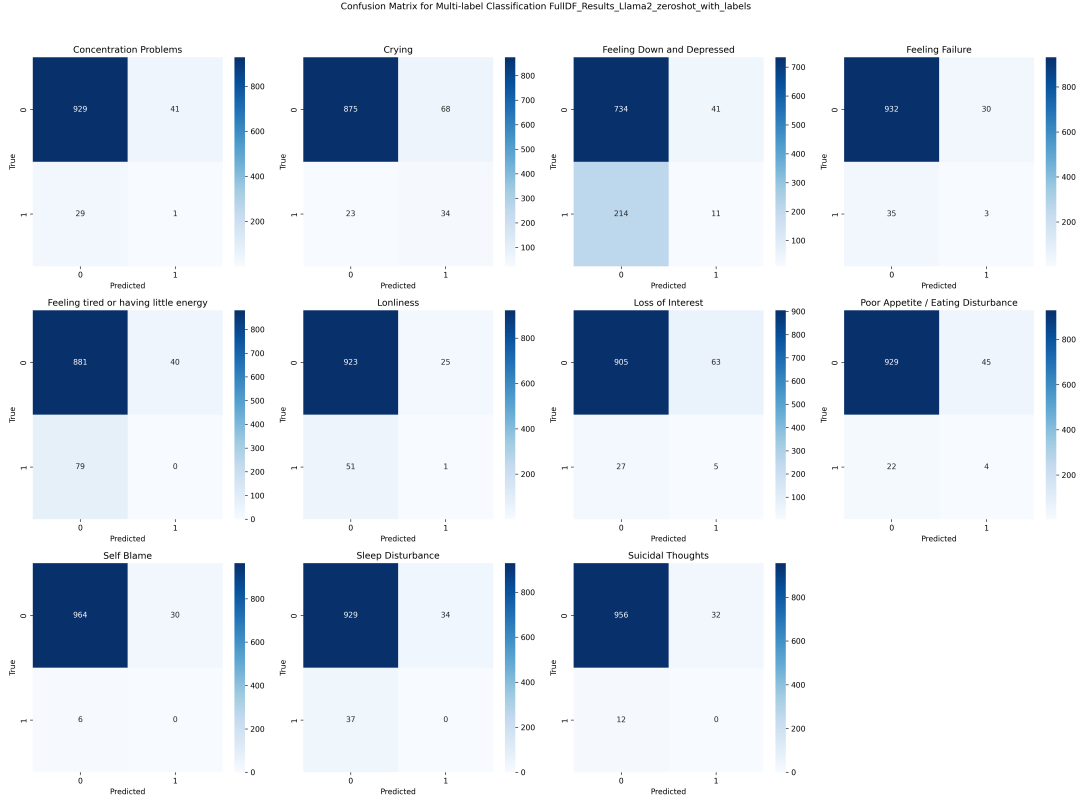


Figure 13: Confusion matrices for Llama2-7B in the zero-shot with labels setting.

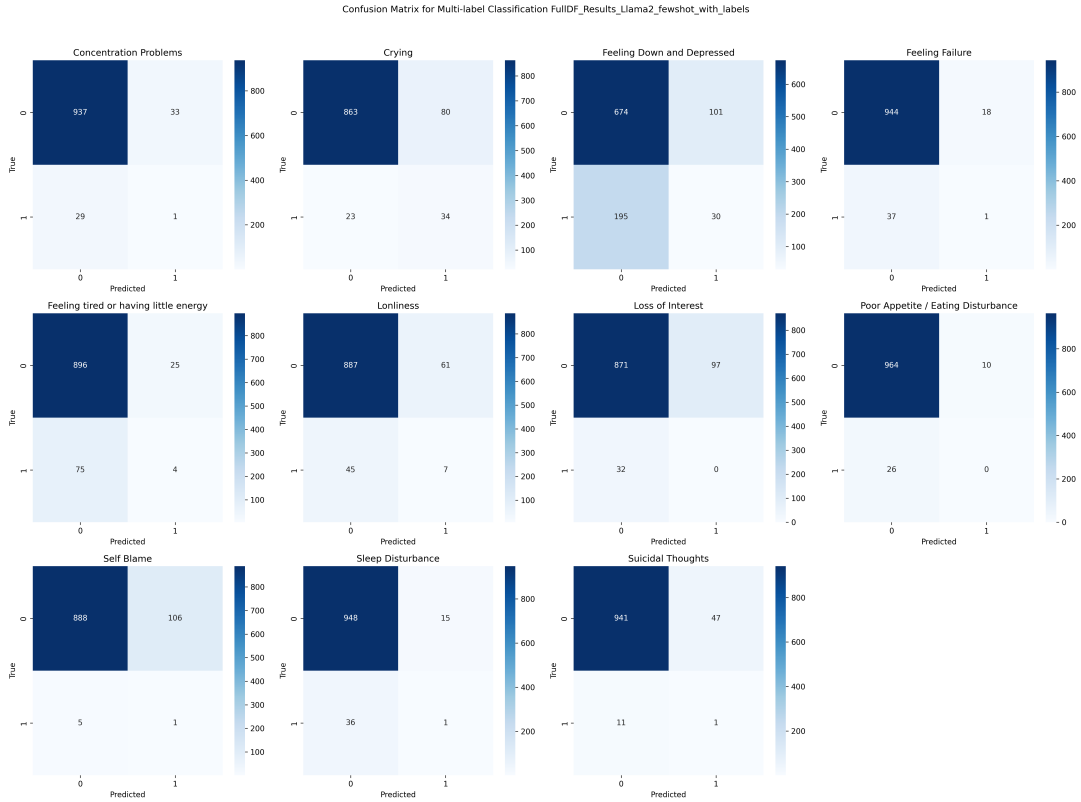


Figure 14: Confusion matrices for LLaMA-7B in the few-shot setting.

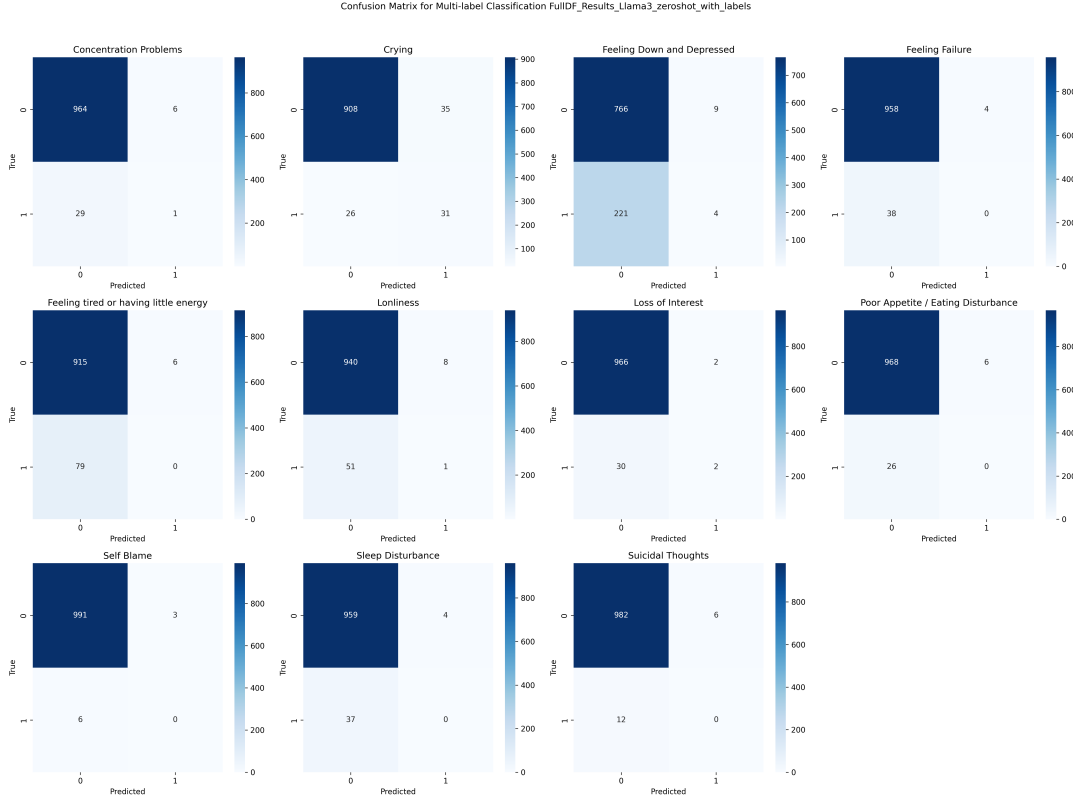


Figure 15: Confusion matrices for LLaMA3-8B in the zero-shot with labels setting.

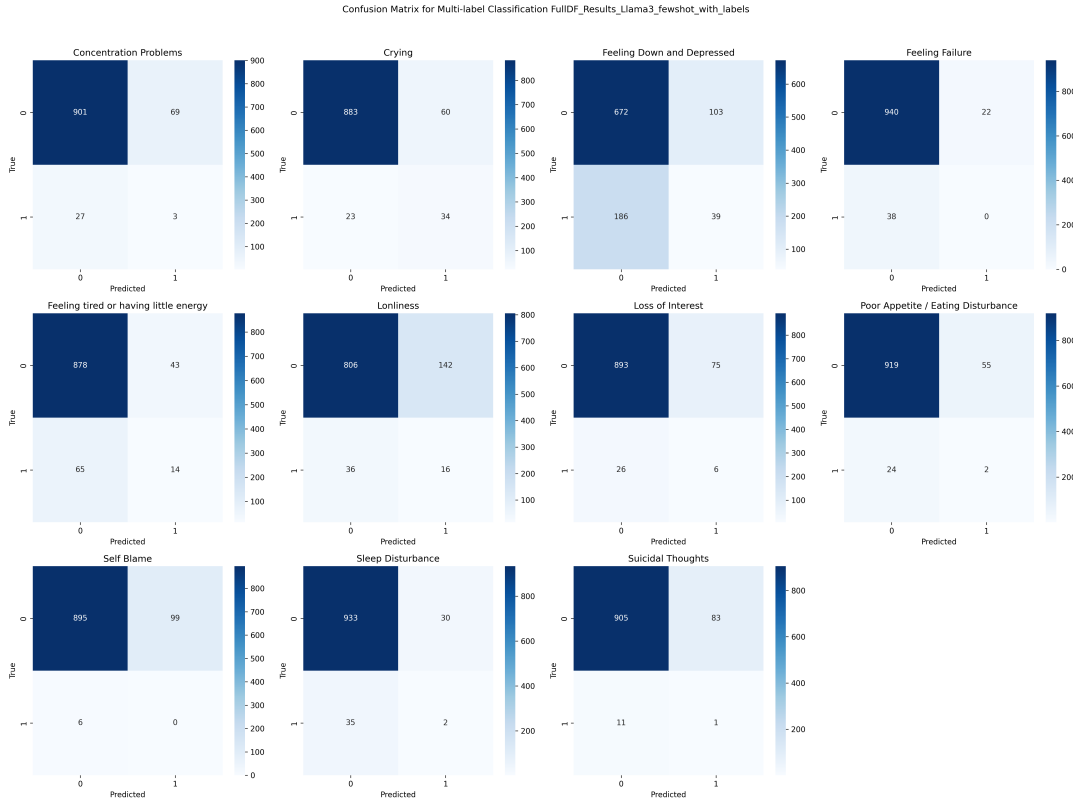


Figure 16: Confusion matrices for LLaMA3-8B in the few-shot setting.