

Article

Not peer-reviewed version

Leveraging Medical Knowledge Graphs and Large Language Models for Enhanced Mental Disorder Information Extraction

[Chaelim Park](#) , [Hayoung Lee](#) , [Okran Jeong](#) *

Posted Date: 4 June 2024

doi: 10.20944/preprints202406.0162.v1

Keywords: Depression; Knowledge Graph; Zero-shot Information Extraction; Large Language Models; DSM-5; Mental Health; Entity Linking



Preprints.org is a free multidiscipline platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Article

Leveraging Medical Knowledge Graphs and Large Language Models for Enhanced Mental Disorder Information Extraction

Chaelim Park, Hayoung Lee and Ok-ran Jeong *

School of Computing, Gachon University, 1342 Sujeong-gu, Seongnam-si 13120, Gyeonggi-do, Republic of Korea

* Correspondence: orjeong@gachon.ac.kr

Abstract: The accurate diagnosis and effective treatment of mental health disorders such as depression remain challenging owing to the complex underlying causes and varied symptomatology. Traditional information extraction methods struggle to adapt to evolving diagnostic criteria such as the Diagnostic and Statistical Manual of Mental Disorders fifth edition (DSM-5) and to contextualize rich patient data effectively. This study proposes a novel approach for enhancing information extraction from mental health data by integrating medical knowledge graphs and large language models (LLMs). Our method leverages the structured organization of knowledge graphs specifically designed for the rich domain of mental health combined with the powerful predictive capabilities and zero-shot learning abilities of LLMs. This research enhances the quality of knowledge graphs through entity linking and demonstrates superiority over traditional information extraction techniques, making a significant contribution to the field of mental health. It enables a more fine-grained analysis of the data and the development of new applications. Our approach redefines the manner in which mental health data are extracted and utilized. By integrating these insights with existing healthcare applications, the groundwork is laid for the development of real-time patient monitoring systems. The performance evaluation of this knowledge graph highlights its effectiveness and reliability, indicating significant advancements in automating medical data processing and depression management.

Keywords: depression; knowledge graph; zero-shot information extraction; large language models; DSM-5; mental health; entity linking

1. Introduction

Depression is a major public health concern worldwide. The "World Mental Health Report" (2020) ¹ of the World Health Organization (WHO) revealed a significant rise in mental health conditions: depression cases reached 264 million (a 28% increase from 2019) and anxiety disorders rose to 374 million (26% increase). This surge highlights the growing societal and economic impacts of mental health, which was previously considered a private matter. In collaboration with the International Labor Organization, the WHO estimates that mental health problems incur annual economic losses of approximately \$1 trillion ². These figures highlight the crucial role of mental health beyond social welfare.

Given the complex and multifaceted nature of depression, successful treatment and management require personalized approaches. Systematic data analysis and comprehension are vital for medical professionals and researchers. However, current biomedical databases often rely on manual information extraction from medical literature, leading to time-consuming and inefficient processes [1]. To overcome the limitations of previous methods, large language models (LLMs), which have shown excellent performance, have been actively utilized in many areas, such as information extraction [2] and prompt tuning [3]. In particular, recent studies have focused on zero-shot information extraction

¹ World mental health report: Transforming mental health for all; <https://www.who.int/publications/i/item/9789240049338>

² WHO and ILO call for new measures to tackle mental health issues at work; <https://www.who.int/news/item/28-09-2022-who-and-ilo-call-for-new-measures-to-tackle-mental-health-issues-at-work>

(ZeroIE) [4] and relation triple extraction (ZeroRTE) [5,6], which enable pretrained models to extract information from new data.

ZeroIE and ZeroRTE play pivotal roles in depression research owing to several key advantages. First, they reduce the data-labeling costs. Depression-related data are vast and manual labeling is expensive and inefficient. ZeroIE leverages pretrained models to extract valuable information from unlabeled data, thereby saving time and resources [7]. Second, ZeroIE and ZeroRTE are versatile and scalable. Depression is a complex illness characterized by a wide range of symptoms and evolving diagnostic criteria. ZeroIE models can extract information from diverse domains and contexts, enabling their application to new symptoms or diagnostic criteria without requiring model retraining. Third, ZeroIE and ZeroRTE facilitate rapid information extraction. Timely access to the latest information is critical for depression research and treatment. ZeroIE models can process and extract information in real time, thereby empowering healthcare professionals and researchers to respond promptly. Finally, ZeroIE and ZeroRTE support personalized treatment. Personalized approaches are essential for effective depression management. ZeroIE enables the rapid extraction and analysis of patient-specific information, which can facilitate the development of personalized treatment plans that optimize outcomes [8].

This study proposes a novel method for enhancing information extraction from mental disorder data by integrating medical knowledge graphs and LLMs. Our approach leverages the structured organization of knowledge graphs tailored to the rich mental health domain. We combine these knowledge graphs with the robust predictive and zero-shot learning capabilities of LLMs to extract and organize depression-related information efficiently. Specifically, we develop a knowledge graph that comprehensively maps the symptoms and diagnostic criteria associated with depression based on the Diagnostic and Statistical Manual of Mental Disorders fifth edition (DSM-5) document, which is a medical resource that describes the criteria for designating anxiety disorders for major depressive disorder and bipolar depression [9].

Furthermore, rigorous performance evaluation experiments were conducted to validate the effectiveness and reliability of the knowledge graph. This research significantly contributes to advancing the understanding of depression management and treatment, while also redefining the analysis and use of data in the mental health field. Our approach outperforms traditional information extraction methods through advanced entity linking techniques and contributes to deeper understanding and management of complex mental health issues such as depression.

The key contributions of this research are as follows:

- Automated and enhanced information extraction accuracy: By leveraging LLMs for zero-shot learning, we automate the traditional manual information extraction process, thereby enabling faster and more accurate data handling.
- Improved data accuracy through entity linking: By employing entity linking techniques, we ensure accurate connections between textual information and entities within the knowledge graph, leading to increased data consistency and reliability.
- Expanded practical applications in medicine: The developed knowledge graph has broader applicability beyond depression, potentially encompassing other mental disorders.

This research paves the way for integration with various medical applications including real-time patient monitoring systems. The remainder of this paper is organized as follows: Section 2 reviews related studies and discusses recent advancements and trends in medical knowledge graphs and zero-shot information extraction techniques. Section 3 details the proposed methodology, with a focus on constructing medical knowledge graphs and using LLMs for entity linking. Section 4 presents the experimental results and demonstrates the effectiveness of the proposed method. Section 5 evaluates the validity of the methodology based on the findings and discusses its potential applications in the medical field. Finally, Section 6 concludes the paper and proposes directions for future research.

2. Related Works

2.1. Medical Knowledge Graphs

Medical knowledge graphs have garnered significant attention owing to their excellent performance in intelligent healthcare applications. As diverse medical departments proliferate within hospitals, numerous medical knowledge graphs are being constructed for various diseases. For instance, [10] built a graph encompassing over one million clinical concepts from 20 million clinical notes, covering drugs, diseases, and procedures, while [11] constructed a graph detailing 156 diseases and 491 symptoms based on emergency room visits. In addition, [12] developed an EMR-based medical knowledge network (EMKN) with 67,333 nodes and 154,462 edges, focusing on symptom-based diagnostic models to explore the application and performance of knowledge graphs.

Various datasets have led to diverse methods of constructing medical knowledge graphs. The authors of [13] utilized logistic regression, naïve Bayes classifiers, and Bayesian networks with noisy OR gates to automatically generate knowledge graphs automatically using maximum likelihood estimation. In addition, a combination of bootstrapping and support vector machines was applied [14] to extract relationships among entities to build an obstetrics- and gynecology-related knowledge graph. Recent studies have increasingly focused on automated entity and relation extraction using deep-learning techniques [15,16].

Existing medical knowledge graphs have provided insights into individual disease domains and expanded their application areas for diagnosis and recommendation. However, constructing medical knowledge graphs requires expert review and labeling, which, while detailed, demand considerable time and resources. This requirement makes it challenging to extend knowledge graphs to new or different diseases, particularly in the mental health sector, where research progress has been slow owing to the complexity and sensitivity of the data. Our study aims to overcome these limitations by automating the extraction of entities and relationships using high-performance LLMs.

2.2. Zero-Shot Information Extraction

As AI models are trained and output based on datasets, their performance depends on the structure of the model as well as the quality of the data. Expert-reviewed labeling of large amounts of data is an essential process; however, it is labor intensive and time consuming. Therefore, to reduce the time and labor spent on datasets, substantial work has been conducted on extracting relations [17] and arguments [18] from a few resources based on zero-/few-shot techniques [19].

Particularly in medicine, where clinical notes, medical reports, and patient information often lack the necessary annotations, the cost of using AI models is prohibitive, thereby delaying their application and limiting their adaptability to other medical fields. Recent advancements include the application of zero-/few-shot information extraction techniques in the medical field, simplifying medical reports using ChatGPT [20] and extracting information from radiology reports [21].

Our study builds on these methodologies, targeting depression-related information extraction with a focus on the accurate linking of entities and relationships using detailed annotation guidelines and zero-shot learning.

2.3. Entity Linking

Entity linking plays a crucial role in preprocessing and preparing model inputs by identifying key entities in the text data and linking them to appropriate identifiers in knowledge bases [22]. This process is vital for converting text mentions into structured data, reducing redundancy, and enhancing data consistency. Recent trends have highlighted the use of LLMs such as GPT-3 for performing entity linking [23,24], utilizing their extensive general knowledge and language comprehension capabilities to identify and link entities accurately, even in complex contexts.

Our research employs the gpt-3.5-turbo-instruct model [25] for effective entity identification and linking during the preprocessing stage, which significantly enhances the accuracy and relevance of the

extracted information and contributes to the efficiency and precision of the information extraction and knowledge graph construction processes.

This approach offers a high level of automation and scalability, especially when working with limited data, and provides a foundational basis for exploring the applicability of LLMs to extend the scope of entity linking across various data types.

3. Proposed Method

This section describes each component of the proposed framework and the data processing in detail. This process includes the schema definition and annotation guideline setting, information extraction and entity linking, guideline-based model usage, triple extraction, final output, and knowledge graph structure.

We propose a novel zero-shot approach that integrates medical knowledge graphs and LLMs to enhance information extraction from mental health disorder data. The data processing pipeline of the proposed framework is illustrated in Figure 1. To improve the extraction efficiency, we use LLMs to extract the desired information from a document corpus and identify and link relevant entities based on schema definitions and annotation guidelines. For zero-shot extraction, we provide high-quality gold annotations to a guideline-based model using the results of information extraction and entity linking, and convert them into a final knowledge graph through triple extraction. Therefore, the proposed framework focuses on effectively automating and optimizing complex data processing tasks.

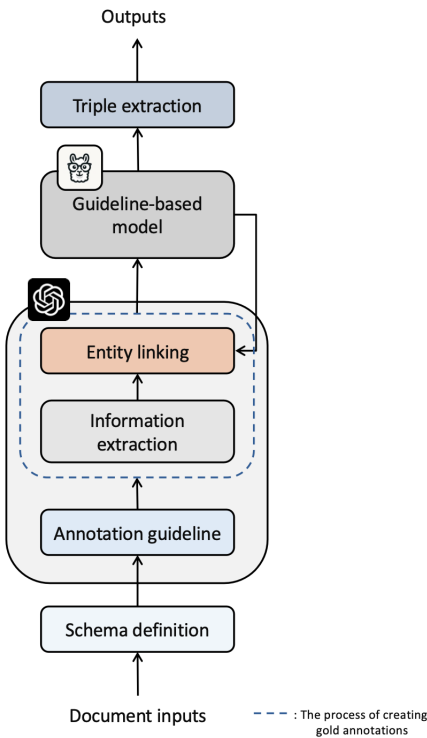


Figure 1. Flowchart of the guideline-based model process for information extraction and entity linking.

3.1. Schema Definition and Annotation Guidelines

We propose a zero-shot-based information extraction method. We provide the model with the required schema and annotation guidelines required to extract the correct information at the document level without any pretrained relevant information [26]. Thus, we aim to achieve successful zero-shot extraction by providing guidelines based on the performance results of the LLM model.

Schema Definition: The schema definition is an essential initial step in setting the standards for the information extraction process. It clearly defines the types of entities and relationships to be extracted from the documents, thereby enabling structured and consistent data processing and establishing the semantic structure of the data through entities and relationships [27].

Entity Definition: Each entity represents a specific concept or object, and entity types and attributes are structured using Python's @dataclass. These entities function as key components in the data extraction and analysis.

```

1 @dataclass
2 class Entity:
3     """
4     Represents an entity in a document, which could be a person, location, organization
5     , or any other named object.
6     The 'name' attribute holds the canonical name of the entity as identified in the
7     document text.
8
9     Attributes:
10         name (str): The canonical name of the entity.
11     """
12     name: str

```

Relationship Definition: Relationships specify the interactions between entities and provide connections among them. The relationship types and examples are defined precisely to ensure consistency in the annotation guidelines.

```

1 @dataclass
2 class Triple:
3     """
4     Represents a triple consisting of two entities and a relation in a semantic graph
5     derived from textual data.
6     A triple is typically used in knowledge representation and information extraction
7     tasks to model factual statements.
8
9     Attributes:
10         Subject (Entity): The subject entity of the triple,
11         usually the one performing or causing an action.
12         Predicate (str): The relationship that connects
13         the subject and object entities.
14         Object (Entity): The object entity of the triple,
15         usually the one experiencing or affected by the action.
16     """
17     Subject: Entity
18     Predicate: str
19     Object: Entity

```

Annotation Guidelines and Model Utilization: The annotation guidelines are based on the schema definition to ensure that the data are processed consistently and accurately within the documents. To enable the LLM to perform a specific task, we provide natural language descriptions of that task, known as prompts, along with the input data. This allows the model to recognize the in-context information within the guidelines and generate the appropriate output based on the format [28,29]. This study uses the gpt-3.5-turbo-instruct model [25] to generate guidelines, offering significant benefits for information extraction, particularly with mental health data.

The model can generate consistent guidelines across different datasets, ensuring consistency in the data annotation process, thereby enhancing the accuracy of data analysis. In addition, the ability of the model to generate effective guidelines for new and unseen data enhances its adaptability to zero-shot or few-shot learning scenarios [30]. Finally, the GPT model leverages advanced natural language processing (NLP) capabilities to provide high-quality data annotation guidelines, contributing to the reliability of the data analysis.

3.2. Information Extraction and Entity Linking

Information Extraction: This stage also utilizes the gpt-3.5-turbo-instruct model [25], leveraging its NLP capabilities to identify key information and entities within the documents. The model excels at recognizing meaningful patterns and structures in text, particularly within specialized domains such as mental health text, ensuring high accuracy in information extraction [31]. This is achieved as follows:

- **Context-aware recognition:** The model analyzes the surrounding text to determine the precise scope and types of entities. For instance, it can distinguish whether "depression" refers to a clinical condition or a general emotional state [32].
- **Entity classification:** The extracted information is categorized based on predefined classes (e.g., symptoms, diagnoses, and treatments) aligned with the established schema.

Entity Linking: Entity linking is crucial to guarantee the quality and accuracy of depression-related data annotations. This study employs the NLP capabilities of the gpt-3.5-turbo-instruct model [25] to perform entity-linking tasks effectively. The model identifies relevant entities within the data and accurately maps their relationships. These high-quality annotations, termed "gold annotations," are then fed into the guideline-based model described in the following section [33]. This step is instrumental in enabling more refined data analysis and information extraction, which are essential for the efficient processing of complex mental health data.

3.3. Guideline-Based Model

We utilize a pretrained Gollie model [34] to analyze the depression data. This model is strong in zero-shot information extraction using annotation guidance and can effectively handle NLP and code-based input and output, making it well-suited for processing medical data such as depression data [35]. The model was trained on large text datasets collected from various domains including news, biomedicine, and social media.

Notably, the training process incorporated normalization techniques, such as dropout and batch normalization, enabling the model to learn data variations and prevent overfitting. This approach allowed the Gollie model [34] to reflect real-world language-usage patterns effectively and adapt to a broader range of data. Utilizing a pretrained model facilitates swift and accurate depression data analysis, bypassing the complexities of the initial model setup and lengthy training times. The model leverages the extensive language knowledge acquired during training to identify and analyze meaningful patterns and relationships within the depression data efficiently.

3.4. Triple Extraction and Output

Triple Extraction: The triple extraction process utilizes the gold annotations generated by the LLM model above to convert the interrelationships between the entities extracted from the document into the triple form [36]. These triples consist of three elements: subject, predicate, and object. They semantically link each piece of information to create a machine-readable representation of the relationships. For instance, the relationship between "major depressive disorder" (subject) and "difficulty thinking" (object) can be expressed as a triple with the predicate "is characterized by."

The results of the triple extraction are stored as structured data, which can be used as the basis for knowledge graphs. The resulting knowledge graph can help researchers and healthcare professionals to access the required information rapidly by visually representing complex data relationships and making them easy to navigate. This process maximizes the value of the data and plays an important role in adding depth to medical decision-making, work, and education.

Output: The extracted triples are integrated into the knowledge graph, making the information readily usable in various applications. This knowledge graph plays a crucial role in visualizing complex relationships and enabling easy access to information [37]. Figure 2 shows a sample of the knowledge graph developed in this study. It illustrates the interrelationships and characteristics of

major depressive disorder and disruptive mood dysregulation disorder, as defined in the DSM-5. This visual representation demonstrates how the proposed method for constructing a knowledge graph can be applied to capture the complexities of mental health issues.

This section details the triple extraction process and the methods for constructing and utilizing the knowledge graph. It highlights the methodological contributions of this study by demonstrating how data structuring can be applied across various fields, particularly for managing complex medical information.

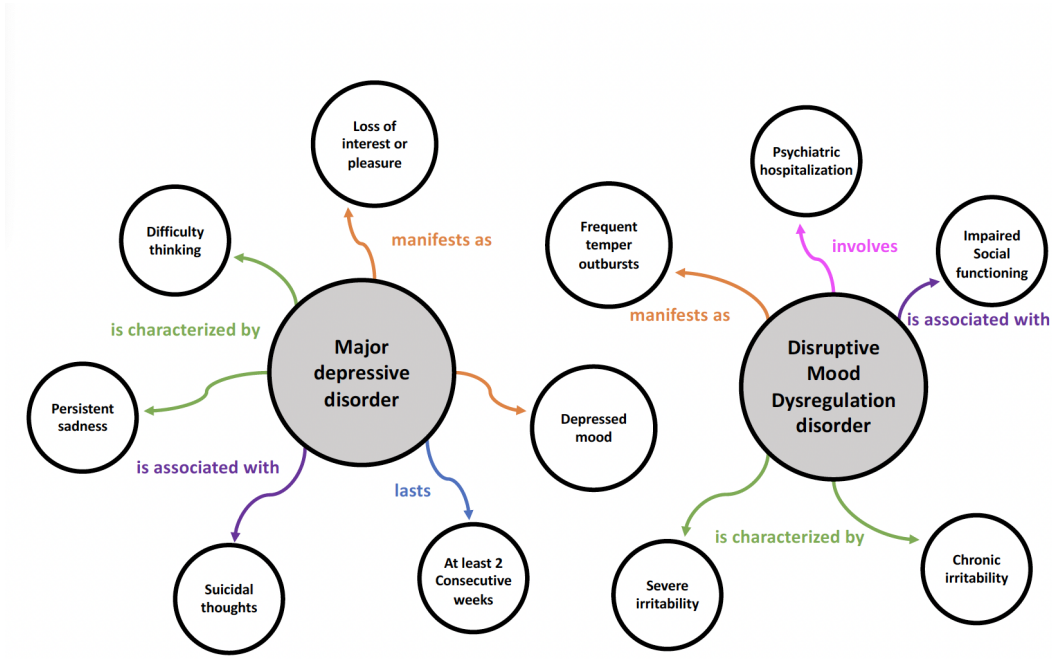


Figure 2. Illustrative representation of the constructed knowledge graph comparing major depressive disorder and disruptive mood dysregulation disorder.

3.5. Structure of the Knowledge Graph

As a result of the zero-shot information extraction process proposed in this study, a medical knowledge graph was constructed by systematically categorizing complex depression symptoms and relationships. This graph encompasses 381 nodes and 505 relationships and provides a detailed representation of the clinical characteristics of depression. The knowledge graph comprises two primary node types: subject and object.

- **Subject nodes:** These nodes represent specific medical conditions or pathological states. They categorize various forms of depression according to the DSM-5 criteria, including major depressive disorder, persistent depressive disorder, premenstrual dysphoric disorder, and substance/medication-induced depressive disorder. Each subject node is labeled with the name of the disorder and a unique identifier.
- **Object nodes:** These nodes depict the symptoms or characteristics that a subject node may exhibit. Examples include emotional or behavioral responses such as depressed mood, fatigue, and loss of interest or pleasure. Each object node includes the name of the symptom and a unique identifier.
- **Interaction between nodes:** Within the knowledge graph, each subject node is linked to one or more object nodes that represent the manifestation of disease traits or symptoms. For instance, the "manifests as" relationship indicates how major depressive disorder might manifest as irritability, while the "is characterized by" relationship suggests that it may be characterized by a depressed mood.

This structured organization and defined relationships allow for an in-depth analysis and understanding of the various symptoms of depression and their interactions. Researchers, physicians, and treatment specialists can use this information to gain a deeper understanding of the causes, manifestations, and characteristics of the diseases. This information will enable the formulation of more effective diagnostic and treatment plans. Table 1 displays a subset of the data used in the knowledge graph, encompassing 488 unique relationships.

Table 1. Overview of relationships in depressive disorders.

Number	Subject	Object	Relation
1	Major depressive disorder	Irritability	manifests as
2	Major depressive disorder	Depressed mood	manifests as
3	Major depressive disorder	Loss of interest or pleasure	manifests as
4	Major depressive disorder	Changes in sleep patterns	manifests as
5	Major depressive disorder	Decreased energy levels	manifests as
...	...	...	...
483	Unspecified depressive disorder	Appetite change	lasts
484	Unspecified depressive disorder	Weight change	lasts
485	Unspecified depressive disorder	Sexual interest or desire	lasts
486	Unspecified depressive disorder	Sleep disturbance	includes
487	Unspecified depressive disorder	Psychomotor changes	includes

4. Experiments

This section details the evaluations to assess the performance of the proposed approach. Two primary experiments were performed.

- **Zero-shot information extraction on healthcare datasets:** This experiment evaluated the effectiveness of zero-shot information extraction on various biomedical datasets.
- **Zero-shot relationship extraction at the document level:** This experiment focused on extracting relationships from unstructured data at the document level using a zero-shot approach.

4.1. Datasets and Settings

We employed several biomedical datasets to evaluate the effectiveness of zero-shot information extraction from healthcare data. In addition, the DocRED [38] and Re-DocRED datasets [39] were used for document-level information extraction, with a specific focus on extracting and systematizing depression-related information through complex relationship extraction.

Table 2. Summary of statistics for all datasets.

Category	Dataset Name	Source	Document Count	Primary Entity Type	Entity Count
Biomedical Dataset	BC5-Chemical [40]	PubMed Abstracts	1,500	Chemicals (Drugs)	N/A
	BC5-Disease [40]	PubMed Abstracts	1,500	Diseases	N/A
	NCBI-Disease [41]	PubMed Abstracts	793	Diseases	6,892
	BC2GM [42]	PubMed Abstracts	N/A	Genes and Alternative Gene Products	N/A
	JNLPBA [43]	PubMed Abstracts	N/A	Proteins, DNA, RNA, Cell Lines, Cell Types	N/A
Document-level Information Extraction	DocRED [38]	Wikipedia	9,228	Relationships	57,263 Triples
	Re-DocRED [39]	Wikipedia	11,854	Relationships	70,608 Triples

4.1.1. BioCreative Datasets

- **BC5-Chemical and BC5-Disease [40]:** Derived from the BioCreative V chemical-disease relation corpus, these datasets focus on exploring interactions between drugs and diseases. Each dataset includes 1,500 PubMed abstracts (evenly split) for training, development, and testing. We used a preprocessed version by Crichton et al., focusing on the named entity recognition (NER) of chemicals and diseases without relationship labels.
- **NCBI Disease [41]:** Provided by the National Center for Biotechnology Information (NCBI), this dataset includes 793 PubMed abstracts with 6,892 disease mentions linked to 790 unique disease entities. We used a preprocessed version by Crichton et al. for training, development, and testing splits.
- **BC2GM [42]:** Originating from the BioCreative II gene mention corpus, this dataset consists of sentences from PubMed abstracts with manually tagged genes and alternative gene entities. Our study focused on gene entity annotations using a version of the dataset separated for development by Crichton et al.
- **JNLPBA [43]:** Designed for applications in molecular biology, this dataset focuses on NER for entity types such as proteins, DNA, RNA, cell lines, and cell types. We focused on entity mention detection without differentiating between entity types, using the same splits as those of Crichton et al.

4.1.2. Document-Level Relationship Extraction Datasets

- **DocRED [38]:** This dataset was constructed using Wikipedia and Wikidata for document-level relationship extraction. It contains 9,228 documents and 57,263 relationship triples, covering 96 predefined relationship types. DocRED was used to evaluate the ability to extract relationships from complex texts spanning multiple sentences.
- **Re-DocRED [39]:** Re-DocRED, which is an expanded version of DocRED, includes additional positive cases (11,854 documents and 70,608 relationship triples) and incorporates relationship types and scenarios that are not addressed by DocRED. This dataset is useful for research aimed at identifying diverse and in-depth relationship patterns within documents.

Table 2 summarizes the statistics for all of the datasets. This structured approach allowed a comprehensive evaluation of the proposed zero-shot information extraction methodology, particularly in the context of mental health data. This ensured a robust assessment and valuable insights into its applicability.

4.2. Experimental Results

This section presents the findings of the experiments evaluating our zero-shot information extraction approach.

4.2.1. Zero-Shot Information Extraction on Healthcare Datasets

Table 3. Performance Comparison of Zero-Shot Information Extraction on Healthcare Datasets

	Supervised Evaluation		Zero-shot Evaluation	
	BioBERT	PubMedBERT	In-context learning	Proposed method
NCBI	89.1	97.8	51.4	85.4
BC5-disease	84.7	85.6	73.0	87.3
BC5-chem	92.9	93.3	43.6	88.5
BC2GM	83.8	84.5	41.1	67.2
JNLPBA	78.6	79.1	48.3	49.7

Our study represents significant advancements in zero-shot information extraction. Unlike traditional methods that require training data, our approach demonstrates the ability to make accurate

predictions without training data. Our method outperformed existing methods when applied to five medical datasets (NCBI Disease [41], BC5-Chemical, BC5-Disease [40], BC2GM [42], and JNLPBA [43]) Notably, it achieved significant performance gains on the NCBI Disease, BC5-Disease, and BC5-Chemical datasets. These results suggest that our research has the potential to push the boundaries of current zero-shot learning [44] capabilities in the medical domain significantly.

Furthermore, compared to traditional supervised learning models ,such as BioBERT [45] and PubMedBERT [46], our approach maintained competitive performance while requiring significantly less data. These benchmark models were selected based on their relevance to this study. They are pretrained models that are established in biomedical text analysis, making them directly relevant to our objectives. Selecting these domain-specific models helped to demonstrate the advantages and practicality of our zero-shot approach.

These findings suggest that zero-shot information extraction has the potential to revolutionize the automatic analysis of medical data. This study provides a practical implementation of this innovative approach, paving the way for future research and broader application possibilities. This progress emphasizes the transformative impact of advanced AI methodologies in handling complex healthcare data, potentially leading to more efficient and accessible medical analytics.

4.2.2. Zero-Shot Relationship Extraction at the Document Level

Table 4. Performance comparison of document-level zero-shot relationship extraction methods.

	Existing Large Language Model			Proposed Method	
	LLaMA2-7B	Flan-T5-XXL	LLaMA2-13B	without entity linking	with entity linking
DocRED	1.2	4.4	4.0	7.8	9.8
RE-DocRED	1.9	4.3	3.5	7.5	9.2

Our investigation of document-level relationship extraction demonstrates the potential of zero-shot approaches in advancing this field. The relevant results are presented in Table 3. Implementing entity linking significantly improved the performance com-pared to its absence. For example, on the DocRED dataset [38], the F1 score increased from 7.803 (without entity linking) to 9.844 (with entity linking). Similar improvements were observed for Re-DocRED [39] (7.527 to 9.150).

These results highlight the importance of entity linking in enhancing the overall annotation accuracy, particularly in zero-shot learning scenarios, where model adaptability is crucial. Our model not only competed well with established models such as LLaMA2-7B, LLaMA2-13B [47], and Flan-T5-XXL [48], but also exhibited consistent competitive strength across the datasets. Notably, the inclusion of annotation and entity-linking processes significantly improved the performance of our model. This approach enables more accurate capturing of the contextual nuances and complex interactions between entities within documents.

Annotations help the model to develop a deeper understanding of specific cases within datasets, foster better connections between entities, and enhance its grasp of their interrelationships within documents. This method is particularly effective for zero-shot learning. Moreover, the entity-linking process strengthens the semantic connections and significantly improves the accuracy of the extracted information. This method is especially crucial for complex datasets containing various entities and relationships. By incorporating these additional steps, our proposed model achieved higher F1 scores than existing models. This highlights the importance of additional information in extracting complex relationships at the document level, thereby validating the suitability of our model for constructing depression-related knowledge graphs.

The capability of the model to extract and link entities and their relationships within complex medical data accurately is essential for systematically mapping the interactions among various symp-toms, diagnoses, and treatment methods associated with depression. Our approach integrates this information, providing medical professionals and researchers with rich resources to improve the under-

standing and management of depression. Therefore, the technological innovations and performance of our model are expected to contribute significantly to in-depth research on data-driven treatment strategies for depression.

5. Discussion

This study has investigated the effectiveness of entity linking combined with zero-shot information extraction using LLMs in the medical data domain. The findings highlight the critical role of entity linking in this context while also revealing some methodological limitations that offer opportunities for future research and applications in healthcare.

- **Performance evaluation and interpretation**

The results confirm that incorporating entity linking with LLMs significantly enhances zero-shot information extraction for medical data. The performance of our model surpassed that of traditional zero-shot approaches, demonstrating that these techniques can compete with conventional supervised learning methods. This is particularly valuable in healthcare, where data labeling is expensive and data privacy is paramount.

- **Importance of entity linking**

Entity linking plays a vital role in ensuring data consistency and boosting the model performance. In this study, this went beyond simple identification tasks. By significantly improving the overall data accuracy, entity linking underscores its importance in maintaining the integrity and usefulness of medical information systems.

- **Methodological limitations and future directions**

This study used a limited set of datasets, which potentially affected the generalizability of the findings. Future studies should address this issue by exploring a broader range of medical datasets and incorporating a wider variety of entity types. This will help to validate and extend the applicability of the proposed method.

- **Potential applications in healthcare**

The constructed medical knowledge graph serves as a critical tool for the systematic analysis of complex medical data and disease states. It has the potential to be integrated into real-time patient management systems to improve both diagnosis and ongoing patient care. Furthermore, the model of synergized LLMs and knowledge graphs suggests the potential benefits of integrating LLMs with knowledge graphs [49]. This approach can leverage the NLP capabilities of LLMs to interpret complex medical data and provide more accurate disease diagnosis and treatment predictions. This holds promise for more precise analysis of the diverse manifestations of depression and the development of effective personalized treatment plans.

A deeper understanding of depression and other complex health conditions can be gained by advancing the methodologies outlined in this study. This will provide richer resources for healthcare professionals and researchers, ultimately improving diagnostic and treatment strategies. The synergy between LLMs and knowledge graphs not only fosters richer data interaction, but also lays the groundwork for transformative changes in medical research and practice, paving the way for innovative healthcare solutions and improved patient outcomes.

6. Conclusions

This study has proposed a novel approach that integrates medical knowledge graphs with LLMs to enhance information extraction for mental health disorders. Our methodology effectively addresses complex mental health conditions, such as depression, by leveraging the structural advantages of knowledge graphs and the robust predictive capabilities of LLMs. This approach significantly improves the accuracy and consistency of information extraction, particularly using entity linking and zero-shot information extraction techniques.

We acknowledge the limitations of this study, particularly the limited dataset. Future research plans include validating the versatility of our methodology across diverse medical datasets and expanding the types of entities involved. These efforts are crucial to strengthen the validity of our approach further and explore its practical applicability in the healthcare sector.

In conclusion, this study demonstrates significant advancements in the field of mental health through the use of medical knowledge graphs and LLMs for information extraction. It provides a powerful tool that can contribute to the diagnosis and treatment of various diseases, marking a notable step forward in integrating advanced AI technologies into healthcare.

Author Contributions: Conceptualization, C.P., H.L. and O.J.; methodology, C.P.; software, C.P.; validation, C.P.; formal analysis, C.P., H.L. and O.J.; investigation, C.P.; resources, C.P.; data curation, C.P.; writing—original draft preparation, C.P.; writing—review and editing, C.P., H.L. and O.J.; visualization, C.P.; supervision, O.J.; project administration, C.P., H.L. and O.J.; funding acquisition, O.J. All authors have read and agreed to the published version of the manuscript.

Funding: This work was supported by Gachon University research fund of 2023 (GCU-202300690001).

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The data presented in this study are available in this article.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Gutierrez, Bernal Jimenez and McNeal, Nikolas and Washington, Clay and Chen, You and Li, Lang and Sun, Huan and Su, Yu. Thinking about gpt-3 in-context learning for biomedical ie? think again. *arXiv preprint arXiv:2203.08410*, 2022.
2. Wang, Yuqing and Zhao, Yun and Petzold, Linda. Are large language models ready for healthcare? a comparative study on clinical language understanding. In *Machine Learning for Healthcare Conference*, pages 804–823, 2023, PMLR.
3. Li, Qing and Wang, Yichen and You, Tao and Lu, Yantao. BioKnowPrompt: Incorporating imprecise knowledge into prompt-tuning verbalizer with biomedical text for relation extraction. *Information Sciences*, 617, pages 346–358, 2022, Elsevier.
4. Kartchner, David and Ramalingam, Selvi and Al-Hussaini, Irfan and Kronick, Olivia and Mitchell, Cassie. Zero-Shot Information Extraction for Clinical Meta-Analysis using Large Language Models. In *The 22nd Workshop on Biomedical Natural Language Processing and BioNLP Shared Tasks*, pages 396–405, 2023.
5. Chia, Yew Ken and Bing, Lidong and Poria, Soujanya and Si, Luo. RelationPrompt: Leveraging prompts to generate synthetic data for zero-shot relation triplet extraction. *arXiv preprint arXiv:2203.09101*, 2022.
6. Wang, Chenguang and Liu, Xiao and Chen, Zui and Hong, Haoyun and Tang, Jie and Song, Dawn. Zero-shot information extraction as a unified text-to-triple translation. *arXiv preprint arXiv:2109.11171*, 2021.
7. Li, Junpeng and Jia, Zixia and Zheng, Zilong. Semi-automatic data enhancement for document-level relation extraction with distant supervision from large language models. *arXiv preprint arXiv:2311.07314*, 2023.
8. Gyrard, Amelie and Boudaoud, Karima. Interdisciplinary iot and emotion knowledge graph-based recommendation system to boost mental health. *Applied Sciences*, 12(19), pages 9712, 2022, MDPI.
9. Svenaeus, Fredrik. Diagnosing mental disorders and saving the normal: American Psychiatric Association, 2013. Diagnostic and statistical manual of mental disorders, American Psychiatric Publishing: Washington, DC. 991 pp., ISBN: 978-0890425558. Price: \$122.70. *Medicine, Health Care and Philosophy*, 17, pages 241–244, 2014, Springer.
10. Finlayson, Samuel G and LePendur, Paea and Shah, Nigam H. Building the graph of medicine from millions of clinical narratives. *Scientific data*, 1(1), pages 1–9, 2014, Nature Publishing Group.
11. Rotmensch, Maya and Halpern, Yoni and Tlimat, Abdulkhakim and Horng, Steven and Sontag, David. Learning a health knowledge graph from electronic medical records. *Scientific reports*, 7(1), pages 5994, 2017, Nature Publishing Group UK London.
12. Zhao, Chao and Jiang, Jingchi and Xu, Zhiming and Guan, Yi. A study of EMR-based medical knowledge network and its applications. *Computer methods and programs in biomedicine*, 143, pages 13–23, 2017, Elsevier.

13. Rotmensch, Maya and Halpern, Yoni and Tlimat, Abdulhakim and Horng, Steven and Sontag, David. Learning a health knowledge graph from electronic medical records. *Scientific reports*, 7(1), pages 5994, 2017, Nature Publishing Group UK London.
14. Zhang, Kunli and Li, Kaixiang and Ma, Hongchao and Yue, Donghui and Zhuang, Lei. Construction of MeSH-like obstetric knowledge graph. In *2018 International Conference on Cyber-Enabled Distributed Computing and Knowledge Discovery (CyberC)*, pages 160–1608, 2018, IEEE.
15. He, Kai and Yao, Lixia and Zhang, JiaWei and Li, Yufei and Li, Chen. Construction of genealogical knowledge graphs from obituaries: Multitask neural network extraction system. *Journal of medical Internet research*, 23(8), pages e25670, 2021, JMIR Publications Toronto, Canada.
16. Sun, Haixia and Xiao, Jin and Zhu, Wei and He, Yilong and Zhang, Sheng and Xu, Xiaowei and Hou, Li and Li, Jiao and Ni, Yuan and Xie, Guotong and others. Medical knowledge graph to enhance fraud, waste, and abuse detection on claim data: Model development and performance evaluation. *JMIR Medical Informatics*, 8(7), pages e17653, 2020, JMIR Publications Inc., Toronto, Canada.
17. Sainz, Oscar and de Lacalle, Oier Lopez and Labaka, Gorka and Barrena, Ander and Agirre, Eneko. Label verbalization and entailment for effective zero-and few-shot relation extraction. *arXiv preprint arXiv:2109.03659*, 2021.
18. Sainz, Oscar and Gonzalez-Dios, Itziar and de Lacalle, Oier Lopez and Min, Bonan and Agirre, Eneko. Textual entailment for event argument extraction: Zero-and few-shot with multi-source learning. *arXiv preprint arXiv:2205.01376*, 2022.
19. Wei, Xiang and Cui, Xingyu and Cheng, Ning and Wang, Xiaobin and Zhang, Xin and Huang, Shen and Xie, Pengjun and Xu, Jinan and Chen, Yufeng and Zhang, Meishan and others. Zero-shot information extraction via chatting with chatgpt. *arXiv preprint arXiv:2302.10205*, 2023.
20. Jeblick, Katharina and Schachtner, Balthasar and Dextl, Jakob and Mittermeier, Andreas and Stüber, Anna Theresa and Topalis, Johanna and Weber, Tobias and Wesp, Philipp and Sabel, Bastian Oliver and Ricke, Jens and others. ChatGPT makes medicine easy to swallow: An exploratory case study on simplified radiology reports. *European radiology*, pages 1–9, 2023, Springer.
21. Hu, Danqing and Liu, Bing and Zhu, Xiaofeng and Lu, Xudong and Wu, Nan. Zero-shot information extraction from radiological reports using ChatGPT. *International Journal of Medical Informatics*, 183, pages 105321, 2024, Elsevier.
22. Al-Moslimi, Tareq and Ocaña, Marc Gallofré and Opdahl, Andreas L and Veres, Csaba. Named entity extraction for knowledge graphs: A literature overview. *IEEE Access*, 8, pages 32862–32881, 2020, IEEE.
23. Peeters, Ralph and Bizer, Christian. Using chatgpt for entity matching. In *European Conference on Advances in Databases and Information Systems*, pages 221–230, 2023, Springer.
24. Pan, Shirui and Luo, Linhao and Wang, Yufei and Chen, Chen and Wang, Jiapu and Wu, Xindong. Unifying large language models and knowledge graphs: A roadmap. *IEEE Transactions on Knowledge and Data Engineering*, 2024, IEEE.
25. Ye, Junjie and Chen, Xuanning and Xu, Nuo and Zu, Can and Shao, Zekai and Liu, Shichun and Cui, Yuhan and Zhou, Zeyang and Gong, Chao and Shen, Yang and others. A comprehensive capability analysis of gpt-3 and gpt-3.5 series models. *arXiv preprint arXiv:2303.10420*, 2023.
26. Wang, Xiaoying and Peng, Baolin and Li, Kun and Zhou, Jingyan and Meng, Helen. Sgp-tod: Building task bots effortlessly via schema-guided llm prompting. *arXiv preprint arXiv:2305.09067*, 2023.
27. Wei, Jason and Bosma, Maarten and Zhao, Vincent Y and Guu, Kelvin and Yu, Adams Wei and Lester, Brian and Du, Nan and Dai, Andrew M and Le, Quoc V. Finetuned language models are zero-shot learners. *arXiv preprint arXiv:2109.01652*, 2021.
28. Zhou, Wenxuan and Zhang, Sheng and Gu, Yu and Chen, Muhao and Poon, Hoifung. Universalner: Targeted distillation from large language models for open named entity recognition. *arXiv preprint arXiv:2308.03279*, 2023.
29. Košprdić, Miloš and Prodanović, Nikola and Ljajić, Adela and Bašaragin, Bojana and Milosevic, Nikola. From Zero to Hero: Harnessing Transformers for Biomedical Named Entity Recognition in Zero-and Few-shot Contexts. Available at SSRN 4463335, 2023.

31. Chen, Yi and Jiang, Haiyun and Liu, Lemao and Shi, Shuming and Fan, Chuang and Yang, Min and Xu, Ruifeng. An empirical study on multiple information sources for zero-shot fine-grained entity typing. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 2668–2678, 2021.
32. Carta, Salvatore and Giuliani, Alessandro and Piano, Leonardo and Podda, Alessandro Sebastian and Pompianu, Livio and Tiddia, Sandro Gabriele. Iterative zero-shot llm prompting for knowledge graph construction. *arXiv preprint arXiv:2307.01128*, 2023.
33. McCusker, Jamie. LOKE: Linked Open Knowledge Extraction for Automated Knowledge Graph Construction. *arXiv preprint arXiv:2311.09366*, 2023.
34. Sainz, Oscar and García-Ferrero, Iker and Agerri, Rodrigo and de Lacalle, Oier Lopez and Rigau, German and Agirre, Eneko. Gollie: Annotation guidelines improve zero-shot information-extraction. *arXiv preprint arXiv:2310.03668*, 2023.
35. Li, Peng and Sun, Tianxiang and Tang, Qiong and Yan, Hang and Wu, Yuanbin and Huang, Xuanjing and Qiu, Xipeng. Codeie: Large code generation models are better few-shot information extractors. *arXiv preprint arXiv:2305.05711*, 2023.
36. Papaluca, Andrea and Krefl, Daniel and Rodriguez, Sergio Mendez and Lensky, Artem and Suominen, Hanna. Zero-and Few-Shots Knowledge Graph Triplet Extraction with Large Language Models. *arXiv preprint arXiv:2312.01954*, 2023.
37. Wu, Xuehong and Duan, Junwen and Pan, Yi and Li, Min. Medical knowledge graph: Data sources, construction, reasoning, and applications. *Big Data Mining and Analytics*, 6(2), pages 201–217, 2023, TUP.
38. Yao, Yuan and Ye, Deming and Li, Peng and Han, Xu and Lin, Yankai and Liu, Zhenghao and Liu, Zhiyuan and Huang, Lixin and Zhou, Jie and Sun, Maosong. DocRED: A large-scale document-level relation extraction dataset. *arXiv preprint arXiv:1906.06127*, 2019.
39. Tan, Qingyu and Xu, Lu and Bing, Lidong and Ng, Hwee Tou and Aljunied, Sharifah Mahani. Revisiting DocRED—Addressing the False Negative Problem in Relation Extraction. *arXiv preprint arXiv:2205.12696*, 2022.
40. Li, Jiao and Sun, Yueping and Johnson, Robin J and Sciaky, Daniela and Wei, Chih-Hsuan and Leaman, Robert and Davis, Allan Peter and Mattingly, Carolyn J and Wieggers, Thomas C and Lu, Zhiyong. BioCreative V CDR task corpus: A resource for chemical disease relation extraction. *Database*, 2016, Oxford Academic.
41. Doğan, Rezarta Islamaj and Leaman, Robert and Lu, Zhiyong. NCBI disease corpus: A resource for disease name recognition and concept normalization. *Journal of biomedical informatics*, 47, pages 1–10, 2014, Elsevier.
42. Smith, Larry and Tanabe, Lorraine K and Ando, Rie Johnson nee and Kuo, Cheng-Ju and Chung, I-Fang and Hsu, Chun-Nan and Lin, Yu-Shi and Klinger, Roman and Friedrich, Christoph M and Ganchev, Kuzman and others. Overview of BioCreative II gene mention recognition. *Genome biology*, 9, pages 1–19, 2008, Springer.
43. Collier, Nigel and Ohta, Tomoko and Tsuruoka, Yoshimasa and Tateisi, Yuka and Kim, Jin-Dong. Introduction to the bio-entity recognition task at JNLPBA. In *Proceedings of the International Joint Workshop on Natural Language Processing in Biomedicine and its Applications (NLPBA/BioNLP)*, pages 73–78, 2004.
44. Gutierrez, Bernal Jimenez and McNeal, Nikolas and Washington, Clay and Chen, You and Li, Lang and Sun, Huan and Su, Yu. Thinking about gpt-3 in-context learning for biomedical ie? think again. *arXiv preprint arXiv:2203.08410*, 2022.
45. Lee, Jinhyuk and Yoon, Wonjin and Kim, Sungdong and Kim, Donghyeon and Kim, Sunkyu and So, Chan Ho and Kang, Jaewoo. BioBERT: A pre-trained biomedical language representation model for biomedical text mining. *Bioinformatics*, 36(4), pages 1234–1240, 2020, Oxford University Press.
46. Gu, Yu and Tinn, Robert and Cheng, Hao and Lucas, Michael and Usuyama, Naoto and Liu, Xiaodong and Naumann, Tristan and Gao, Jianfeng and Poon, Hoifung. Domain-specific language model pretraining for biomedical natural language processing. *ACM Transactions on Computing for Healthcare (HEALTH)*, 3(1), pages 1–23, 2021, ACM New York, NY.
47. Touvron, Hugo and Martin, Louis and Stone, Kevin and Albert, Peter and Almahairi, Amjad and Babaei, Yasmine and Bashlykov, Nikolay and Batra, Soumya and Bhargava, Prajjwal and Bhosale, Shruti and others. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023.
48. Chung, Hyung Won and Hou, Le and Longpre, Shayne and Zoph, Barret and Tay, Yi and Fedus, William and Li, Yunxuan and Wang, Xuezhi and Dehghani, Mostafa and Brahma, Siddhartha and others. Scaling instruction-finetuned language models. *Journal of Machine Learning Research*, 25(70), pages 1–53, 2024.

49. Pan, Shirui and Luo, Linhao and Wang, Yufei and Chen, Chen and Wang, Jiapu and Wu, Xindong. Unifying large language models and knowledge graphs: A roadmap. *IEEE Transactions on Knowledge and Data Engineering*, 2024, IEEE.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.