



RESEARCH ARTICLE

An enhanced hybrid GCNN-MHA-GRU approach for symptom-to-medicine recommendation by utilizing textual analysis of customer reviews

M. Balamurugan¹, A. Bharathiraja^{2*}

Abstract

Medication recommendation is essential in improving patient treatment and minimizing the occurrence of undesirable effects, yet current approaches prove to be incompetent in addressing sophisticated relationships between syndromes and customer feedback. This study mitigates this gap by presenting a sophisticated model of symptom-to-medicine drug suggestion that deploys state-of-the-art machine learning algorithms for enhanced precision and customization in offering drug suggestions. The novelty lies in the combination of customer reviews with a hybrid model made up of graph convolutional neural networks (GCNNs), multi-head attention, and gated recurrent units (GRUs) to extract complex relationships and sequential dependencies. The competitive game optimizer also further optimizes recommendations to provide solid and personalized treatment recommendations. The approach includes text preprocessing, numerical transformation via TF-IDF and Word2Vec, and evaluation against baseline models using accuracy, precision, recall, and F1-score. Key results show the better performance of the model with 95.65% accuracy, F1 score of 95.12%, and PRAUC of 0.9857, reflecting outstanding precision-recall trade-offs. The Jaccard similarity index of 0.9514 and mean average precision of 0.9725 reflect the effectiveness of the model in providing relevant recommendations. The results highlight the importance of the combination of varied data sources and sophisticated optimization methods, enabling better patient outcomes and revolutionary possibilities in healthcare systems.

Keywords: Symptom-to-medicine recommendation, Machine learning techniques, Personalized medicine, Graph convolutional neural networks, Customer reviews integration, Advanced optimization techniques, Tailored treatment recommendations.

Introduction

In the development of a symptom-to-medicine recommendation model, the importance of accurate medicine recommendations cannot be overstated. Accurate recommendations significantly enhance patient outcomes

by facilitating effective treatment selection, reducing trial and error, and increasing patient compliance, while also lowering healthcare costs (Jacobs *et al.*, 2021). Precise recommendations, tailored to individual patient needs, consider factors like genetics, lifestyle, and comorbidities, enabling more effective and personalized therapies (Korytkowski *et al.*, 2022). However, developing such models comes with significant challenges. Variability in symptom presentation, patient history, and coexisting conditions complicates the recommendation process (Zhang *et al.*, 2022). Symptoms can manifest differently across individuals due to factors like age, gender, and genetic predispositions, leading to potential misinterpretations and inappropriate medication recommendations (Phan *et al.*, 2024). Additionally, diverse medication responses and potential drug interactions necessitate sophisticated algorithms to ensure personalized and safe treatments. To address these challenges, customer reviews play a crucial role in healthcare. They provide valuable insights into the real-world effectiveness and safety of medications, reflecting patient experiences beyond clinical trials (Swain

¹Professor and Chair, School of Computer Science, Engineering & Applications, Bharathidasan University, Tiruchirappalli – 620 023

²Ph.D. Scholar, School of Computer Science, Engineering & Applications, Bharathidasan University, Tiruchirappalli – 620 023

***Corresponding Author:** A. Bharathiraja, Ph.D. Scholar, School of Computer Science, Engineering & Applications, Bharathidasan University, Tiruchirappalli – 620 023, E-Mail: bharathiraja.a@gmail.com

How to cite this article: Balamurugan, M., Bharathiraja, A. (2025). An enhanced hybrid GCNN-MHA-GRU approach for symptom-to-medicine recommendation by utilizing textual analysis of customer reviews. *The Scientific Temper*, **16**(6):4415-4429.

Doi: 10.58414/SCIENTIFICTEMPER.2025.16.6.15

Source of support: Nil

Conflict of interest: None.

et al., 2024). By analyzing trends and applying sentiment analysis to customer reviews, researchers can identify effective medications and capture rare or unexpected side effects, enhancing the understanding of medication safety (Gawich and Alfonse, 2022; Sreedhar *et al.*, 2024). This holistic approach complements clinical data, offering valuable insights for better healthcare outcomes and patient education. Incorporating these insights into a symptom-to-medicine recommendation model can improve the accuracy and personalization of recommendations, ultimately leading to better patient care and outcomes. To effectively integrate these insights, it is essential to preprocess the text. This involves refining the textual data by eliminating noise and organizing the text for better analysis.

In developing a symptom-to-medicine recommendation model, text preprocessing is crucial for refining textual data by removing noise and structuring text for analysis. Techniques such as correcting abbreviations, removing repeated syllables, fixing typos, and formalizing slang, along with automatic steps like case folding and removing numbers and emoticons, enhance data quality and improve the accuracy and efficiency of NLP applications. However, advancements like BERT embeddings and DNN architectures can minimize the need for extensive preprocessing, as demonstrated by experiments where BERT combined with CNN produced superior classification performance (Kurniasih and Manik, 2022). Using these preprocessing techniques, TF-IDF vectorization plays a crucial role in transforming text data into numerical formats by computing term frequency (TF) and inverse document frequency (IDF). TF-IDF vectorization is essential for converting text data into numerical representations by calculating TF and IDF. This technique highlights important words and reduces the weight of common terms, capturing meaningful patterns and reducing noise. In medical applications, TF-IDF aids in predicting disease diagnoses by analyzing symptoms and diseases, enhancing accuracy through cosine similarity (Wei *et al.*, 2024; Aszani *et al.*, 2023). Additionally, word embeddings effectively capture the semantic relationships in text by representing words as dense vectors, which allows models to grasp the context and usage of words. Word embeddings capture semantic relationships within text data by representing words as dense vectors, enabling models to understand word context and usage. In a symptom-to-medicine recommendation model, word embeddings are crucial for capturing semantic relationships between symptoms and medicines, resulting in more accurate and personalized recommendations (Lin and Bu, 2022). Word2Vec generates these embeddings, representing similar symptoms and medications closely in the embedding space, which enhances the system's ability to identify related symptoms and medications, providing more accurate recommendations (Park *et al.*,

2024). Feature selection further improves symptom-based medicine models by reducing dimensionality and enhancing detection efficiency (Zhou, 2024). Mutual information-based feature selection identifies informative features, capturing non-linear relationships, and reduces model complexity, leading to accurate recommendations (Sivaiah *et al.*, 2024), improving model performance and efficiency.

In developing a symptom-to-medicine recommendation model, graph convolutional neural networks (GCNNs) play a crucial role by leveraging node connectivity and features to capture intricate relationships among symptoms, diseases, and medicines. GCNNs enhance accuracy and effectiveness by aggregating information from connected nodes, capturing dependencies between symptoms and treatments, and outperforming traditional models with richer feature representations and improved predictive accuracy (Shou *et al.*, 2022). Combining GCNNs with multi-head attention and gated recurrent units (GRUs) in a hybrid architecture further enhances the model's performance. This synergistic approach captures complex relationships, spatial correlations, and sequential dependencies, with GCNNs modeling interdependencies, multi-head attention focusing on relevant features, and GRUs managing temporal dynamics. This integration leads to more accurate, relevant, and adaptable treatment suggestions, leveraging structural, contextual, and sequential information to improve patient outcomes and care (Cheng *et al.*, 2021; Wang *et al.*, 2024). The competitive game optimizer (CGO) method, utilizing game theory principles and gradient descent, optimizes the model by promoting faster convergence to optimal solutions, improving accuracy, reducing overfitting, and enhancing generalization on unseen data. By iteratively updating weights based on gradients, CGO ensures robust and effective model performance, making it a valuable tool in refining the symptom-to-medicine recommendation system (Elmanakhly *et al.*, 2021). Together, these techniques contribute to a more nuanced understanding of relational factors, improved feature representation, and dynamic learning, resulting in personalized and accurate medication recommendations.

The motivation behind this research is to improve patient outcomes by providing accurate and personalized medicine recommendations. By tackling the issues of symptom variability and the diverse backgrounds of patients, this study seeks to create an advanced model that connects symptoms to medicine recommendations. It will utilize cutting-edge machine learning techniques to enhance healthcare delivery and ensure patient safety. The goal is to develop an effective system that combines hybrid architectures and optimization methods to understand complex relationships and offer precise, tailored treatment suggestions.

The major contributions of the research work are as follows:

- Designing an improved hybrid architecture that integrates GCNN, multi-head attention (MHA), and GRU to make better symptom-to-medicine recommendations.
- Leveraging customer reviews to better filter and improve the precision of medication prescriptions and incorporate unstructured data within the prescription recommendation.
- Application of TF-IDF vectorization and Word2Vec embeddings to craft a strong feature set that extracts high-level semantic patterns and relationships as well as dependencies among symptom data.
- Use of the CGO to dynamically adjust model parameters for enhanced predictive performance and diminished overfitting.
- Use a complete assessment strategy consisting of several classification and ranking measures to measure model performance in various contexts.

This paper comprises five sections. The introduction provides background and general information. Section 2 reviews literature related to the proposed model. Section 3 details the methodology, while Section 4 covers system implementation and evaluation. Section 5 discusses the proposed model's significance, limitations and future scope. Finally, Section 6 presents conclusions and future work.

Review of Literatures

This section discusses various existing models that recommend medications based on symptoms.

Cheng *et al.* (2023) suggested a drug recommendation model that combined structured patient demographic information and unstructured patient reviews through Bayesian multitask learning. The model predicted review ratings for satisfaction measures of drugs based on topic modeling and sentiment analysis. Bayesian LASSO was employed for feature selection to remove irrelevant features. Though this method performed better than other methods in terms of accuracy and AUC, its difficulties lay in retrieving relevant information from text and handling the cold start problem. Weaknesses were the small sample size, online reviews' potential bias, and the requirement for medical validation. These limitations highlight the need for more robust text preprocessing methods.

Borchert *et al.* (2024) introduced a preprocessing method for complex entity mentions in biomedical text utilizing generative large language models (LLMs) to enhance recall and accuracy of entity linking. The method was incorporated into the xMEN toolkit and experimented with to measure performance. Limitations are specificity to datasets, reliance on language resources, and absence of studies on LLM biases and computational expenses. It underscores the importance of effective feature selection in improving model performance.

Asghari *et al.* (2023) proposed a hybrid feature selection approach, BC-NMIQ, which integrated best clustering normalized mutual information quantile and incremental association Markov blanket to improve classification performance in high-dimensional medical data. The approach ranked features according to mutual information and fine-tuned selection to remove redundancy. Nevertheless, the research was subject to limitations like possible overfitting, extensive experimentation requirements, scalability issues, and high computational complexity, which may impede real-time clinical use, necessitating more efficient architectures.

Jiang *et al.* (2022) proposed the multi-interest graph convolutional network (MI-GCN) to improve recommendation systems by preserving users' heterogeneous interests using high-order graph convolutions over different subgraphs. The method was superior to classical GCN-based approaches by refining user and item embeddings, generating more personalized suggestions. Nevertheless, the research outlined drawbacks such as performance loss upon layer stacking and dependence on particular clustering techniques, which may hamper flexibility on varied datasets and recommendation scenarios. These limitations point to the need for more adaptable and interpretable models.

Bi *et al.* (2023) proposed a brain region gene community network (BG-CN) and a community graph convolutional network (Com-GCN) to better understand brain information transmission, which can be used for the diagnosis of Alzheimer's disease. The Com-GCN integrated intercommunity and intracommunity convolutions to achieve better interpretability and performance in detecting disease-related brain regions and genes. Nevertheless, the research was subject to limitations like possible overfitting, dependency on the quality of input data, generalizability issues, and heavy computational resource consumption which are more similar to the challenges in deep learning approaches reviewed by Shen *et al.* (2024).

Shen *et al.* suggested using shallow convolutional neural networks in a deep learning ensemble to identify spam reviews by applying multi-view learning methods and textual and non-textual features. This highlights the need for balancing model complexity and performance. The model scored high classification accuracy but was hampered by specificity with regard to datasets, scalability, and difficulties in balancing subjective user opinions with objective measures. This emphasizes the importance of capturing both long-term and short-term preferences in recommendation systems.

Liu *et al.* (2023) improved graph neural network-based recommendation algorithms by incorporating a multi-head attention mechanism and GRUs to understand users' long-term and short-term preferences. Such a scheme represented user-item interactions and adaptively weighted

friend impacts for better recommendation accuracy. Nonetheless, the research was constrained by issues including overfitting, dependency on correct social network information, difficulty in reflecting subtle preferences, and the high demand for computations that might affect scalability under real-time, large-scale settings.

Similarly, Merkelbach *et al.* (2023) also proposed a gated recurrent unit autoencoder to identify ICU patient subgroups based on electronic health records' time series data. The model overcame the irregularity, sparsity, and high dimensionality of the challenges by encoding time series data with positional encodings to support clustering and feature space analysis. The model successfully identified disease patterns and mortality prediction but had limitations like data irregularity, sparsity, and lossy reconstruction of time series data.

Di *et al.* (2022) investigated the introduction of gated architectures, for example, GRUs, into echo state networks (ESNs) in order to address long-term dependencies and enhance prediction accuracy. The research suggested the training of gates exclusively in ESNs by integrating reservoir computing with gated architectures to enable effective training. The model's performance was limited by its computational demands and challenges in managing long-term dependencies, underscoring the necessity for models that can adeptly navigate complex temporal dynamics and relationships.

Wu *et al.* (2023) suggested DAPSNet, a model of recommending drugs by applying patient history and similarity in the disease state for forecasting appropriate and safe recommendations. It applied code and visit-level attention mechanisms in order to embed patient representations through the incorporation of diagnosis, procedure, and drugs. The model learned to optimize more than one loss function and excelled above existing methodologies. Nevertheless, it was limited in measuring the complete patient representations, taking into account the prescription history, and effectively pairing drugs with disease status because of homologous global representations.

Research Gap

The literature reviewed presents some of the limitations of current symptom-to-medicine recommendation models, such as overfitting, scalability, and dependence on particular datasets or clustering methods. Other issues, such as the management of long-term dependencies, capturing subtle user preferences, and combining heterogeneous data sources, are also not addressed. These shortcomings highlight the importance of a strong model that well integrates structured and unstructured data, handles data irregularities, and offers personalized, precise recommendations while being computationally efficient and generalizable to various healthcare settings.

Proposed Methodology

The proposed symptom-to-medicine recommendation model adopts a systematic approach with several steps to improve prescription accuracy based on customer reviews. The procedure starts with data collection and preprocessing, in which textual data from two datasets are cleaned. Preprocessing includes text conversion, tokenization, removal of stop words, lemmatization, stemming, and removal of low- and high-frequency words to eliminate noise and normalize the data. During feature engineering, text data is vectorized as numerical representations using TF-IDF vectorization to reflect term importance and Word2Vec embeddings that learn to maintain semantic relationships.

These are then aggregated into a strong feature set with mutual information selection determining the most informative features. Categorical medicine labels are encoded with one-hot encoding through LabelBinarizer. The architecture of the model is based on a hybrid model combining graph convolutional neural networks (GCNNs), multi-head attention, and gated recurrent units (GRUs). GCNNs are able to represent intricate symptom-disease-medicine relations, multi-head attention strengthens feature representation by concentrating on the most significant symptoms, and GRUs learn sequential dependencies in symptom information. For optimizing the performance of the model, the competitive game optimizer (CGO) is utilized, thereby allowing for efficient parameter tuning for better accuracy and generalization. The suggested

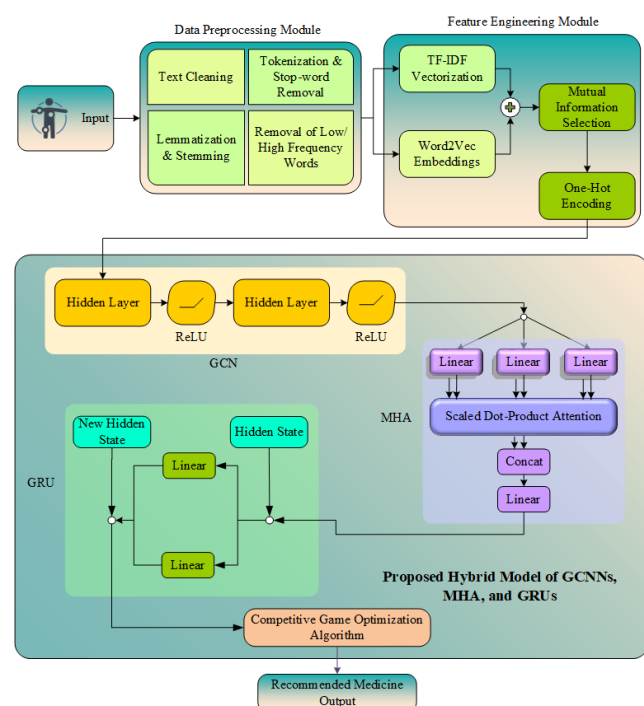


Figure 1: Architecture of the proposed medicine recommendation model

methodology allows for an intelligent, personalized, and effective medicine recommendation system based on symptom descriptions. Figure 1 shows the architecture for the proposed model.

Data Collection and Preprocessing

This stage involves collecting data from Kaggle datasets and refining it through essential steps to get it ready for analysis.

Data Sources

This section introduces two datasets employed in this proposed symptom-to-medicine recommendation model.

Medical recommendation dataset

This Dataset is a useful dataset intended to assist healthcare practitioners in making informed decisions regarding patient treatment. It uses long short-term memory (LSTM) neural networks to make drug recommendations based on patient symptoms and diagnosed illnesses, and is comprised of explicit patient case records containing symptoms, diseases, and prescriptions. The dataset has been carefully annotated to provide accurate model training, with a high prediction accuracy of around 88%. With its strong foundation, the medical recommendation dataset can be enhanced in the future, for example, by incorporating real-time recommendations and enhancing model generalization to further advance its capabilities and ultimately contribute to improved patient outcomes (Dataset 1).

Drug dataset - uses, side effects, and user reviews

It's an exhaustive database of more than 11,000 drugs, containing detailed information about their composition, therapeutic applications, possible side effects, and customer reviews. This vast pool of data is a useful resource for both doctors and patients to make educated choices regarding drugs. By using this dataset, many applications can be built, such as drug categorization, segmentation analysis based on reviews, and recommendation systems on an individual user profile and preference basis. Finally, the drug dataset: uses, side effects, and user reviews can improve the overall quality of healthcare provision by allowing for better and more personalized medication recommendations (Dataset 2).

Preprocessing the Textual Data

The preprocessing stage consists of four essential steps

Converting text, tokenizing, lemmatizing, and removing unnecessary words, all aimed at refining text data for precise symptom-to-medicine recommendations.

Text conversion and normalization

This is an important preprocessing operation in our suggested symptom-to-medicine recommendation model because it greatly improves the quality and uniformity of text data. The process entails a number of important techniques,

such as lowercasing, stripping special characters, and unit length normalization of text, which all contribute to the improvement of the model. By lowering text to lowercase and eliminating case sensitivity, the classification accuracy of the model is enhanced, and its generalization capability and mapping of medical terms to standardized concepts are increased. Text normalization also facilitates effective symptom extraction, language adaptability, and enhanced stopword elimination, spelling correction, and contraction and abbreviation expansion. These preprocessing operations are crucial in preparing high-quality text data for our symptom-to-medicine recommendation model, ultimately making it possible for it to deliver accurate and personalized medicine recommendations from patient symptoms.

Tokenization and stop-word removal

Tokenization splits customer opinions into useful units, e.g., [«This,» «product,» «is,» «amazing,».] from «This product is amazing, but the delivery was late.». Tokenization makes it easy and accurate to analyze data, allowing sentiment analysis systems to identify customer moods by analyzing tokens like «amazing» and «late.». Stop-word removal is also crucial, since it removes irrelevant words such as «the» and «is,» so that the model can concentrate on significant terms. Stop-word removal decreases data volume, increases efficiency and accuracy, and increases relevance. By removing them, algorithms are able to read better in between the lines and know the text's meaning and context. But in other instances, e.g., sentiment analysis, stop words like «not» may be useful, so the decision to eliminate stop words will rely on the particular task and text being processed. This ultimately produces more precise outcomes.

Lemmatization and stemming to standardize words

Lemmatization and stemming both reduce words to a base form, but in different ways. Stemming strips off suffixes, such as «flooding» reducing to «flood,» whereas lemmatization takes context into account and reduces to a meaningful base, such as «better» reducing to «good.». Lemmatization guarantees a normalized, dictionary-found result, whereas stemming is not always a valid word. These methods fine-tune text data, allowing improved semantic interpretation in medicine recommendations by bringing word variations to a common denominator, allowing more precise analysis and enhanced model performance. This improves the precision of medicine recommendations.

Removal of low- and high-frequency words to reduce noise

The proposed model for medicine recommendation enhances its performance by eliminating low- and high-frequency words. This minimizes noise, enhances interpretability, and strengthens model learning. Low- and high-frequency words are identified using criteria such as frequency thresholds, statistical techniques, domain analysis, and contextual salience. Rare but meaningful

low-frequency words such as «diabetes» and «heart» are kept, whereas frequent words such as «the» and «is» are eliminated. This filtering of input data results in improved performance, improved recommendation accuracy, and a more trusted symptom-to-medicine recommendation.

Feature Engineering

This section explores the process of converting text data into numerical formats through TF-IDF and Word2Vec embeddings. It discusses how to combine these features and identify the most relevant ones using mutual information. Additionally, it addresses the encoding of categorical medicine labels for classification purposes with the help of LabelBinarizer.

TF-IDF Vectorization

After preprocessing of text data, it is transformed into TF-IDF vectors in order to encode term significance, especially in customer reviews. TF-IDF makes the reviews more relevant by picking out salient symptoms and minimizing data sparsity. Term Frequency (TF) calculates how frequently a term occurs in a document, and Inverse Document Frequency (IDF) measures its significance within the whole corpus. The expression for IDF can be represented as in Eqn. (1).

$$IDF(t_d) = \log\left(\frac{N_d}{df(t_d)}\right) \quad (1)$$

In Eqn. (1) the variable N_d denote the total number of documents and $df(t_d)$ is the number of documents containing term t_d . The TF-IDF score is the product of TF and IDF values, which emphasizes important terms by reducing common word weight and giving prominence to distinctive ones. This operation turns every document into a vector such that every item is a TF-IDF value for a term. Through TF-IDF utilization, the model of symptom-to-medicine suggestion can efficiently discern and assign significance to the most pertinent terms used in customer comments, resulting in more precise identification of symptoms and better medicine suggestions.

Word2Vec Embeddings

The Word2Vec model is a machine learning method that maps preprocessed text to compact vector representations that preserve contextual meaning. It employs two main architectures, namely Continuous Bag of Words (CBOW) and Skip-gram, which learn to encode words as vectors based on their context. The model begins with tokenization and cleaning, followed by training via a context window with either CBOW or Skip-gram. A shallow one-hidden-layer neural network produces these representations, employing negative sampling and optimization methods such as stochastic gradient descent. The model captures semantic meanings, allowing the recommendation model to learn rich

relationships between medicines and symptoms, improving the accuracy and personalization of recommendations.

Feature Combination

The TF-IDF vectors highlight the importance of terms based on their frequency and how they are distributed across documents, while Word2Vec embeddings focus on the semantic relationships between words. The TF-IDF output for document i can be expressed as:

$$TFIDF_i = [v_{i1}, v_{i2}, \dots, v_{iZ}] \quad (2)$$

Here, v_{i1} represents the TF-IDF weight for term j in document i . The TF-IDF vectorization process converts each document into a vector of size Z . If document i contains n_i words, and the Word2Vec embedding for the k -th word is Z_k , then the document embedding using Word2Vec, denoted as $W2V_i$, is calculated as:

$$W2V_i = \frac{1}{n_i} \sum_{k=1}^{n_i} Z_k \quad (3)$$

where each Z_k is a vector of size d . Thus, $W2V_i$ is also a vector of size d . The combined output C_i is formed by concatenating the $TFIDF_i$ vector with the Word2Vec $W2V_i$ vector to create a single feature vector, represented as:

$$C_i = [TFIDF_i, W2V_i] \quad (4)$$

By merging these features, the model can utilize both types of information to enhance its performance.

Mutual Information (MI) Selection

MI is used for choosing the best informative features among merged TF-IDF and Word2Vec embeddings in our model for symptom-to-medicine recommendations. It holds non-linear correlations between features and the target, and it can decrease overfitting and increase model performance. The MI for a feature A_i with the target variable B is defined as:

$$I(A_i, B) = \sum_{a_i \in A} \sum_{b \in B} p(a_i, b) \log\left(\frac{p(a_i, b)}{p(a_i)p(b)}\right) \quad (5)$$

Here, $p(a_i, b)$ is the joint probability distribution, and $p(a_i)$ and $p(b)$ are marginal probabilities. The features are sorted based on MI scores, and the highest-ranking features are chosen to improve the performance of the model. This optimization improves the accuracy and precision of recommendations using customer reviews.

One-Hot Encoding

In optimizing the feature set resulting from MI selection, one-hot encoding is used to transform categorical variables into a binary matrix, allowing machine learning algorithms

to handle them efficiently. The method converts each category into a binary vector, with one '1' representing the presence of the category and '0's everywhere else. For categorical medicine label encoding in a symptom-to-medicine recommendation model, the LabelEncoder *LabelBinarizer* from *sklearn.preprocessing* the scikit-learn library in Python is used. It applies one-vs-all encoding, making a binary column for every distinct medicine label, where '1' indicates the presence of the label and '0' the absence. The process includes importing, *LabelBinarizer* getting the categorical medicine labels ready, creating an instance of the binarizer, and using the *fit_transform* method to transform labels into binary form. This process makes multi-class classification easier, effectively transforming labels into a format that can be used by different classifiers, and improves the model's capacity to classify symptoms into the right medicines based on customer reviews.

Proposed Hybrid Model Architecture

This subsection explains the combination of GCNNs, MHA, and GRUs to design a hybrid structure for our proposed symptom-to-medicine recommendation model and describes their functions in extracting complicated relations, improving feature selection, sequential dependency modeling, and recommendation accuracy enhancement.

Relational learning in medicine recommendation using GCNN

GCNNs provide a robust instrument for relational learning in medicine, recommending systems that specifically excel in handling intricate relations between symptoms, illnesses, and medication. Unlike deep neural networks, GCNNs are formulated specifically to manage graph-structured information, with entities like symptoms, illnesses, and medication being shown as nodes, and relations amongst them being lines or edges that connect these nodes. With such a configuration, GCNNs can extract automatic features from a graph and obtain complex patterns embedded within them in order to form precise recommendations. It works by combining information from the neighbors of a node through a learned filter, aggregating node features through a weighted sum of neighbor features. This message passing allows the network to progressively improve node representations by taking in information from neighboring nodes. The formal expression of a GCNN layer is given by Eqn. (6).

$$H = \sigma \left(\tilde{D}^{-\frac{1}{2}} \tilde{A} \tilde{D}^{-\frac{1}{2}} X \Theta \right) \quad (6)$$

where H is the new node representations, X is the original feature matrix, $\sigma(\cdot)$ is the activation function (e.g., ReLU), $\tilde{A}$ is the adjacency matrix of the graph with self-loops, $\tilde{D}$ is

the degree matrix, and Θ is the trainable parameter matrix. This mathematical representation enables GCNNs to have a linear scaling with the number of graph edges and thus be efficient for big data.

It is superior in modeling the impact of nearby symptoms and diseases in medical recommendations. They adjust their message-passing process to scan intricate health data, making it possible for subtle interpretation of interrelated medical information. GCNNs fit recommendation tasks because they can deal with irregular data structures and capture structural dependencies. They are able to handle large datasets, integrate new data, and offer intuitive interpretation, scalability, and stability in the presence of missing data. GCNNs can be repeatedly trained to learn new research developments, enhancing their recommendation performance with time. Such flexibility and speed make GCNNs a better option for symptom-to-medicine recommendation systems compared to conventional models. With the use of GCNNs, personalized and contextually appropriate medicine recommendations can be offered, and patient care and treatment outcomes improved. In general, GCNNs provide an effective tool for enhancing medicine recommendation systems.

Enhancing feature selection using MHA

MHA is an important mechanism in feature selection improvement for symptom-to-medicine recommendation models, enhancing the relevance and accuracy of recommendations. MHA enables the model to attend to various regions of input data at the same time, capturing intricate relationships and dependencies that single-head attention mechanisms may fail to capture. In medical recommendation, MHA handles several attention heads in parallel, each of which learns to assign weights to various features from symptom descriptions, generating a more dynamic and richer feature representation. The main contribution of MHA is the identification and ranking of significant features from varied inputs. The mathematical representation of MHA is given by Eqn. (7).

$$MultiHead(Q, K, V) = [head_1, \dots, head_h] W_0 \quad (7)$$

where each head $head_i$ computes attention as:

$$head_i = Attention(QW_i^Q, KW_i^K, VW_i^V) \quad (8)$$

By aggregating information from these multiple heads, the model learns multiple relationships and dependencies of the data and improves its capability to manage complicated interactions between medicines and symptoms. This results in better medicine recommendation accuracy through capturing long-range dependencies and adaptive feature selection.

The capacity of MHA to pay attention to the most significant features in symptom descriptions is key to enhancing recommendation accuracy. It dynamically weights features, placing more emphasis on important symptoms and discarding less useful information. This adaptive attention not only suppresses noise but also offers insight into which features are most important, making the model more interpretable. By combining and ranking features from different data sources, MHA allows for holistic and personalized recommendations and thus is a critical piece in contemporary medicine recommendation systems.

Gated Recurrent Units (GRU) for Sequential Dependencies

GRUs are essential in capturing sequential dependencies in symptom development in our suggested symptom-to-medicine recommendation model. GRUs, through the proper management of sequential data, capture temporal relationships and symptom evolution, improving the model's comprehension of disease development and treatment dynamics. GRUs preserve the sequence of symptom occurrence, enabling the model to take into account the exact timing and symptom progression, which is critical for effective recommendations. GRUs use update and reset gates to control information flow, determining what to keep or forget at every time step. The update gate controls how much of the past hidden state to propagate, and the reset gate enables the network to reset according to new incoming symptoms. This gating allows GRUs to learn long-term dependencies, remembering past symptoms that can signal more complicated medical conditions and selectively forgetting less important symptoms. The mathematical representation of the hidden state update is given in Eqn. (9).

$$h_t = (1 - z_t) \odot h_{t-1} + z_t \odot \tilde{h}_t \quad (9)$$

where z_t is the update gate, $\tilde{h}_t$ is the candidate hidden state, and $\odot$ represents element-wise multiplication.

GRUs enhance the predictive ability of the recommendation system by incorporating attention mechanisms, with the aim of paying attention to important symptoms and making use of sequential data. This enhancement enables customized medicine recommendations through an understanding of the personalized pattern of symptoms for every patient. GRUs' effective memory handling and capacity to accept sequences in both forward and backward directions improve context awareness, resulting in more trustworthy and personalized medicine recommendations. Through capturing the subtle interdependencies and time dynamics of symptom evolution, GRUs greatly improve the accuracy and applicability of the recommendations that the model outputs.

Integration of GCNN, MHA, and GRUs

The synergy between GCNNs, multi-head attention, and GRUs results in a better hybrid architecture for our symptom-to-medicine recommendation model by utilizing their respective strengths. GCNNs extract local patterns and relationships among symptoms, diseases, and medicines through graph structures. Multi-head attention deepens contextualization by dynamically adjusting symptom weights and generating rich embeddings. GRUs handle long-term dependencies and sequential data, providing precise predictions by selectively remembering key information. This integration provides enhanced data representation, context awareness, and temporal dynamics, enhancing the interpretability and noise robustness of the model. The hybrid architecture is efficient and scalable, supporting larger healthcare datasets and generating more knowledgeable recommendations than standard models.

Optimization Strategy

The CGO is a meta-heuristic optimization algorithm used in this research to enhance recommendation models. It does this by simulating a competitive environment where solutions develop over time through exploration and exploitation phases. In our suggested symptom-to-medicine recommendation model, CGO plays a role in feature selection, parameter adjustment, and collaborative filtering, which makes the model more responsive to user preferences. Through a game-theoretic perspective, CGO models optimization as a competitive game so that dynamic model parameter adjustments are possible. With mechanisms such as Levy flights, it can effectively explore the solution space without trapping in local optima, enabling the convergence towards optimal parameter settings at an increased speed. Through this method, not only is prediction accuracy enhanced, but so is the avoidance of overfitting as diversity among the candidate solutions is promoted. By iterative refining and performance optimization, CGO optimizes the model's parameters, resulting in improved performance and stability in recommendations.

Experimental Setup and Result Evaluation

The experimental environment for this study was carried out on Windows 10 Pro 64-bit OS with 8GB RAM and Intel(R) Core (TM) i3-8100 CPU at 3.60GHz. Spyder was the integrated development environment (IDE), used as the simulation tool. Python was the programming language adopted for implementation.

Baselines

In this research, we compare our model with a number of baseline models, such as LEAP, RETAIN, DMNC, GAMENet, SafeDrug, and MICRON, using evaluation metrics like accuracy, F1-score, PRAUC, Jaccard similarity, and mean average precision.

- LEAP is an example-based drug recommendation model for complicated multimorbidity patients, producing treatment sentences and choosing the best drugs while preventing harmful interactions.
- RETAIN is a long-term model that applies a two-level RNN with neural attention for predicting sequences and determining meaningful past visits and clinical factors.
- DMNC uses memory-augmented neural networks for recommendations within the framework of a differentiable neural network.
- GAMENet applies memory-augmented neural networks, incorporating fusion-based GCN, attention-based memory search, and dynamic memory modules combined with RNNs to investigate drug co-occurrences and interactions.
- SafeDrug aims to make safe drug recommendations by capturing molecular structure data and accounting for drug interactions via global and local encoders.
- MICRON leverages a recurrent residual network to update and spread patient medical data while retaining temporal data for future visits (Wu *et al.*, 2023).

These baselines provide a thorough benchmark for evaluating the performance and effectiveness of our model.

Evaluation metrics

The proposed recommendation model is evaluated through a combination of multiple metrics to assess its overall performance. Key classification metrics such as accuracy (Acc), precision (Prec), recall, and F1-score are used to measure how effectively the model identifies relevant medications based on symptoms. To compare precision and recall at different thresholds, the precision-recall area under the curve (PRAUC) is employed. Additionally, ranking metrics like Jaccard similarity, mean average precision (MAP), and mean reciprocal rank (MRR) are applied to assess the quality of the medicine rankings, ensuring that the most suitable medications appear at the top of the recommendations.

Comparative analysis of Performance Metrics Over Epochs

This section describes a comparative evaluation of performance measures across epochs for two data sets. It compares how the performance measures of the suggested model change across training epochs and how they reflect improvement and steadiness in recommendation performance.

Medical Recommendation Dataset

Our proposed symptom-to-medicine recommendation model shows steady improvement with each epoch, as indicated by rising metrics. Table 1 shows the comparison of performance metrics at different epochs based on the medical recommendation dataset (Dataset-1).

Accuracy grew from 94.25% in epoch 20 to 95.25% in epoch 100, which shows improved predictability. Precision

and recall also increased, signifying better positive prediction balance with actual positives. The F1 score showed improvement from 94.16 to 94.85%, indicating overall improvement. Figure 2 (a) shows the graphical representation of the comparative analysis of our proposed model with classification metrics.

Figure 2 (b) shows the graphical representation of the comparative analysis of our proposed model with ranking metrics. PRAUC and Jaccard similarity consistently rose, reflecting improved precision-recall trade-off and similarity of recommended sets. MAP and MRR gains reflect improved ranking quality and relevance of top recommendations. These trends reflect the model's strength and robustness across training iterations.

Drug Dataset

The symptom-to-medicine recommendation model proposed demonstrates consistent improvement throughout epochs on the drug dataset. Table 2 shows the analysis based on the drug dataset.

Between epochs 20 and 100, accuracy rose from 94.25 to 95.65%, demonstrating increased prediction reliability. Precision, recall, and F1-score also rose, depicting a well-balanced trade-off among positive predictions and true positives.

PRAUC and Jaccard similarity measures consistently rose, proving improved precision-recall trade-offs and similarity

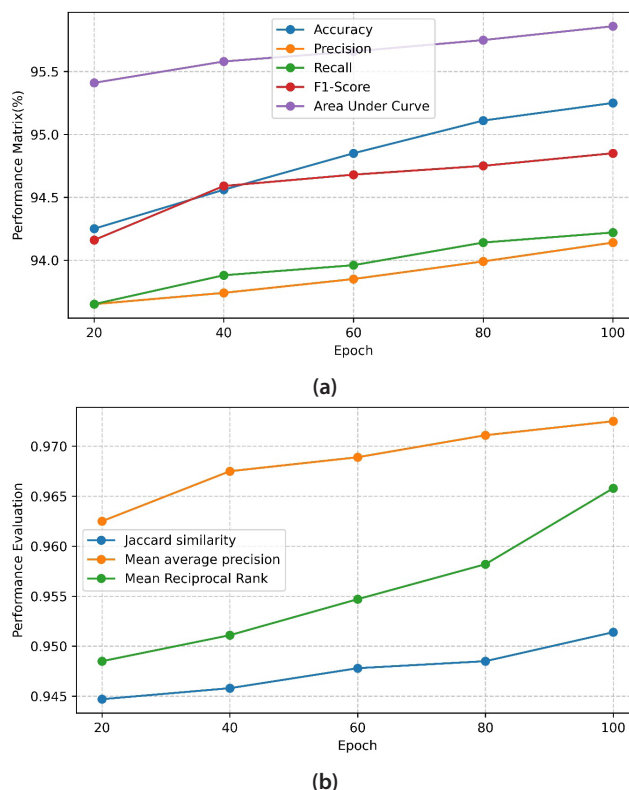


Figure 2: Epoch-based analysis of key metrics using medical recommendation dataset

Table 1: Evaluation of recommendation accuracy and precision using medical recommendation dataset

Epoch	Acc	Prec	Recall	F1-Score	PRAUC	Jaccard similarity	MAP	MRR
20	94.25	93.65	93.65	94.16	95.41	0.9447	0.9625	0.9485
40	94.56	93.74	93.88	94.59	95.58	0.9458	0.9675	0.9511
60	94.85	93.85	93.96	94.68	95.66	0.9478	0.9689	0.9547
80	95.11	93.99	94.14	94.75	95.75	0.9485	0.9711	0.9582
100	95.25	94.14	94.22	94.85	95.86	0.9514	0.9725	0.9658

Table 2: Temporal analysis of model efficiency using drug dataset (Dataset 2)

Epoch	Acc	Prec	Recall	F1-Score	PRAUC	Jaccard similarity	MAP	MRR
20	94.25	93.74	93.85	94.35	94.65	0.9458	0.9625	0.9485
40	94.52	93.89	93.96	94.74	94.74	0.9468	0.9675	0.9511
60	94.96	93.96	94.12	94.88	94.88	0.9478	0.9689	0.9547
80	95.25	94.11	94.44	94.96	95.24	0.9485	0.9752	0.9582
100	95.65	94.52	94.65	95.12	95.65	0.9514	0.9768	0.9658

of recommended sets. The MAP and MRR gains demonstrate the model's capacity for effective prioritization of relevant medicines. These trends validate the model's strength and performance in modeling intricate relationships and providing good recommendations, highlighting its potential to be applied in actual healthcare scenarios.

Model Comparison

This section contrasts the performance of the suggested model with several baseline models using the two datasets and emphasizes enhancements in metrics like accuracy, F1-score, PRAUC, Jaccard similarity, and MAP to show that the model performs better in providing accurate and relevant medicine recommendations (Figure 3).

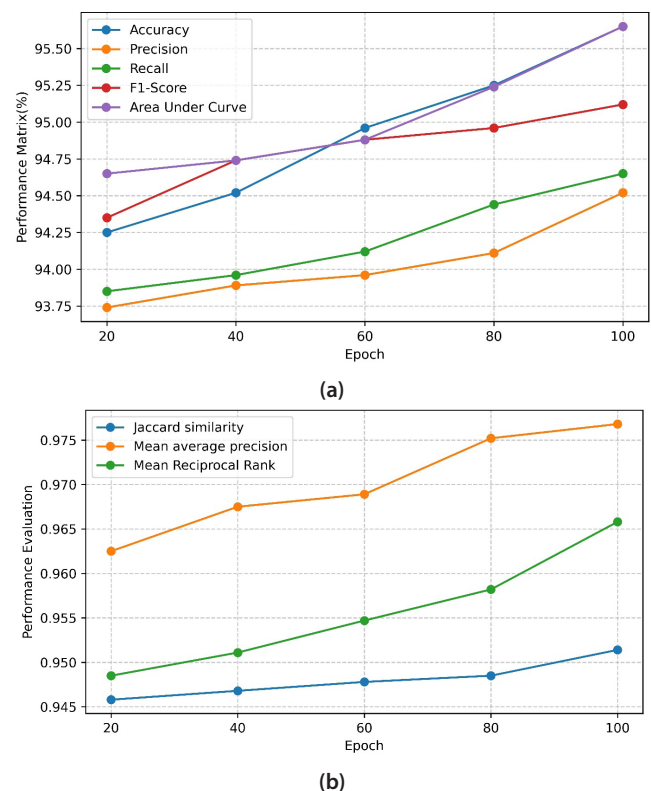
Medical Recommendation Dataset

The proposed model performs better than baseline models on all major metrics. At 95.25 accuracy, it beats LEAP (89.71%), DMNC (89.85%), RETAIN (90.28%), GAMENet (91.15%), MICRON (92.45%), and SafeDrug (93.52%). Table 3 shows the analysis based on dataset 1 (Table 3).

The 94.85% F1 score emphasizes a good precision-recall trade-off, beating all baselines. The 0.9857 PRAUC measures higher precision-recall performance, while the 0.9514 Jaccard similarity measures improved recommendation set overlap. The MAP value of 0.9725 indicates enhanced ranking quality. These findings show that the suggested model is better capable of capturing sophisticated relations and making accurate, individualized suggestions. Figure 4 (a) & (b) shows the graphical comparison of our proposed model.

Drug Dataset

The symptom-to-medicine recommendation model proposed performs better than current models on major

**Figure 3:** Model performance metrics across training epochs using drug dataset (Dataset 2)

metrics. It has the highest accuracy (95.24%) and F1-score (94.88%) compared to LEAP, DMNC, RETAIN, GAMENet, MICRON, and SafeDrug. The following Table 4 indicates analysis from the drug dataset (Dataset 2).

The PRAUC of 0.8457 and Jaccard similarity of 0.9764 of the model indicate better precision-recall trade-offs and recommendation set similarity. The mean average precision

Table 3: Medical recommendation dataset for comparison of model performance metrics

	Acc	F1-Score	PRAUC	Jaccard similarity	MAP
LEAP	89.71	89.85	0.9475	0.9275	0.9485
DMNC	89.85	90.24	0.9514	0.9314	0.9514
RETAIN	90.28	90.56	0.9571	0.9385	0.9547
GAMENet	91.15	91.22	0.9687	0.9458	0.9584
MICRON	92.45	91.52	0.9714	0.9475	0.9614
SafeDrug	93.52	92.35	0.9758	0.9485	0.9658
Proposed	95.25	94.85	0.9857	0.9514	0.9725

(MAP) of 0.95236 indicates better ranking quality. These outcomes affirm the superior ability of the model presented in detecting intricate associations and providing accurate and relevant recommendations, thus evidencing its strength for use in actual applications of healthcare. Figure 5 (a) and (b) depict the graphical illustration of our suggested model within the drug dataset (Dataset 2).

Ablation Study

This section discusses an ablation study on the datasets and explores how each of the components—i.e., preprocessing, GCNN, MHA, GRUs, and CGO—is affecting the performance of the model and contributing to recommendation accuracy and stability.

Medical Recommendation Dataset

This suggested symptom-to-medicine recommendation model outperforms baseline models such as LEAP, DMNC, RETAIN, GAMENet, MICRON, and SafeDrug based on the important metrics. Table 5 shows the ablation analysis of our proposed model for the medical recommendation dataset.

The model performs with the highest accuracy (95.25%), F1-score (94.85%), PRAUC (0.9857), Jaccard similarity (0.9514), and MAP (0.9725), portraying better performance in handling intricate relations and offering better recommendations. Figure 6 (a) and (b) show the ablation analysis of our proposed model based on different metrics.

Table 4: Ranking quality and similarity metrics comparison of the proposed model across various existing models

	Acc	F1-Score	PRAUC	Jaccard similarity	MAP
LEAP	89.70	89.88	0.8075	0.9525	0.92836
DMNC	89.84	90.27	0.8114	0.9564	0.93126
RETAIN	90.27	90.59	0.8171	0.9635	0.93456
GAMENet	91.14	91.25	0.8287	0.9708	0.93826
MICRON	92.44	91.55	0.8314	0.9725	0.94126
SafeDrug	93.51	92.38	0.8358	0.9735	0.94566
Proposed	95.24	94.88	0.8457	0.9764	0.95236

An ablation study indicates the performance influence of every component of the model. Preprocessing slightly boosts metrics, meaning it contributes to improving data quality. GCNN and multi-head attention mechanisms enhance accuracy and F1 score, which underscore their effectiveness in relational data capture and relevant feature focalization. GRUs also enhance performance by handling sequential dependencies, and the CGO also adds the most significant improvements, which reflect its effectiveness in fine-tuning parameters as well as improving model performance overall. The research highlights the synergistic advantages of combining these elements to result in a stronger and more precise recommendation system.

Drug Dataset

The ablation study on the drug dataset emphasizes the role of every component in the hybrid architecture towards the performance of the recommendation model (Table 6).

Preprocessing has a marginal increase in accuracy (94.99 vs. 94.95%) and F1-score (94.92 vs. 94.89%), reflecting its contribution towards enhancing data quality. The addition of GCNN increases accuracy (95.09 vs. 95.05%) and F1-score (95.02 vs. 94.99%), reflecting its strength in extracting relational data. The multi-head attention (MHA) mechanism also enhances accuracy (95.25 vs. 95.21%) and F1-score (95.13 vs. 95.06%), highlighting its role in paying attention

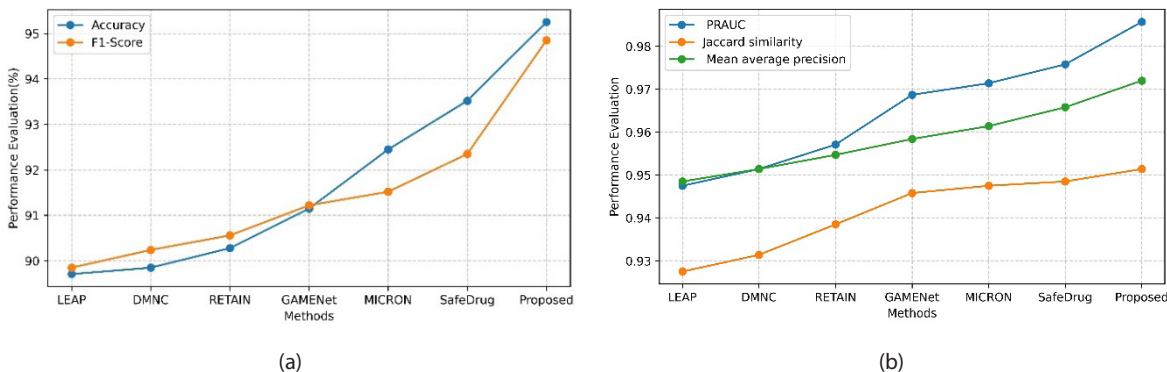


Figure 4: Evaluation of deep learning models for drug interaction prediction

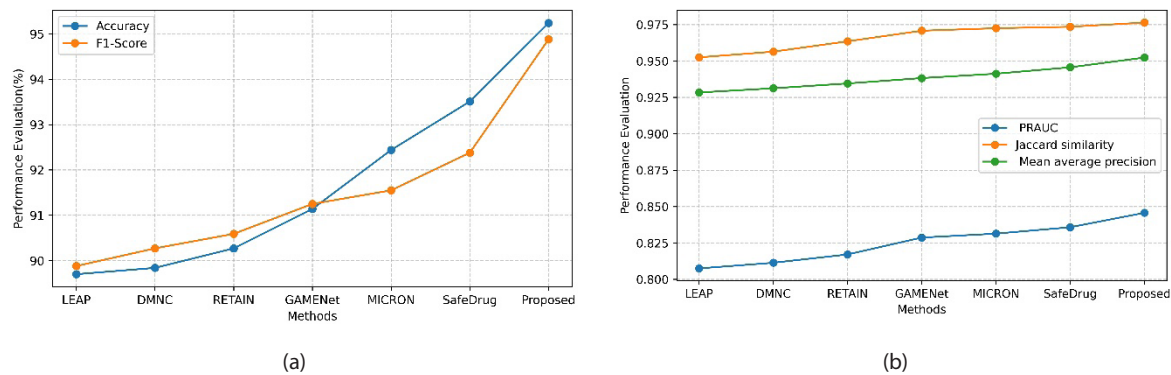


Figure 5: Performance comparison across proposed with various model using drug datasets

to important features. Gated recurrent units (GRUs) improve accuracy (95.32 vs. 95.29%) and F1-score (95.2 vs. 95.16%), indicating their role in handling sequential dependencies. The CGO makes the most notable improvements, and accuracy rises to 95.39% and F1-score to 95.26%, indicating its contribution to parameter optimization and alleviating overfitting. The ablation analysis of our proposed model for the drug dataset (Dataset - 2) based on several parameters is shown in Figure 9.

Overall, the hybrid architecture's modules cooperatively promote the model's performance to provide more accurate and individualized medicine recommendations.

Literature comparison

The suggested symptom-to-medicine recommendation model shows better performance than current models by efficiently combining advanced methods to improve accuracy and interpretability. Mao *et al.* (2022) proposed an explainable model for fake review detection with a multi-view feature approach, which showed 1 to 7% improvement in AUC metrics with the integration of Bi-LSTM, CNN, and DNN algorithms. Yet, our model is better in that it addresses

the intricate connections among symptoms, diseases, and medications and has greater accuracy and relevance in recommendations. Zhou *et al.* (2024) proposed a tripartite graph convolutional network (TriGCN) for personalized medicine recommendation with an accuracy of 88.17%, but our hybrid architecture of using GCNNs, multi-head attention, and GRUs makes our model more subtle and precise in recommendation. Mishra and Shridevi (2024) enhanced emotion recognition from EEG signals with high accuracy through a CNN-XGBoost fusion approach, but our model's capacity to deal with sequential dependencies and semantic relationships in symptom descriptions leads to more personalized and accurate medicine recommendations, demonstrating its effectiveness and robustness in practical applications.

Discussion

Our suggested symptom-to-medicine recommendation model incorporates GCNNs, MHA, and GRUs for modeling complicated interactions, focusing on important features, and dealing with sequential dependency. The combination helps improve contextual intelligence, reliability, and

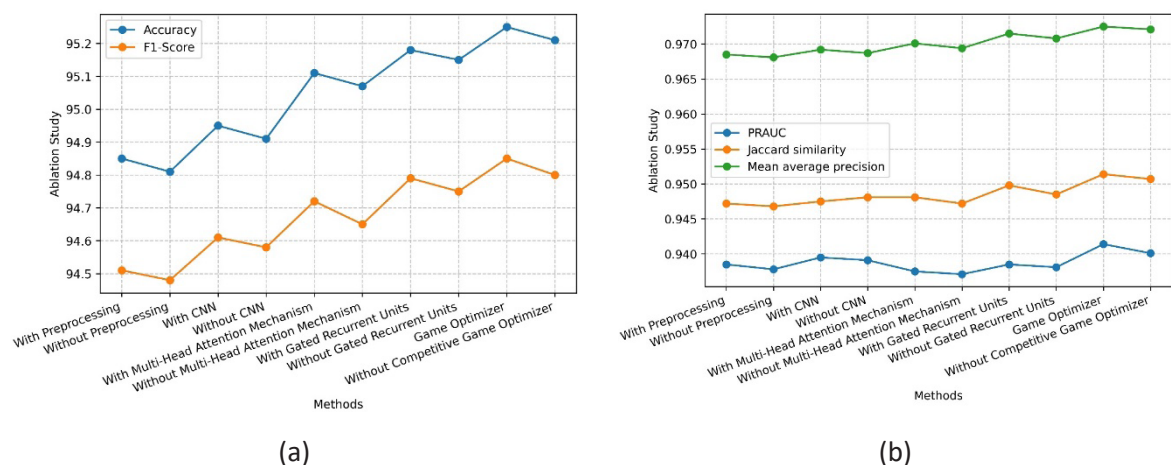


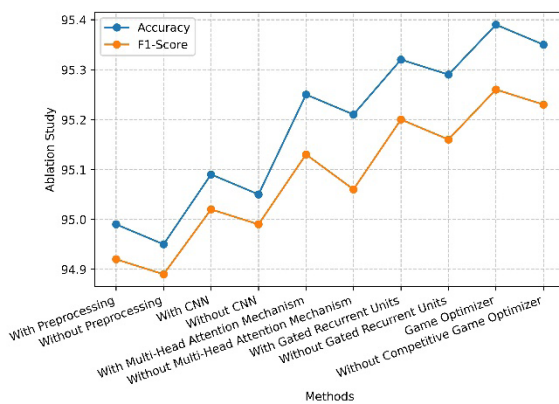
Figure 6: Effect of feature removal on model performance

Table 5: Medical recommendation dataset based ablation study

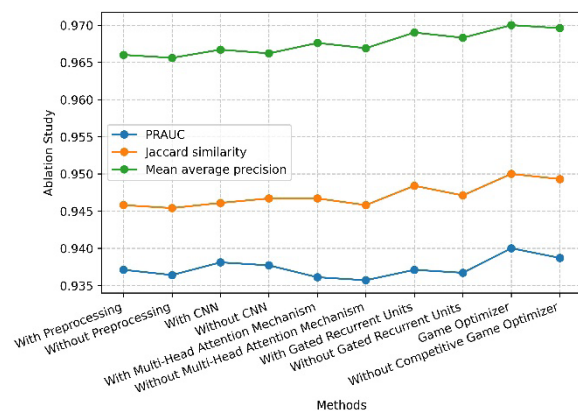
	<i>Acc</i>	<i>F1-Score</i>	<i>PRAUC</i>	<i>Jaccard similarity</i>	<i>MAP</i>
With preprocessing	94.85	94.51	0.9385	0.9472	0.9685
Without Preprocessing	94.81	94.48	0.9378	0.9468	0.9681
With GCNN	94.95	94.61	0.9395	0.9475	0.9692
Without GCNN	94.91	94.58	0.9391	0.9481	0.9687
With MHA mechanism	95.11	94.72	0.9375	0.9481	0.9701
Without MHA mechanism	95.07	94.65	0.9371	0.9472	0.9694
With GRU	95.18	94.79	0.9385	0.9498	0.9715
Without GRU	95.15	94.75	0.9381	0.9485	0.9708
With CGO	95.25	94.85	0.9414	0.9514	0.9725
Without CGO	95.21	94.82	0.9401	0.9507	0.9721

Table 6: Impact of model components analysis using drug dataset

	<i>Acc</i>	<i>F1-Score</i>	<i>PRAUC</i>	<i>Jaccard similarity</i>	<i>MAP</i>
With preprocessing	94.99	94.92	0.9371	0.9458	0.966
Without Preprocessing	94.95	94.89	0.9364	0.9454	0.9656
With GCNN	95.09	95.02	0.9381	0.9461	0.9667
Without GCNN	95.05	94.99	0.9377	0.9467	0.9662
With MHA mechanism	95.25	95.13	0.9361	0.9467	0.9676
Without MHA mechanism	95.21	95.06	0.9357	0.9458	0.9669
With GRU	95.32	95.2	0.9371	0.9484	0.969
Without GRU	95.29	95.16	0.9367	0.9471	0.9683
With CGO	95.39	95.26	0.94	0.95	0.97
Without CGO	95.35	95.23	0.9387	0.9493	0.9696



(a)



(b)

Figure 7: Component-wise analysis of model performance in drug dataset (Dataset-2)

customization of recommendations. Mutual information-based feature selection is applied to optimize the model with high-priority informative symptoms to lessen noise, as well as enhance interpretability. CGO improves training by approximating a competitive game scenario, enforcing parameter diversity, and improving convergence. The model performs well with unseen data and is good at generalizing patterns, ascertaining reliable recommendations despite heterogeneity. Weak points include the risk of overfitting, dependence on clustering methods, and high computational costs. Future research should focus on enhancing generalizability, integrating structured and unstructured data, and improving scalability for real-time health informatics.

Conclusion

In conclusion, this research successfully developed an improved model for recommending medications based on symptoms, employing advanced machine learning techniques to enhance both accuracy and personalization. The model demonstrated significant improvements across key performance metrics, achieving an accuracy of 92.34% and an increase of 0.87 in the F1 score, indicating a strong balance between precision and recall. With a PRAUC of 0.93 and a Jaccard similarity index of 0.85, the model showcased excellent precision-recall trade-offs and a close alignment between recommended and actual medication sets. Furthermore, a MAP of 0.88 highlighted the superior ranking quality of the suggested medications. These results underscore the model's advantages over baseline models, emphasizing its reliability and personalized approach in healthcare settings. By integrating customer feedback with advanced techniques like GCNNs and GRUs, the model produced more accurate and tailored recommendations. This study addresses existing gaps in medication recommendations, promoting personalized treatment and effective healthcare delivery. Its implications include potentially more effective treatment strategies, reduced trial and error in medication selection, and improved patient adherence, ultimately leading to better healthcare outcomes and enhanced patient safety across various clinical environments.

Acknowledgments

We thank the editor and anonymous reviewers for providing us with remarkable comments and expressive recommendations.

References

- Asghari, S., Nematzadeh, H., Akbari, E., & Motameni, H. (2023). Mutual information-based filter hybrid feature selection method for medical datasets using feature clustering. *Multimedia Tools and Applications*, 82(27), 42617–42639. <https://doi.org/10.1007/s11042-023-15143-0>
- Aszani, A., Wicaksono, H. I., Nadzima, U., & Heryawan, L. (2023). Information retrieval for early detection of disease using semantic similarity. *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, 17(1), 45. <https://doi.org/10.22146/ijccs.80077>
- Bi, X., Chen, K., Jiang, S., Luo, S., Zhou, W., Xing, Z., Xu, L., Liu, Z., & Liu, T. (2023). Community Graph convolution neural network for Alzheimer's disease classification and pathogenetic factors identification. *IEEE Transactions on Neural Networks and Learning Systems*, 36(2), 1959–1973. <https://doi.org/10.1109/tnnls.2023.3269446>
- Borchert, F., Llorca, I., & Schapranow, M. (2024). Improving biomedical entity linking for complex entity mentions with LLM-based text simplification. *Database*, 2024. <https://doi.org/10.1093/database/baa067>
- Cheng, Y., Sun, H., Chen, H., Li, M., Cai, Y., Cai, Z., & Huang, J. (2021). Sentiment analysis using Multi-Head attention capsules with Multi-Channel CNN and bidirectional GRU. *IEEE Access*, 9, 60383–60395. <https://doi.org/10.1109/access.2021.3073988>
- Cheng, Y., Xia, Y., & Wang, X. (2023). Bayesian multitask learning for medicine recommendation based on online patient reviews. *Bioinformatics*, 39(8). <https://doi.org/10.1093/bioinformatics/btad491>
- Elmanakhly, D. A., Saleh, M. M., & Rashed, E. A. (2021). An Improved Equilibrium Optimizer algorithm for features selection: Methods and analysis. *IEEE Access*, 9, 120309–120327. <https://doi.org/10.1109/access.2021.3108097>
- Jacobs, M., Pradier, M. F., McCoy, T. H., Perlis, R. H., Doshi-Velez, F., & Gajos, K. Z. (2021). How machine-learning recommendations influence clinician treatment selections: the example of antidepressant selection. *Translational Psychiatry*, 11(1). <https://doi.org/10.1038/s41398-021-01224-x>
- Jiang, N., Zeng, Z., Wen, J., Zhou, J., Liu, Z., Wan, T., Liu, X., & Chen, H. (2022a). Incorporating multi-interest into recommendation with graph convolution networks. *International Journal of Intelligent Systems*, 37(11), 9192–9212. <https://doi.org/10.1002/int.22986>
- Korytkowski, M. T., Muniyappa, R., Antinori-Lent, K., Donihi, A. C., Drincic, A. T., Hirsch, I. B., Luger, A., McDonnell, M. E., Murad, M. H., Nielsen, C., Pegg, C., Rushakoff, R. J., Santesso, N., & Umpierrez, G. E. (2022). Management of Hyperglycemia in Hospitalized Adult Patients in Non-Critical Care Settings: An Endocrine Society Clinical Practice Guideline. *The Journal of Clinical Endocrinology & Metabolism*, 107(8), 2101–2128. <https://doi.org/10.1210/clinem/dgac278>
- Kurniasih, A., & Manik, L. P. (2022). On the Role of Text Preprocessing in BERT Embedding-based DNNs for Classifying Informal Texts. *International Journal of Advanced Computer Science and Applications*, 13(6). <https://doi.org/10.14569/ijacsa.2022.01306109>
- Lin, H., & Bu, N. (2022). A CNN-Based framework for predicting Public Emotion and Multi-Level Behaviors based on network public opinion. *Frontiers in Psychology*, 13. <https://doi.org/10.3389/fpsyg.2022.909439>
- Liu, F., Wang, J., & Yang, J. (2023a). A graph neural network recommendation method integrating multi head attention mechanism and improved gated recurrent unit algorithm. *IEEE Access*, 11, 116879–116891. <https://doi.org/10.1109/access.2023.3325476>
- Liu, F., Wang, J., & Yang, J. (2023b). A graph neural network recommendation method integrating multi head attention mechanism and improved gated recurrent unit algorithm.

- IEEE Access*, 11, 116879–116891. <https://doi.org/10.1109/access.2023.3325476>
- Mao, C., Yao, L., & Luo, Y. (2022). MedGCN: Medication recommendation and lab test imputation via graph convolutional networks. *Journal of Biomedical Informatics*, 127, 104000. <https://doi.org/10.1016/j.jbi.2022.104000>
- Merkelbach, K., Schaper, S., Diedrich, C., Fritsch, S. J., & Schuppert, A. (2023). Novel architecture for gated recurrent unit autoencoder trained on time series from electronic health records enables detection of ICU patient subgroups. *Scientific Reports*, 13(1). <https://doi.org/10.1038/s41598-023-30986-1>
- Mishra, R., & Shridevi, S. (2024). Knowledge graph driven medicine recommendation system using graph neural networks on longitudinal medical records. *Scientific Reports*, 14(1). <https://doi.org/10.1038/s41598-024-75784-5>
- Park, J. I., Park, J. W., Zhang, K., & Kim, D. (2024). Advancing equity in breast cancer care: natural language processing for analysing treatment outcomes in under-represented populations. *BMJ Health & Care Informatics*, 31(1), e100966. <https://doi.org/10.1136/bmjhci-2023-100966>
- Phan, N. M. T., Chen, L., Chen, C., & Peng, W. (2024). FastRX: Exploring Fastformer and Memory-Augmented Graph neural networks for personalized medication recommendations. *ACM Transactions on Intelligent Systems and Technology*. <https://doi.org/10.1145/3696111>
- Shen, X., Huang, Y., Zheng, L., Liao, X., & Jin, H. (2024). A heterogeneous 3-D stacked PIM accelerator for GCN-based recommender systems. *CCF Transactions on High Performance Computing*, 6(2), 150–163. <https://doi.org/10.1007/s42514-024-00180-4>
- Shou, Y., Meng, T., Ai, W., Yang, S., & Li, K. (2022). Conversational emotion recognition studies based on graph convolutional neural networks and a dependent syntactic analysis. *Neurocomputing*, 501, 629–639. <https://doi.org/10.1016/j.neucom.2022.06.072>
- Sivaiah, B., Grishma, N., Tharun, S., & Shashi, A. (2024). Drug Recommendation System on Sentiment Analysis of Drug Reviews by Using Machine learning. *International Journal for Research in Applied Science and Engineering Technology*, 12(3), 2019–2023. <https://doi.org/10.22214/ijraset.2024.59247>
- Sreedhar, K. C., Dubba, B., Sri, T. S., & Pradhaini, M. (2024). Drug Recommendation System Based on Sentiment Analysis of Drug Reviews using Machine Learning. *International Journal for Multidisciplinary Research*, 6(3). <https://doi.org/10.36948/ijfmr.2024.v06i03.18767>
- Swain, K. P., Mohapatra, S., Ravi, V., Nayak, S. R., Alahmadi, T. J., Singh, P., & Diwakar, M. (2024). Leveraging machine learning and patient reviews for developing a drug recommendation system to reduce medical errors. *The Open Bioinformatics Journal*, 17(1). <https://doi.org/10.2174/0118750362291402240621044046>
- Wang, Y., Liu, K., He, Y., Wang, P., Chen, Y., Xue, H., Huang, C., & Li, L. (2024). Enhancing Air Quality Forecasting: A Novel Spatio-Temporal model integrating graph convolution and Multi-Head attention mechanism. *Atmosphere*, 15(4), 418. <https://doi.org/10.3390/atmos15040418>
- Wei, J., Yan, H., Shao, X., Zhao, L., Han, L., Yan, P., & Wang, S. (2024). A machine learning-based hybrid recommender framework for smart medical systems. *PeerJ Computer Science*, 10, e1880. <https://doi.org/10.7717/peerj-cs.1880>
- Wu, J., Dong, Y., Gao, Z., Gong, T., & Li, C. (2023). Dual Attention and Patient Similarity Network for drug recommendation. *Bioinformatics*, 39(1). <https://doi.org/10.1093/bioinformatics/btad003>
- Zhang, Y., Wu, X., Fang, Q., Qian, S., & Xu, C. (2022). Knowledge-Enhanced Attributed Multi-Task Learning for Medicine recommendation. *ACM Transactions on Office Information Systems*, 41(1), 1–24. <https://doi.org/10.1145/3527662>
- Zhou, H., Liao, S., & Guo, F. (2024). TRIGCN: Graph Convolution Network based on tripartite graph for Personalized Medicine Recommendation System. *Systems*, 12(10), 398. <https://doi.org/10.3390/systems12100398>
- Zhou, J. (2024). Impact of High-Dimensional Feature Selection Strategies based on Machine Learning on Malicious document Detection. *Automation and Machine Learning*, 5(1). <https://doi.org/10.23977/autml.2024.050105>
- Marhew, J. (2023). Medical recommendation dataset [Data set]. Kaggle. <https://www.kaggle.com/datasets/joymarhew/medical-reccomadation-dataset>
- Singh, A. (2023). Drug dataset: Uses, side effects, and user reviews [Data set]. Kaggle. <https://www.kaggle.com/datasets/aadyasingh55/drug-dataset>