



Effects of Online Medical Teams on Patients’ Choices for Doctor Selection: A Hybrid Deep Learning Framework

Yongbo Ni ^{1*} and Donghui Yang ¹

¹ Department of Management Science and Engineering, Southeast University, Nanjing, China.
*Corresponding author:nybgyc@163.com

Abstract. Amidst the overwhelming online medical service information, patients without professional medical knowledge often struggle to identify suitable doctors in online healthcare communities. The advent of online medical teams (OMTs) as a new source of publicly available information, offers patients novel avenues to understand the features of doctors. Therefore, in this study, we aim to explore the effect of OMTs information on patients’ choices for doctor selection, thereby better assisting patients in selecting a suitable doctor. We first integrate the OMTs information with the online reviews, disease descriptions, and doctor profiles to build the multi-source and multi-type medical data inputs. Based on these inputs, we develop a hybrid deep learning framework to uncover the effects of various factors on predicting patients’ choices for doctor selection. The results indicate that OMTs information can significantly enhance the probability of predicting patients’ choices for doctor selection. Surprisingly, the effect of OMTs information on patients’ preference features is greater than that of patients’ disease features in the process of doctor selection.

Keywords: Online Healthcare Community, Online Medical Team, Doctor Selection, Deep Learning.

NOMENCLATURE

Parameters	Descriptions
PDF	Patients’ disease features
PPF	Patients’ preference features
DDF	Disease-type features a doctor specializes in
DPF	Doctor attribute features related to PPF
EDDF	Enhanced patients’ awareness of the target doctor’s DDF
EDPF	Enhanced patients’ awareness of the target doctor’s DPF

1 Introduction

The swift advancement of information technology has led a significant number of patients to turn to online healthcare communities (OHCs) for health services [1]. However, the plethora of health service options can be overwhelming for patients lacking professional medical knowledge to select a suitable doctor [2]. To assist patients in

efficiently selecting online doctors, current research mainly relies on publicly available information on OHCs, such as online reviews and consultation records to enhance the awareness of patients' features[3]. These studies primarily investigate the effect of publicly available information from two aspects: patients' disease features (PDF) and patients' preference features (PPF) [4,5]. From the perspective of PDF, research has utilized data from patients' online disease descriptions and doctors' consultation records, suggesting that a higher degree of similarity between the PDF and the disease-type features a doctor specializes in (DDF) increases the likelihood of the patient selecting that doctor [6]. From the perspective of PPF, studies have leveraged patients' online reviews and doctor attribute features (DPF) such as professional titles to assess their effects on patients' choices for doctor selection [7]. The majority of existing research has focused on the effect of a single feature on doctor selection, however, patients often consider both PDF and PPF when seeking online health services.

In addition, the rapid expansion of OHCs has witnessed a transition in their service paradigm from a simple one-to-one mode to a multi-to-one online medical teams (OMTs) mode [8]. Doctors from diverse departments can form OMTs, surmounting geographical constraints to engage in online medical collaborations[9,10]. Within these OMTs, doctors can synthesize their collective experiences, foster enhanced communication, and deepen their understanding of patients' conditions[11]. Given the significant effectiveness of OMTs, existing research has carried out extensive research on their performance, organizational structure [12,13]. However, they have predominantly centered on the team as a whole, with inefficient attention given to the effect of the OMTs information on patients' choices for individual doctors. The OMTs information serves as an additional source of online publicly available information, providing patients with novel avenues to assess the competencies of individual doctors [14]. On one hand, patients can gain insight into the professional ability of a target doctor by examining the disease issues addressed by the OMT that the doctor is in, thereby enhancing their awareness of the DDF of the target doctor (EDDF). On the other hand, the reputation and feedback of other team members can indirectly shape patients' favorability towards the target doctor, enhancing their awareness of the target doctor's DPF (EDPF).

In response to the above discussion, this study aims to investigate the effect of OMTs information on patients' choices for doctor selection, focusing on addressing the following three research questions:

RQ1. Without considering the OMTs information, how do the PDF and PPF impact patients' choices for doctor selection?

RQ2. How does the supplementary OMTs information impact patients' choices for doctor selection?

RQ3. How do the EDDF and EDPF impact patients' choices for doctor selection, respectively?

As shown in Fig.1, to address the aforementioned research questions, this study compiles multi-source and multi-type data from Haodf-China's largest OHCs, including patients' online reviews, disease descriptions, doctor profiles, and OMTs information. By constructing a hybrid deep learning framework, we investigate the effects of the above factors on patients' choices for doctors selection.

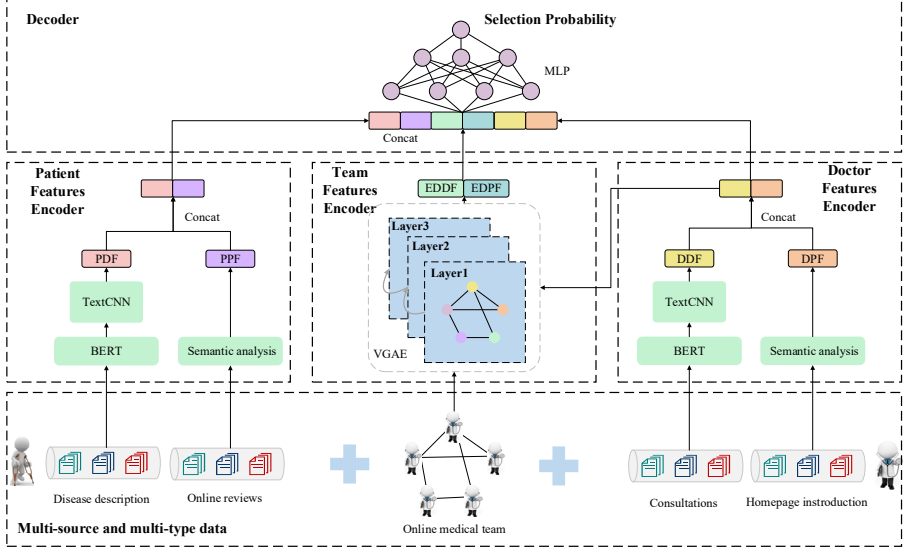


Fig. 1. Illustration of the main process for the hybrid deep learning framework.

2 Methodology

2.1 Representation of Disease Features

This study constructs a BERT-TextCNN model (see Fig. 2) to analyze the PDF of patients. Firstly, this study uses the BERT model to learn patients' disease descriptions and represent them with a word embedding matrix ($\mathbf{M}_p$). Next, convolution operations are performed on $\mathbf{M}_p$. Each convolution kernel k starts from position l and slides according to its size s to extract text features f_l^k .

$$f_l^k = \sigma(\mathbf{W}_k \mathbf{M}_p(l:l+s-1) + b_k) \quad (1)$$

Then, a maximum pooling operation is performed on f_l^k to extract the key feature information: $e_k = \max\{f_1^k, f_2^k, \dots, f_{n-s+1}^k\}$, and the acquired information is fused $\mathbf{e}_p = [e_1, e_2, \dots, e_{n_l}]$. Finally, a fully connected layer is constructed to analyze the PDF.

$$\text{PDF} = \sigma(\mathbf{W}_p \mathbf{e}_p + b_p) \quad (2)$$

Where σ uses the Softmax function to determine the probability of a patient belonging to different diseases. Based on the predicted disease probabilities, a cross-entropy loss function is used to optimize the model parameters.

$$loss = \sum_{p \in P} L_p \log \text{PDF} + (1 - L_p) \log(1 - \text{PDF}) \quad (3)$$

Where L_p represents the true probability distribution. Since the text form of patients' disease descriptions is quite similar to doctors' consultation records, this study also utilizes BERT-TextCNN model to analyze doctors' consultation records to obtain the disease-type features a doctor specializes in (DDF).

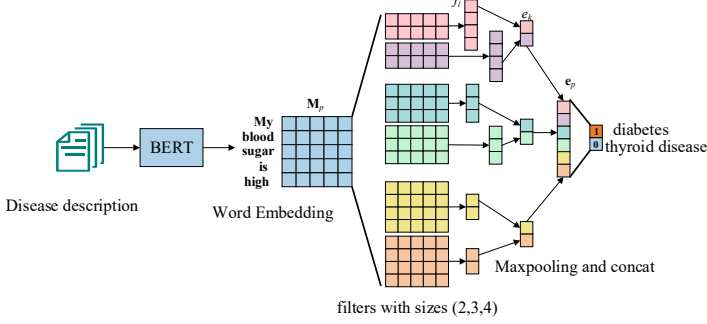


Fig. 2. Process for PDF learning.

2.2 Representation of Preference Features

As shown in Fig.3, to analyze the patients' preference features (PPF), this study employs semantic analysis methods to mine and analyze patients' online reviews on OHCs. This study first preprocesses the online reviews by removing punctuation and irrelevant information. Following this, the Jieba segmentation tool is applied to perform word segmentation on the text and to discard irrelevant stop words. Subsequently, the Term Frequency-Inverse Document Frequency (TF-IDF) [15] technique is employed to evaluate the significance of each word. The resulting *tf-idf* score is calculated as the product of TF and IDF.

$$tf-idf = \frac{n_{t,d}}{N_d} \times \log\left(\frac{N_D}{n_t + 1}\right) \quad (4)$$

Based on the *tf-idf* values, this study selects high-frequency keywords to learn the PPF.

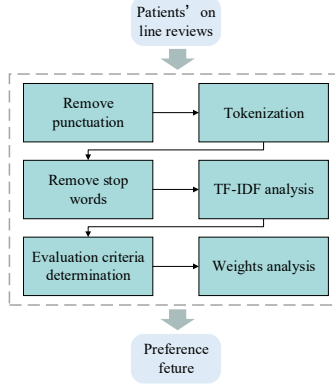


Fig. 3. Process for PPF learning.

In this study, the keywords extracted are clustered based on their semantic similarity to establish evaluation criteria[16]. The similarity between any two keywords is represented as S_{ij} , and the predefined similarity threshold is denoted as $\mathcal{S}$. If $S_{ij} \geq \mathcal{S}$, the two keywords are considered to have high similarity and can be categorized under the same evaluation criterion. The doctors' attribute features (DPF) aligned with patient preferences are subsequently derived from the doctors' profile information, in accordance with the established preference criteria.

$tf-idf$ reflects the importance of each keyword, while the accumulation of $tf-idf$ values of keywords contained within each criterion in patient reviews reflects the degree of patient preference for that evaluation criterion. Let $tf-idf_{pi}$ denote the cumulative $tf-idf$ value of the p -th patient review under the i -th evaluation criterion; then, the preference degree of p -th patient for the i -th criterion can be expressed as:

$$w_{pi} = \frac{tf-idf_{pi}}{\sum_{i=1}^I tf-idf_{pi}} \quad (5)$$

2.3 Representation of Enhanced Doctor Feature Through OMTs Information

The process of learning enhanced feature representations for individual doctors within OMTs is illustrated in Fig. 4. The input to the Variational Graph Auto-Encoder (VGAE) [17] consists of two parts: the network structure and the node feature vector matrix. The network structure represents the collaborative relationships among doctors, while the node features are represented by DDF and DPF ($\mathbf{X}=[\text{DDF}, \text{DPF}]$). The VGAE employs a graph neural network as the encoder to learn the corresponding latent variables $\mathbf{Z}$.

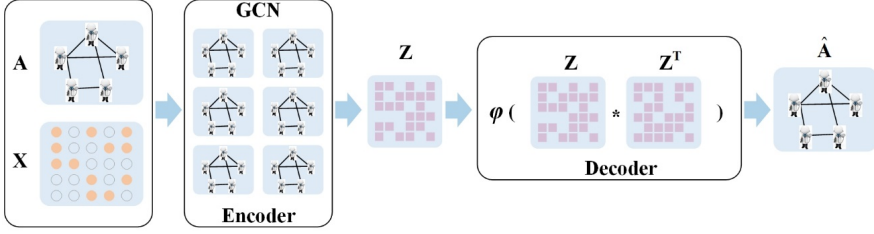


Fig. 4. Process for EDDF and EDPF learning.

$$q(\mathbf{Z} | \mathbf{X}, \mathbf{A}) = \prod_{i=1}^N q(z_i | \mathbf{X}, \mathbf{A}), \text{ with } q(z_i | \mathbf{X}, \mathbf{A}) = N(z_i | \mu_i, \text{diag}(\sigma_i^2)) \quad (6)$$

Where the model assumes the latent variable $\mathbf{Z}$ follows a standard Gaussian distribution, with μ and σ representing the mean and standard deviation, respectively. The μ and σ are learned through two layers of Graph Convolutional Network-GCN($\mathbf{X}, \mathbf{A}$). $\mu = \text{GCN}_{\mu}(\mathbf{X}, \mathbf{A})$ and $\log \sigma = \text{GCN}_{\sigma}(\mathbf{X}, \mathbf{A})$. The GCN($\mathbf{X}, \mathbf{A}$) is defined as

$$\text{GCN}(\mathbf{X}, \mathbf{A}) = \tilde{\mathbf{A}} \text{ReLU}(\tilde{\mathbf{A}} \mathbf{X} \mathbf{W}_0) \mathbf{W}_1 \quad (7)$$

Based on the obtained μ and σ , the latent variable z can be represented, and the graph structure can further be reconstructed through the decoder.

$$z = \mu + \varepsilon \odot \sigma, \varepsilon \in N(0, 1) \quad (8)$$

$$p(\mathbf{A} | \mathbf{Z}) = \prod_{i=1}^N \prod_{j=1}^N p(A_{ij} | z_i, z_j), \text{ with } p(A_{ij} = 1 | z_i, z_j) = \sigma(z_i^T z_j) \quad (9)$$

The loss function for the entire model can be defined as:

$$\text{loss} = \mathbb{E}_{q(\mathbf{Z} | \mathbf{X}, \mathbf{A})} [\log p(\mathbf{A} | \mathbf{Z}) - \text{KL}[q(\mathbf{Z} | \mathbf{X}, \mathbf{A}) \| p(\mathbf{Z})]] \quad (10)$$

Where $\text{KL}[q(\bullet) \| p(\bullet)]$ denotes the Kullback-Leibler divergence between distributions $p(\bullet)$ and $q(\bullet)$. Based on the aforementioned model, this study selects the optimized latent variable $\mathbf{Z}$ to represent the enhanced features of doctors within the team.

3 Experiments

3.1 Dataset

We use the online consultation records in the endocrinology department in Haodaifu platform as an example to implement our experiments. The platform will recommend the top 100 relevant doctors to patients. However, because the platform only provides patients' online reviews and consultation information from the past 90 days, we can only obtain valid information for 66 doctors. Additionally, based on the doctors' team affiliations, we acquired information on 32 doctors not listed in the platform's recom-

mendation, totaling 98 valid doctors. From the information of these 98 doctors, we obtained online reviews and consultation records from 2021 patients in the past 90 days. Finally, we select 70% of the data as training set and use 30% of the data as test set.

3.2 Data Preprocessing

Table 1. Vector representations of some keywords

words	1	2	...	51	52	...	99	100
hyperthyroidism	-0.437	0.591	...	-0.106	-0.265	...	0.325	0.122
diabetes	-0.156	0.500	...	0.006	-0.060	...	0.132	0.556
pregnancy	-0.984	1.321	...	0.290	-0.521	...	0.227	-0.104
obesity	-0.316	0.339	...	0.069	-0.076	...	0.016	-0.062
adrenal	-0.326	0.357	...	0.214	0.033	...	0.045	-0.053

Table 2. Eight major diseases in the endocrinology department.

No.	Diseases	Descriptions
F1	thyroid diseases	hyperthyroidism, hypothyroidism, thyroiditis
F2	diabetes	type II diabetes, type II diabetes, gestational diabetes
F3	obesity	obesity, fat
F4	pituitary diseases	pituitary microadenoma, high pituitary prolactin
F5	hematological diseases	hyperlipidemia, hyperkalemia, hyperphosphatemia
F6	adrenal diseases	suprarenoma, low norepinephrine
F7	bone diseases	osteoporosis, low bone density
F8	reproductive	menoxenia, polycystic ovarian syndrome

Preprocessing of PDF and DDF. First, we apply a 100-dimensional vector to represent the relevant keyword embeddings learned by the BERT model to encode the critical information from patients’ disease descriptions and doctors’ consultation records. Table 1 lists the vector representations of some common words. Based on the learned word embeddings, we utilize TextCNN to characterize PDF of patients and DDF of doctors. The diseases encountered during endocrinology are systematically classified into eight primary categories (refer to Table 2). Drawing from these disease representations, the PDF of patients is delineated in Table 3. Specifically, P1 and P2 exhibit a higher risk of thyroid diseases, P3 is at a higher risk of hematological diseases. Similarly, representations of DDF for doctors can be obtained.

Table 3. Representations of PDF for patients.

Patients	F1	F2	F3	F4	F5	F6	F7	F8
P1	0.9910	0.0002	0.0004	0.0008	0.0016	0.0022	0.0010	0.0028
P2	0.9575	0.0156	0.0006	0.0009	0.0136	0.0087	0.0001	0.0030

P3	0.0378	0.1501	0.0254	0.0231	0.6189	0.0833	0.0388	0.0226
P4	0.1146	0.1127	0.5256	0.0591	0.0671	0.0430	0.0283	0.0496
P5	0.0163	0.8818	0.0068	0.0043	0.0737	0.0088	0.0057	0.0026

Preprocessing of PPF and DPF . Patients’ online reviews are divided into a series of keywords, then, the TF-IDF technology is applied to analyze the importance of each keyword. Next, these keywords are clustered based on similarity to form evaluation criteria. We set S as 0.3, the similarity between “professional” and “experienced” is 0.67 which is higher than S (see Fig.5). Thus, these two words can be clustered into the same evaluation criteria. Based on the similarity, these keywords can be categorized into four main criteria: Leechcraft, Attitude, Online communication skill, and Recovery (see Table 4). Then, as shown in Table 5, we can see the following insights into patient preferences: P1 assigns the greatest importance to the doctors’ online communication skill. In contrast, P2 and P4 prioritize the doctors’ service attitude more significantly.

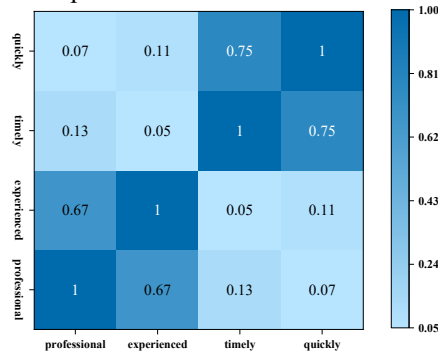


Fig. 5. Similarity between each keyword.

Table 4. Illustrations of evaluation criteria.

Criteria	High-frequency words
Leechcraft	professional, skillful, trustworthy, exquisite, experienced
Attitude	amiable, kindhearted, careful, patient, responsible, serious, earnest, amiableness, perfunctory
Online communication skill	meticulous, detailed, timely, clear, quickly, unambiguous, perspicuous
Recovery	reduce, recovery, cure, eliminate

Table 5. Patients’ preferences on the evaluation criteria.

Patients	Leechcraft	Recovery	Attitude	Online communication skill
P1	0.139	0.136	0.291	0.434
P2	0.195	0.192	0.409	0.204
P3	0.665	0.164	0.00	0.171
P4	0.241	0.000	0.505	0.254
P5	0.162	0.160	0.340	0.338

In terms of leechcraft, doctor's professional title can be used as a proxy for indicating the doctor's professional competence. The professional titles include physician, attending physician, associate chief physician, and chief physician, with corresponding scores of 0, 0.33, 0.67, and 1. For evaluation criteria of recovery and attitude, we utilizes the curative effectiveness and attitude from the doctor's homepage information, with values scaled between 0 and 1. In terms of online communication skill, we chose the average response time as the score. The average response time includes: slow, slightly slow, normal, fast, and very fast, with standardized scores of 0, 0.25, 0.5, 0.75, and 1. Utilizing these scoring guidelines, we derive the scores of each doctor Table 6.

Table 6. Doctors' scores on the evaluation criteria.

Doctors	Leechcraft	Recovery	Attitude	Online communication skill
D1	0.67	0.99	1	1
D2	0.67	1	1	0.75
D3	1	1	0.99	1
D4	1	0.96	0.98	0.5
D5	0.33	1	1	1

Preprocessing of EDDF and EDPF. Based on the information about the OMTs, we construct a collaboration relationship network among the 98 doctors. Additionally, we combine the data of DDF and DPF to create an initial representation of the doctor features for team feature learning. Leveraging the relationships within the OMTs and the initial features of doctors, we employed the Variational Graph Auto-Encoder (VGAE) method to learn patients' cognition regarding the doctor features (EDDF and EDPF). The initial features and enhanced features of doctors are shown in Table 7.

Table 7. Initial features and enhanced features of doctors.

Features	ID	DDF(10 ^{^-2})/ EDDF(10 ^{^-3})								DPF(10 ^{^-2})/ EDPF(10 ^{^-3})			
Initial	3	71.42	8.54	0.83	1.00	8.73	5.95	0.81	2.71	0.67	0.99	1	1
	5	66.88	12.41	1.23	1.31	8.84	5.84	0.88	2.60	1	0.96	0.93	1
	8	99.06	0.06	0.04	0.09	0.22	0.20	0.09	0.25	1	0.98	0.98	0.75
Learned	3	9.88	3.22	-7.96	-5.26	2.55	0.35	-8.21	-5.37	1.81	2.50	2.29	2.67
	5	9.75	0.32	-5.22	-5.32	0.92	0.79	-9.06	-3.83	1.47	0.21	2.67	1.95
	8	19.24	-9.68	-8.30	-2.88	-0.61	-8.78	-3.58	2.09	-0.71	1.79	-2.27	-2.30

3.3 Results

Metrics. We evaluate the performance of the model using three metrics: HR@K, NDCG@K, and MRR [2]. The definitions of these three metrics are as follows:

$$HR@K = \frac{\sum_{p \in P} Hit_p @K}{|P|} \quad (11)$$

where P represents all patients in the test dataset; $Hit_p @K$ indicates whether the doctors consulted by patient P appear in the top-K predicted doctor list.

$$NDCG@K = \frac{NDCG_p@K}{|P|} \quad (12)$$

$$NDCG_p@K = \frac{DCG_p@K}{IDCG_p@K} \quad (13)$$

$$DCG_p@K = \sum_{i=1}^K \frac{2^{rel_i} - 1}{\log_2(i+1)} \quad (14)$$

Where $rel_i \in \{0,1\}$ denotes whether the doctor at the i -th position in the prediction list is the one actually chosen by the patient. $DCG_p@K$ represents the discounted cumulative gain of all predicted doctors in the list, while $IDCG_p@K$ denotes the maximum possible value of $DCG_p@K$ among all candidates.

$$MRR = \frac{1}{|P|} \sum_{p \in P} \frac{1}{rank_p} \quad (15)$$

where $rank_p$ indicates the position of the doctor actually chosen by the patient.

Results analysis and comparison. To solve the problems raised in Section 1, we established six experiments (See Table 8) to analyze the effects of different features.

Table 8. Results of the six experiments.

Models		M1	M2	M3	M4	M5	M6
Fea- tures	PDF	▲		▲	▲	▲	▲
	PPF		▲	▲	▲	▲	▲
	DDF	▲		▲	▲	▲	▲
	DPF		▲	▲	▲	▲	▲
	EDDF				▲	▲	
	EPDF				▲		▲
Re- sults	HR@5	0.2526	0.2420	0.3287	0.4079	0.3295	0.3866
	HR@10	0.4637	0.4602	0.6194	0.6259	0.6198	0.6244
	NDGC@5	0.1244	0.1273	0.1729	0.2843	0.2036	0.2257
	NDGC@10	0.2284	0.2073	0.2421	0.3850	0.2951	0.3403
	MRR	0.1693	0.1443	0.1945	0.2928	0.2017	0.2267

Answer to RQ1: Experiments M1 and M2 explore patients' choices for doctor selection from the perspective of disease and preference, respectively. As shown in Table 8, when considering only disease features (PDF and DDF), the values for HR@5, HR@10, NDGC@10, and MRR are all higher than those that only consider preference features (PPF and DPF). The NDGC@5 value for M1 is slightly lower than that for M2. Overall, without considering the effect of OMTs information on patient decision-making, the effect of disease features on patients' choices for doctor selection is greater than that of preference features.

Answer to RQ2: In Experiment M4, we expanded upon the parameters of Experiment M3 by adding supplementary information (EDDF and EDPF) on OMTs. A comparison between M4 and M3 shows that M4 significantly outperforms M3 in the $HR@K$, $NDCG@K$, and MRR metrics. Specifically, M4 exhibits a 7.92% increase in $HR@5$, an 11.14% increase in $NDCG@5$, and a 9.83% increase in MRR compared to M3. The information for EDDF and EDPF is learned from publicly available information on OMTs. This indicates that publicly available information on OMTs can enhance patients' understanding of doctor features. By incorporating these supplementary data, the success rate in predicting patients' choices for doctor selection has been effectively increased.

Answer to RQ3: To further explore the impact of OMTs information, this study conducts separate analyses to uncover the effects of EDDF and EDPF on patients' choices for doctor selection. As shown in Table 8, Experiment M5 demonstrates that after incorporating additional EDDF information, the model shows improvements in $HR@K$, $NDCG@K$, and MRR when compared to M3. Similarly, Experiment M6 also shows higher values in $HR@K$, $NDCG@K$, and MRR compared to M3 after adding the EDPF information. This indicates that, regardless of whether it is from the perspective of patients' disease features or their preference features, publicly available information on OMTs enhances patients' understanding of doctor features, thereby improving the accuracy of predicting patients' choices for doctor selection.

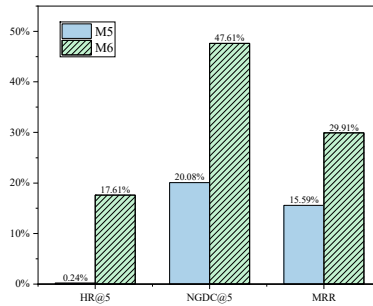


Fig. 6. Comparison among M3, M5, and M6.

As illustrated in Fig. 6, using M3 as a baseline, M5 shows improvements of 0.24%, 20.08%, and 15.59% in $HR@5$, $NDCG@5$, and MRR , respectively, while M6 demonstrates increases of 17.61%, 47.61%, and 29.91% in the same metrics. The superior predictive performance of Experiment M6 over M5 indicates that EDPF exerts a more substantial effect on patients' choices for doctor selection than EDDF. This experimental result contradicts the finding in RQ1, where PDF is found to have a greater impact on patients' choices for doctor selection than PPF. The underlying reason for this phenomenon may be that patients initially focus more on selecting doctors who align with their disease features. But once they become informed about doctors' specialties, they place greater emphasis on whether the doctors' various attributes align with their personal preferences. In other words, with existing information such as PDF, PPF, PPF, and DPF, the additional information on OMTs better caters to patients' diverse demand for doctor attributes.

4 Dicsussions and Implications

4.1 Robustness check

To ascertain the robustness of our findings, this study removed data for one of the eight diseases in the dataset and re-conducted the experiments. As demonstrated in Table 9, the experimental results observed after the partial removal of disease data are largely consistent with those obtained in Section 3.2. Notably, the effect of PDF on patients’ choices for doctor selection remained more significant than that of PPF, irrespective of the OMTs information. Similarly, incorporating OMTs information significantly enhanced the accuracy of predicting the probabilities of patients’ choices for doctor selection, with the effect of EDDF being more pronounced than that of EDPF.

Table 9. Results of the robustness check.

Metrics	M1	M2	M3	M4	M5	M6
HR@5	0.2445	0.2328	0.2956	0.3832	0.3029	0.3723
HR@10	0.5912	0.5622	0.6241	0.6715	0.6241	0.6277
NDGC@5	0.1451	0.1360	0.1588	0.1837	0.1663	0.1799
NDGC@10	0.1509	0.1499	0.1522	0.1872	0.1633	0.1718
MRR	0.2242	0.2163	0.2467	0.2788	0.2573	0.2605

4.2 Comparision to Traditional Methods

In predicting the probability of patients’ choices for doctor selection, many traditional methods primarily analyze the similarity of features between doctors and patients. This study, however, integrates multi-source information such as PDF, PPF, DDF, DPF, EDDF, and EDPF to construct a deep learning model for predicting the probability of patients’ choices for doctor selection. To validate the effectiveness of the method, a comparison is made between this method and traditional methods. We compare our method with concept-based auto-assignment (CBAA) method [6], content-based (CB) method [18], graph computing and LDA topic model (GC-LDA) [19].

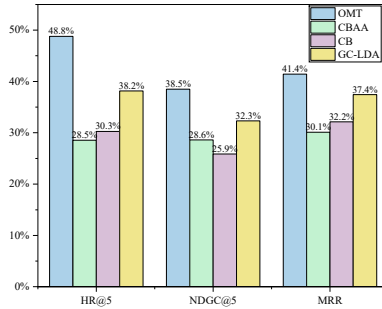


Fig. 7. Comparison with OMT and the other three methods.

Fig. 7 presents a comparison of results between the method described in this paper (OMT), and the aforementioned three methods. Among these, the CBAA and CB methods yield relatively lower values for HR@5, NGDC@5, and MRR metrics. The GC-LDA method performs slightly better than both CBAA and CB in these three metrics. In contrast, the OMT method significantly outperforms the traditional methods in all three metrics, with an average improvement of 16.4%, 9.5%, and 8.2% in HR@5, NGDC@5, and MRR. From these results, it can be concluded that the deep learning-based approach, which integrates additional OMTs information effectively enhances the prediction of the probability of patients' choices for doctor selection.

4.3 Implications

The findings of this study have guiding implications for the operation and management of actual online health platforms. Firstly, based on the conclusions of RQ1, it is necessary for doctors to enhance their professional capabilities and clearly demonstrate their areas of expertise in consultation records, allowing patients to understand the types of diseases they specialize in and their treatment capabilities. After meeting the aforementioned conditions, according to the conclusions of RQ2 and RQ3, doctors should also pay attention to meeting patients' preference features and strengthening their reputation and credibility. Furthermore, under permissible conditions, they should enhance team cooperation with other doctors to enhance patient cognition and understanding of them. Secondly, for online health platforms, in order to assist patients in selecting online doctors, it is essential to consider both patients' disease features and preference features while also incorporating OMTs information.

5 Conclusion

In this study, we develop a hybrid deep learning framework to uncover the effects of OMTs information on patients' choices for doctor selection. The results shows OMTs information is beneficial for enhancing the prediction of patients' choices for doctor selection. Specially, we also find the effect of OMTs information on patients' preference features is greater than that of patients' disease features in doctor selection. The features of OMTs are varied. The diversity of OMTs (e.g., gender, professional titles) and organizational patterns (e.g., led by experts, led by general doctors) also have an impact on patient selection of doctors. In future work, we plan to interagetd these OMTs' features to investigate their effects on patients' choices for doctor selection.

References

1. X. Chen, H. Wang, X. Li, Doctor recommendation under probabilistic linguistic environment considering patient's risk preference, *Annals of Operations Research* (2022). <https://doi.org/10.1007/s10479-022-04843-9>.

2. H. Yuan, W. Deng, Doctor recommendation on healthcare consultation platforms: an integrated framework of knowledge graph and deep learning, *Internet Research* 32 (2) (2022) 454-476. <https://doi.org/10.1108/intr-07-2020-0379>.
3. L. Jiang, J. Hou, X. Ma, P.A. Pavlou, Punished for success? A natural experiment of displaying clinical hospital quality on review platforms, *Information Systems Research* (2024). <https://doi.org/10.1287/isre.2021.0630>.
4. T. Zhou, Y. Wang, L. Yan, Y. Tan, Spoiled for choice? Personalized recommendation for healthcare decisions: A multiarmed bandit approach, *Information Systems Research* 34 (4) (2023) 1493-1512. <https://doi.org/10.1287/isre.2022.1191>.
5. Z. Huang, L. Yuan, Enhancing learning and exploratory search with concept semantics in online healthcare knowledge management systems: An interactive knowledge visualization approach, *Expert Systems with Applications* 237 (2024) 121558. <https://doi.org/10.1016/j.eswa.2023.121558>.
6. H. Naderi, B. Kiani, S. Madani, K. Etmnani, Concept based auto-assignment of healthcare questions to domain experts in online Q&A communities, *International Journal of Medical Informatics* 137 (2020) 104108. <https://doi.org/10.1016/j.ijmedinf.2020.104108>.
7. H. Yang, X. Guo, T. Wu, Exploring the influence of the online physician service delivery process on patient satisfaction, *Decision Support Systems* 78 (2015) 113-121. <https://doi.org/10.1016/j.dss.2015.05.006>.
8. X. Ding, X. Zhang, W.H. Wang, X. You, How online healthcare team evolve into organization: A social network analysis, *Digital Health* 10 (2024) 20552076241286634. <https://doi.org/10.1177/20552076241286634>.
9. H. Liu, S.C. Perera, J.-J. Wang, J.M. Leonhardt, Physician engagement in online medical teams: A multilevel investigation, *Journal of Business Research* 157 (2023) 113588. <https://doi.org/10.1016/j.jbusres.2022.113588>.
10. Z. Deng, G. Fan, Z. Deng, B. Wang, Why doctors participate in teams of online health communities? A social identity and brand resource perspective, *Information Systems Frontiers* (2023). <https://doi.org/10.1007/s10796-023-10434-1>.
11. H. Yang, Z. Yan, L. Jia, H. Liang, The impact of team diversity on physician teams' performance in online health communities, *Information Processing & Management* 58 (1) (2021) 102421. <https://doi.org/10.1016/j.ipm.2020.102421>.
12. J.-J. Wang, H. Liu, J. Ye, The impact of individual and team professional capital on physicians' performance in online health-care communities: a cross-level investigation, *Internet Research* 33 (1) (2023) 152-177. <https://doi.org/10.1108/intr-08-2021-0544>.
13. J. Li, H. Wu, Z. Deng, N. Lu, R. Evans, C. Xia, How professional capital and team heterogeneity affect the demands of online team-based medical service, *Bmc Medical Informatics and Decision Making* 19 (2019) 119. <https://doi.org/10.1186/s12911-019-0831-y>.
14. H. Liu, S.C. Perera, J.-J. Wang, Does the physicians' medical team joining behavior affect their performance on an online healthcare platform? Evidence from two quasi-experiments, *Transportation Research Part E-Logistics and Transportation Review* 171 (2023) 103040. <https://doi.org/10.1016/j.tre.2023.103040>.
15. H.C. Wu, R.W.P. Luk, K.F. Wong, K.L. Kwok, Interpreting TF-IDF term weights as making relevance decisions, *Acm Transactions on Information Systems* 26 (3) (2008) 13. <https://doi.org/10.1145/1361684.1361686>.
16. J. Chen, X. Li, Doctors ranking through heterogeneous information: The new score functions considering patients' emotional intensity, *Expert Systems with Applications* 219 (2023) 119620. <https://doi.org/10.1016/j.eswa.2023.119620>.
17. T.N. Kipf, M. Welling, 2016. Variational graph auto-Encoders. Publishing, p. arXiv:1611.07308.

18. J. Liu, C. Li, Y. Huang, J. Han, An intelligent medical guidance and recommendation model driven by patient-physician communication data, *Frontiers in Public Health* 11 (2023) 1098206. <https://doi.org/10.3389/fpubh.2023.1098206>.
19. Q. Meng, H. Xiong, A doctor recommendation based on graph computing and LDA topic model, *International Journal of Computational Intelligence Systems* 14 (1) (2021) 808-817. <https://doi.org/10.2991/ijcis.210205.002>.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (<http://creativecommons.org/licenses/by-nc/4.0/>), which permits any noncommercial use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

