

CARES: A Hybrid caregivers recommendation system using deep learning and knowledge graphs

Qiaoyun Zhang^a, Sze-Han Wang^b, Chung-Chih Lin^c, Chih-Yung Chang^{d,*},
Diptendu Sinha Roy^e

^a Qiaoyun Zhang is with the School of Artificial Intelligence, Chuzhou University, Chuzhou 239000, China

^b Sze-Han Wang is with the Department of Computer Science and Information Engineering, Tamkang University, New Taipei 25137, Taiwan

^c Chung-Chih Lin is with the Department of Computer Science and Information Engineering, Tamkang University, New Taipei 25137, Taiwan

^d Chih-Yung Chang is with the Department of Computer Science and Information Engineering, Tamkang University, New Taipei 25137, Taiwan

^e Diptendu Sinha Roy is with the Department of Computer Science and Engineering, National Institute of Technology Meghalaya, Shillong, 793003, India

ARTICLE INFO

Keywords:

Recommendation

DCN

XGBoost

Knowledge graph

ABSTRACT

Recommendation systems have prospered by leveraging user-item interactions and their features for personalized recommendations. Recent advancements in deep learning further enhance these recommendation systems with powerful backbones for learning from user-item data. However, solely depending on these interactions often leads to the cold-start problem, where items lacking historical data cannot be effectively recommended. Additionally, the issue of high similarity between user and item features frequently goes unresolved. This paper introduces a Hybrid Caregiver Recommendation mechanism, called CARES, designed to recommend suitable caregivers for postpartum women using deep learning and knowledge graphs. Initially, the proposed CARES utilizes Extreme Gradient Boosting (XGBoost) to identify important features, addressing the issue of feature similarity. Then it employs *K*-Means clustering to group postpartum women and caregivers based on similar features. Subsequently, it utilizes a Deep & Cross Network (DCN) to automatically learn feature interactions and constructs knowledge graphs to tackle the cold start problem. The proposed CARES also integrates exploration and exploitation strategies to balance the accuracy and diversity of recommendations. The proposed CARES compares with existing mechanisms on real datasets, and the simulation results demonstrate its effectiveness in terms of precision, recall, and F1-Score.

1. Introduction

With the continuous growth in market trading demands, recommendation systems play an important role in daily transactions. Especially in markets with sustained demand growth, these systems help users make choices from a wide range of products or services, such as e-commerce (e.g. Alibaba), multi-media (e.g. Movies), and the caregiver recommendation system discussed in this paper. In the early stages, recommendation systems primarily relied on subjective judgments, resulting in inefficiency and lack of accuracy.

For more information, please refer to the relevant sections under submission guidelines for the journal in the Guide for Authors.

* Correspondence author at. Department of Computer Science and Information Engineering, Tamkang University, New Taipei 25137, Taiwan.

E-mail address: cychang@mail.tku.edu.tw (C.-Y. Chang).

<https://doi.org/10.1016/j.iot.2025.101769>

Available online 19 September 2025

2542-6605/© 2025 Elsevier B.V. All rights are reserved, including those for text and data mining, AI training, and similar technologies.

However, the rapid development of machine learning [1] and deep learning [2] techniques in data analysis has revolutionized the field. These advancements enable the creation of highly specialized recommendation systems and enhance their commercial value.

There are two most classical methods in the field of recommendation systems, including collaborative filtering [3–5] and content-based methods [6–8]. Collaborative filtering recommends items to users based on their historical behavior toward similar items. However, this approach suffers from the cold start problem, where new items without historical data cannot be recommended. Conversely, content-based methods recommend items by calculating the similarity between them, directly providing users with preferences. Although this strategy can avoid the cold-start problem for new items, accurately defining item similarity remains a challenge. To address these issues, some researchers [9,10] have utilized knowledge graphs as supplementary information for recommendations. Recommendation systems that incorporate knowledge graphs can not only alleviate the cold start problem for new items but also provide a more precise measure of similarity between items by analyzing knowledge graphs. Knowledge graph embedding techniques aimed to transform entities and relations into continuous vector representations while preserving graph semantics and structure [11–13], thereby enabling recommendation models to capture richer user-item relationships and mitigate sparsity issues.

Recently, with the rapid advancement of deep learning, researchers have increasingly applied neural architectures to recommendation systems [14]. Various frameworks have been proposed to model user behavior and decision-making processes, employing convolutional neural networks, recurrent networks, attention mechanisms, and hybrid deep structures [15–24]. For example, studies such as [15,16] introduced deep learning-based strategies to improve recommendation accuracy and personalization. However, these models often rely heavily on large-scale historical interaction data, making them vulnerable to the cold-start problem, particularly when encountering new users or recently registered service providers.

To mitigate these limitations, recent research has integrated knowledge graph embeddings [22], Graph Convolutional Networks (GCNs) [17], and more broadly, Graph Neural Networks (GNNs) [21,24,25], to capture high-order semantic relationships between users and items. GNN-based approaches such as DualGNN [25], Self-supervised Multimodal GCN [26], and Heterogeneous Graph-enhanced Course Recommendation (HGCR) [27] have demonstrated strong performance in sparse and multimodal recommendation settings by leveraging both graph topology and feature representations. Furthermore, emerging models like Multi-view GCN [28] and Multi-view Deep Graph Contrastive Collaborative Filtering (MD-GCCF) [29] extended the GNN paradigm to learn from multiple relational views, enhancing robustness and adaptability in complex environments. Despite their promising results, these GNN-based methods typically required high-quality, densely connected graph structures and rely on global message propagation. This often leads to over-smoothing and diminished discriminative capacity of node representations, particularly in large-scale or highly dynamic settings.

To address the above problems, this paper proposes a novel hybrid recommendation framework, called CARES, which goes beyond a mere aggregation of existing techniques. The proposed CARES is built as a modular yet interdependent architecture, where each component not only contributes its strength but also feeds into the next phase, creating a unified pipeline that enhances both effectiveness and interpretability. Specifically, XGBoost [30] is employed not just for classification but to explicitly select important and distinguishing features, effectively mitigating the issue of high feature similarity and enabling more meaningful downstream modeling. These selected features serve as the basis for K-Means clustering, which groups postpartum women and caregivers into semantically coherent clusters, facilitating localized personalization and reducing computational costs.

Subsequently, DCN [31] is applied within each cluster to model high-order nonlinear feature interactions, enabling accurate prediction of caregiver satisfaction scores. To address the cold-start problem, cluster-specific knowledge graphs are constructed to reveal latent relationships and support recommendations even when historical data is sparse, especially for new users or caregivers. Finally, CARES integrates an exploration-exploitation strategy that balances precision-driven recommendations from DCN with diversity-enhancing candidates from the knowledge graph, resulting in a more balanced and adaptive recommendation list. To the best of our knowledge, this is the first staged hybrid recommendation architecture applied in the postpartum care domain, offering both interpretability and generalization. The main contributions of this paper are summarized as follows:

- (1) **Employing K-means clustering to improve the accuracy and efficiency of recommendation:** By applying the K-means clustering technique, postpartum women and caregivers are efficiently grouped based on their features. This approach not only involves targeted training for each group to improve the accuracy of recommendations but also simplifies the identification process of candidate caregivers, effectively reducing computational demands.
- (2) **Integrating machine learning and deep learning techniques to offer a comprehensive recommendation:** The proposed CARES combines machine learning and deep learning techniques, including XGBoost and DCN to provide a comprehensive recommendation mechanism for caregivers.
- (3) **Constructing knowledge graphs to mitigate the cold-start problem:** The proposed CARES addresses the cold-start problem faced by new caregivers and postpartum women through the creation of a knowledge graph. This graph captures potential matches, ensuring effective recommendations even for newcomers. By continually updating the knowledge graph, the system accurately tracks and responds to changes in user interests, maintaining the relevance and timeliness of recommendations.
- (4) **Balancing recommendations through exploration and exploitation strategies:** The proposed CARES employs a dynamic parameter to balance exploration and exploitation. Exploration involves selecting caregivers from the knowledge graph who are not recommended by DCN, enhancing the diversity of options. Conversely, exploitation involves selecting caregivers based on the predictions of DCN, ensuring the accuracy of recommendations. These strategies balance both recommendations' accuracy and diversity.

The remainder of the paper is organized as follows. Section II discusses and compares previous relevant studies. Section III describes the detailed problem descriptions. Section IV presents the detailed proposed CARES mechanism. Section V provides the performance evaluation. The summary and future work are discussed in section VI.

2. Related work

The related work is summarized into two categories: classical recommendation models and deep learning-based recommendation systems.

2.1. Classical recommendation models

In classical recommendation models, two prevalent approaches are the group recommendation method [32–34] and the knowledge graph-based method [9,10]. Specifically, Castro et al. [32] analyzed group interests and user connections, dividing larger groups into subgroups. Each subgroup then selected similar users and items to form candidate datasets. Recommendations were generated through singular value decomposition on these datasets and integrated using a dynamic combination function to produce results. However, this method overlooked certain group features, such as the relationships between references of group members. To address this, Qin et al. [33] proposed a group recommender system based on opinion dynamics that incorporated these relationships using a smart weights matrix to guide the process. However, these approaches suffered from cold start problems, where new items without historical data could not be recommended.

To tackle cold start problems, Chen et al. [9] and Cai et al. [10] constructed a knowledge graph to capture the indirect relationships between users and items. Furthermore, Cai et al. [10] developed a many-objective recommendation algorithm based on the partition detection strategy to enhance model efficiency. This method involves associating individuals with the nearest reference vector to promote diversity and then identifying poorly converged individuals in each partition for removal. However, this multi-objective process might reduce algorithm efficiency, especially in the partition deletion strategy, where continuous assessment and removal of individuals could demand additional processing time, potentially hindering the system's ability to quickly respond to user recommendation needs. The proposed CARES utilizes *K*-Means to cluster the postpartum women and caregivers, and then constructs a knowledge graph for each cluster, significantly reducing the time and computational resources required to build and maintain the knowledge graph.

2.2. Deep learning-based recommendation

Distinct from the classical recommendation models, recent studies adopted a diverse array of deep learning architectures, spanning from autoencoder [15,16], and graph neural networks [17–24] as the main building block to mine complex and nonlinear patterns of user-item interactions. Specifically, Wu et al. [15] developed a Collaborative Neural Social Recommendation (CNSR), which integrated an autoencoder with social correlation regularization for user embedding. A unique neural network with a specialized collaboration layer was developed to model user-item interactions. However, the CNSR struggled with the severe sparsity of historical data and poor content quality. To overcome these issues, Bai et al. [16] utilized stacked denoising autoencoders to enhance feature extraction, specifically addressing the problem of low-quality description content. Additionally, to tackle the sparsity of historical usage data, they shifted the focus from modeling individual services to learning patterns of developers' preferences. However, due to the lack of historical usage data, the new items are hard to recommend using these mechanisms.

To address this issue, many studies [17,24–29] have explored knowledge graph-based and graph neural network-based recommendation frameworks. Specifically, Wang et al. [21] proposed a personalized recommender based on learnable attribute sampling and a heterogeneous graph neural network (LASGRec) to improve the recommender's performance. However, LASGRec struggled with the cross-domain cold start problem. To tackle this, Liu et al. [22] focused on multi-domain item-item recommendation using cross-domain knowledge graph embedding, analyzing associations within the same domain and interactions across diverse domains with the help of a rich knowledge graph. Additionally, Liu et al. [24] employed a GNN to represent users' session information and then utilized graph attention networks to aggregate the social information of users and their friends on social networks to effectively model users' interests. However, these methods did not address the high similarity between the features of users and items and the importance of distinguishing features.

In recent years, more advanced GNN-based models have been developed to improve recommendation accuracy under sparse or heterogeneous data conditions. For example, Wang et al. [25] designed two coupled graph neural networks to capture dual-domain semantics in multimedia recommendations, whereas Kim et al. [26] combined self-supervised signals from multiple modalities to enhance collaborative filtering robustness. Wang et al. [27] utilized heterogeneous graphs and attention mechanisms to support personalized course recommendations but were less generalizable to other domains.

Additionally, Yao et al. [28] integrated attention to fuse information from multiple user-item interaction views, while Li et al. [29] introduced a contrastive learning framework to distinguish subtle differences across user groups and features in multi-view graph spaces. Nevertheless, these approaches often require globally connected and high-quality graphs, and few explicitly address the feature similarity issue or perform automated, interpretable feature selection.

To overcome these problems, the proposed CARES utilizes XGBoost to select important features and then employs a DCN to automatically learn features of interactions between postpartum women and caregivers. It helps capture differences between features. Additionally, it constructs knowledge graphs and utilizes exploration and exploitation strategies to balance accuracy and diversity.

Table 1 presents a comprehensive comparison between the proposed CARES framework and representative related works across multiple dimensions. The “Target” column indicates whether each method is designed for group or individual recommendation scenarios. The “Cold Start Type” column specifies the type of cold-start problem addressed, including user-level (U) or item-level (I). The “Knowledge Graph” column denotes whether the model incorporates structured knowledge graphs, while the “Deep Learning” column outlines the specific deep learning techniques employed. The “Context” column identifies the application scenario of each method, including Individual, Group, Social, or Multi-domain recommendation. This contextual classification is essential for evaluating the generalizability and applicability of the compared systems. Additionally, the “Dataset” column specifies the primary dataset(s) used in each study.

Compared to prior approaches, the proposed CARES exhibits a stronger capacity to handle both user and item cold-start problems through a well-designed, multi-stage hybrid strategy. Specifically, it employs XGBoost to identify robust and distinguishing features, applies K-Means clustering to construct generalized profiles for users and items, and leverages knowledge graphs to infer latent associations for newly registered caregivers or postpartum users. Moreover, CARES is distinguished by its explicit focus on group-based recommendation, its effective handling of cold-start scenarios through knowledge-graph reasoning, and the integration of DCN to enhance the accuracy and robustness of personalized matching in the postpartum care domain.

3. Notations, assumptions and problem descriptions

This section introduces the notations, assumptions, problem descriptions and objective of this paper.

3.1. Notations and assumptions

Typically, after childbirth, Chinese women usually require the services of a postpartum caregiver at home, which include newborn care, postpartum recovery for the mother, preparation of nutritious meals, home management, and cleaning, as well as postnatal education as well as support services.

The proposed CARES has collected a large amount of human-based matched recommendation data for caregivers and postpartum women. Human-based matched recommendation often considers the unique requirements of each postpartum woman and the distinct features of caregivers. It typically relies on subjective judgments to determine if various needs can be met. To improve this process, this paper utilizes AI based on these data to develop a recommendation system. The goal is to improve the quality of matches by achieving higher rating scores given by postpartum women. In contrast to the previously subjective and time-consuming human-based comparisons, using AI not only provides interpretability but also enables more accurate recommendations.

Consider a recommendation mechanism $\mathcal{M}$. Assume that there are n caregivers $C = \{c_1, c_2, \dots, c_n\}$, where $c_i \in C$ denotes the i th postpartum caregiver. Assume that there are m postpartum women $W = \{w_1, w_2, \dots, w_m\}$, where $w_j \in W$ denotes the j -th postpartum woman. The objective of mechanism $\mathcal{M}$ is to recommend the most suitable postpartum caregiver c_i for the postpartum woman w_j from set C , ensuring that w_j gives the highest rating to the postpartum caregiver c_i . Let δ_{ij} ($1 \leq j \leq m, 1 \leq i \leq n$) denote whether caregiver c_i is recommended to the postpartum woman w_j by $\mathcal{M}$. The value of δ_{ij} can be derived from Eq. (1).

Table 1
The comparison of the proposed CARES and related work.

Paper	Target	Cold Start	Knowledge Graph	Deep Learning	Context	Dataset
[3,32–34]	Group	×	×	×	Group	Programmable Web, MovieLens, Douban Movie
[6]	Individual	×	×	×	Individual	Korean Film Council
[9]	Group	I	✓	×	Group	ACMDL, PubMed
[10]	Individual	U	✓	×	Multi-domain	CiaoDVD
[15,16]	Individual	×	×	AutoEncoder	Social	Flixster, Douban, ProgrammableWeb
[17]	Individual	×	✓	GCN	Individual	MovieLens
[20]	Individual	U/I	✓	GNN and Attention	Social	Last-FM, Movielens-1 M, Book-Crossing, Dianping Food
[21]	Individual	U	✓	GNN	Social	MovieLens, Book-Crossing
[22]	Group	U	✓	Graph Embedding	Multi-domain	FB15K-237, KG3Domain
[23]	Individual	I	✓	Graph AutoEncoder	Individual	Mafengwo, CAMRa2011
[24]	Individual	U	✓	GNN	Social	Douban, Epinions
[25]	Individual	×	×	Dual GNN	Multimedia	Movielens, Tiktok
[26]	Individual	U	×	Self-supervised GCN	Social	Amazon
[27]	Individual	×	✓	Heterogeneous GNN + Attention	Education	MOOCube, CR
[28]	Individual	×	×	Multi-view GCN + Attention	Individual	Cora, Citeseer, Pubmed
[29]	Individual	U	×	Multi-view Graph Contrastive	Individual	Amazon Book, Yelp2018, Gowalla
Proposed CARES	Group	U/I	✓	DCN	Group	YesMa, Book-Crossing, Last.FM

$$\delta_{ij} = \begin{cases} 1, & c_i \text{ is recommended to } w_j, \\ 0, & \text{otherwise.} \end{cases} \quad (1)$$

Assume that each postpartum woman w_j is described by a set of multiple features. Let ρ denote the total number of features. The feature set of w_j can be represented as:

$$F_j^w = (\mathcal{F}_{j,1}^w, \mathcal{F}_{j,2}^w, \dots, \mathcal{F}_{j,\rho}^w), \quad (2)$$

where F_j^w denotes the features set of w_j , and $\mathcal{F}_{j,k}^w$ ($1 \leq k \leq \rho$) denotes the k -th feature of w_j , including age, address, number of other children, service duration, infant care, maternal and child health, nutrition, due date, educational level, budget, and income.

Similarly, the feature set of c_i can be represented as:

$$F_i^c = (\mathcal{F}_{i,1}^c, \mathcal{F}_{i,2}^c, \dots, \mathcal{F}_{i,\rho}^c), \quad (3)$$

where F_i^c denotes the features set of c_i , and $\mathcal{F}_{i,k}^c$ ($1 \leq k \leq \rho$) denotes the k -th feature of c_i , including age, service area, experience, level, cooking skills, idle, and rating. Let $s_{i,j}$ denote the satisfaction of postpartum women w_j served by the postpartum caregiver c_i using mechanism $\mathcal{M}$. Let $s_{\mathcal{M}}$ denote the threshold of satisfaction. Let $\hat{\beta}_{ij}$ and β_{ij} denote the predicted and labeled recommendation results for postpartum caregiver c_i to the postpartum woman w_j , respectively. More specifically, the value of β_{ij} is 1 if $\mathcal{M}$ recommends caregiver c_i to the postpartum woman w_j . Otherwise, the value of β_{ij} is 0. On the other hand, the value of label β_{ij} can be represented by Eq. (4).

$$\hat{\beta}_{ij} = \begin{cases} 1, & s_{i,j} \geq s_{\mathcal{M}}, \\ 0, & \text{otherwise.} \end{cases} \quad (4)$$

3.2. Problem descriptions

Consider the confusion matrix $\mathfrak{I}_j = [\mathfrak{I}_{a,b}^j]_{n \times n}$ for postpartum woman w_j using mechanism $\mathcal{M}$, where $\gamma_{a,b}^j$ denotes that postpartum caregiver $c_a \in C$ is suitable for w_j , while the mechanism $\mathcal{M}$ recommends postpartum caregiver $c_b \in C$ for w_j .

Let $TP_j^{\mathcal{M}}$ represent the true positive of $\mathfrak{I}_j$. The value of $TP_j^{\mathcal{M}}$ can be computed using Eq. (5).

$$TP_j^{\mathcal{M}} = \sum_{a=1}^n \mathfrak{I}_{a,a}^j \quad (5)$$

Let $FP_j^{\mathcal{M}}$ and $FN_j^{\mathcal{M}}$ represent the false positive and false negative of $\mathfrak{I}_j$, respectively. The value of $FP_j^{\mathcal{M}}$ and $FN_j^{\mathcal{M}}$ can be computed using Eqs. (6) and (7), respectively.

$$FP_j^{\mathcal{M}} = \sum_{a=1}^n \sum_{b \neq a}^n \mathfrak{I}_{a,b}^j \quad (6)$$

$$FN_j^{\mathcal{M}} = \sum_{b=1}^n \sum_{a \neq b}^n \mathfrak{I}_{a,b}^j \quad (7)$$

Let $\mathbb{U}$ be the union of all $\mathfrak{I}_j$, for $1 \leq j \leq m$. That is,

$$\mathbb{U} = \bigcup_{j=1}^m \mathfrak{I}_j \quad (8)$$

Let $TP^{\mathcal{M}}$, $FP^{\mathcal{M}}$ and $FN^{\mathcal{M}}$ denote true positive, false positive, and false negative of all $\mathbb{U}$, respectively. The value of $TP^{\mathcal{M}}$, $FP^{\mathcal{M}}$ and $FN^{\mathcal{M}}$ can be measured by applying Eqs. (9), (10) and (11), respectively.

$$TP^{\mathcal{M}} = \sum_{j=1}^m TP_j^{\mathcal{M}} \quad (9)$$

$$FP^{\mathcal{M}} = \sum_{j=1}^m FP_j^{\mathcal{M}} \quad (10)$$

$$FN^{\mathcal{M}} = \sum_{j=1}^m FN_j^{\mathcal{M}} \quad (11)$$

Let $\mathcal{P}^{\mathcal{M}}$ and $\mathcal{R}^{\mathcal{M}}$ denote the precision and recall of $\mathbb{U}$, respectively. The values of $\mathcal{P}^{\mathcal{M}}$ and $\mathcal{R}^{\mathcal{M}}$ can be evaluated by applying Eqs. (12) and (13), respectively.

$$\mathcal{P}^{\mathcal{M}} = \frac{TP^{\mathcal{M}}}{TP^{\mathcal{M}} + FP^{\mathcal{M}}} \quad (12)$$

$$\mathcal{R}^{\mathcal{M}} = \frac{TP^{\mathcal{M}}}{TP^{\mathcal{M}} + FN^{\mathcal{M}}} \quad (13)$$

Let $\mathcal{F}1^{\mathcal{M}}$ denote F1-score of $\mathbb{U}$ obtained by the prediction of mechanism $\mathcal{M}$. The value of $\mathcal{F}1^{\mathcal{M}}$ can be calculated by Eq. (14).

$$\mathcal{F}1^{\mathcal{M}} = \frac{2 * \mathcal{P}^{\mathcal{M}} * \mathcal{R}^{\mathcal{M}}}{\mathcal{P}^{\mathcal{M}} + \mathcal{R}^{\mathcal{M}}} \quad (14)$$

C. Objective

Assume that ψ be the set of all possible mechanisms employed to choose the caregivers for $W = \{w_1, w_2, \dots, w_m\}$. The main goal of this paper is to find the most effective mechanism $\mathbb{M} \in \psi$, which satisfies Eq. (15).

F1-Score Objective:

$$\mathbb{M} = \arg \max_{\mathcal{M} \in \psi} \mathcal{F}1^{\mathcal{M}} \quad (15)$$

When optimizing for the first objective, the proposed mechanism CARES incorporates the second objective to balance the recommendation opportunities for those caregivers who are the best candidates that satisfy the first objective. The second objective is to minimize the number of times the same postpartum caregiver is recommended while maintaining diversity and personalization in the recommendation list. Let $f^{\mathcal{M}}(c_i)$ represent the recommendation frequency of postpartum caregiver c_i by applying mechanism $\mathcal{M}$. The second objective of this paper is to identify the most effective mechanism Φ , which has the minimal sum of squares of all caregivers' recommendation frequencies, thereby ensuring effective control over the recommendation frequency of each postpartum caregiver, as shown in Eq. (16).

Fairness Objective:

$$\Phi = \max_{\mathcal{M} \in \mathbb{M}} \frac{(\sum_{i=1}^n f^{\mathcal{M}}(c_i))^2}{n \sum_{i=1}^n f^{\mathcal{M}}(c_i)^2} \quad (16)$$

While achieving the best solution, some constraints should be satisfied. The first constraint is the Postpartum Caregiver Constraint, which guarantees that each postpartum woman w_j employs at most one postpartum caregiver at any given time $t \in T$. Additionally, a postpartum caregiver will not provide services to multiple postpartum women simultaneously at any given time $t \in T$. Let $\delta_{ij}(t)$ denote whether w_j employs c_i at time $t \in T$.

$$\delta_{ij}(t) = \begin{cases} 1, & w_j \text{ employs } c_i \text{ at time } t \in T, \\ 0, & \text{otherwise.} \end{cases} \quad (17)$$

The postpartum caregiver constraint is given in Eq. (18).

(1) Postpartum Caregiver Constraint

$$\begin{aligned} \forall w_j \in W, \forall t \in T \text{ such that } \sum_{i=1}^n \delta_{ij}(t) &\leq 1 \\ \forall c_i \in C, \forall t \in T \text{ such that } \sum_{j=1}^m \delta_{ij}(t) &\leq 1 \end{aligned} \quad (18)$$

Assume that there is a set of requirements defined by postpartum women, denoted by $R = (R_1, \dots, R_q)$ which consists of q distinct requirements. These requirements include working area, living with the postpartum women, cooking proficiently, and fulfilling additional responsibilities. Similarly, assume that there is a set of supports of caregivers, denoted by $S = (S_1, \dots, S_q)$, which includes q types of supports. Let λ_i denote whether the i th support matches the i th requirement. The value of λ_i can be expressed by Eq. (19).

$$\lambda_i = \begin{cases} 1, & \text{if } S_i \text{ matches } R_i, \\ 0, & \text{otherwise.} \end{cases} \quad (19)$$

The second constraint is the Requirement and Support Matching constraint, ensuring that the support provided to the postpartum woman aligns with her requirements throughout the recommendation process. This constraint is mathematically formulated as the product of all λ_i values equaling 1. It can be expressed in Eq. (20).

(2) Requirement and Support Matching Constraint

$$\prod_{i=1}^q \lambda_i = 1 \quad (20)$$

4. The proposed cares mechanism

This paper proposes an innovative caregiver recommendation mechanism, named CARES, which aims to match postpartum women

with suitable caregivers by integrating clustering, deep learning, and knowledge graph techniques. As shown in Fig. 1, the proposed CARES consists of four phases: (1) Data Preprocessing, (2) Feature Selection with XGBoost, (3) Feature Clustering with K-Means, and (4) Ranking with DCN and Knowledge Graph. In the first phase, raw input data from postpartum women and caregivers are cleaned, normalized, and aligned, producing standardized feature vectors as the foundation for subsequent analysis. These processed features are then passed into the second phase, where XGBoost, a gradient-boosting decision tree algorithm, is employed to evaluate feature importance and filter out less relevant features. The resulting subset of filtering features is used as the input for the third phase, where K-Means clustering groups postpartum women into clusters with similar needs and characteristics. Finally, in the ranking phase, the filtered candidate caregiver group serves as inputs to the DCN, which captures both low-order and high-order feature interactions for accurate matching. Simultaneously, the knowledge graph complements this process by inferring hidden relationships, which is especially beneficial for alleviating the cold-start problem. The outputs of DCN and knowledge graph reasoning are jointly integrated into the final ranking module, where an exploration-exploitation strategy balances accuracy with diversity to generate personalized recommendations. The following gives the details of these phases.

4.1. Data preprocessing

The data preprocessing phase aims to clean and normalize raw input data from postpartum women and caregivers to ensure consistency and comparability across all features. This step is essential for preparing high-quality input suitable for subsequent machine-learning tasks.

Let $\mathcal{V}'_i = \{v'_{i,1}, v'_{i,2}, \dots, v'_{i,a}\}$ denote a bag of values in the i th feature $\mathcal{F}_i$. If $v'_{ij} \in \mathcal{V}'_i$ is a numerical value, such as age, number of other children, service duration, budget, and income, Min-Max normalization is applied to ensure comparability. This normalization method rescales feature values to fit within a specified range, usually between 0 and 1, by subtracting the minimum value of the feature and then dividing by the range of the feature values. Let v_{ij} denote the normalized value of v'_{ij} . The value of v_{ij} can be computed from Eq. (21).

$$v_{ij} = \frac{v'_{ij} - \min_{1 \leq k \leq a} v'_{i,k}}{\max_{1 \leq k \leq a} v'_{i,k} - \min_{1 \leq k \leq a} v'_{i,k}} \quad (21)$$

If $v'_{ij} \in \mathcal{V}'_i$ is string value, such as an address, educational level and other non-numerical features, the one-hot encoding normalization is applied. This normalization maps the values to a binary vector, where the length of the vector equals the number of categories in feature $\mathcal{F}_i$. Let $B = (b_1, b_2, \dots, b_a)$ ($b \leq a$) denote the disjoint categories. Let $v_{ij} = (e_{j,1}^i, e_{j,2}^i, \dots, e_{j,b}^i)$ denote the

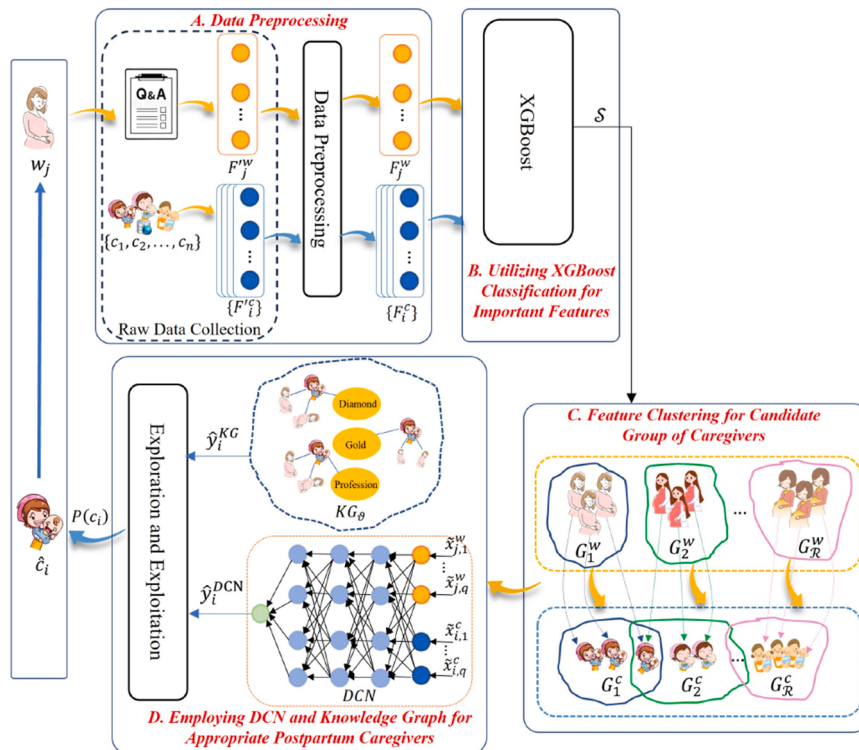


Fig. 1. The architecture of proposed CARES.

normalized vector of the v_{ij} after applying one-hot encoding, where $v_{j,k}^i$ ($1 \leq j \leq a, 1 \leq k \leq b$) denote whether or not v_{ij} is the same as b_k . The value of $v_{j,k}^i$ is calculated according to Eq. (22).

$$v_{j,k}^i = \begin{cases} 1, & v_{ij} = b_k, \\ 0, & \text{otherwise.} \end{cases} \quad (22)$$

This phase serves as the foundation for the proposed CARES pipeline, ensuring that all heterogeneous data whether numeric or categorical are encoded in a standardized format. The resulting preprocessed feature set denoted as $\mathbb{S}$, is then passed to the Feature Selection with XGBoost in the next phase.

4.2. Feature selection with XGBoost

To reduce the dependency on manual domain knowledge and improve model adaptability across datasets, the proposed CARES framework adopts a semi-automated feature selection strategy based on XGBoost. Initially, candidate features are directly extracted from the raw caregiver data without domain-specific filtering or encoding. XGBoost, a gradient-boosting decision tree model [35], is then employed to rank and select the most relevant features for predicting caregiver suitability.

Let $X^{c_i} = (x_{i,1}^c, \dots, x_{i,j}^c, \dots, x_{i,\rho}^c)$ denote the feature vector for caregiver c_i , where $x_{i,j}^c \in \mathbb{S}$, ρ is the number of features. Let y_i^X denote whether the caregiver is classified as satisfactory, where 1 represents satisfactory and 0 represents unsatisfactory. Let $\hat{y}_i^X$ denote the prediction result. The value of $\hat{y}_i^X$ can be derived by Eq. (23) by applying the XGBoost model.

$$\hat{y}_i^X = \sum_{k=1}^K f_k(X^{c_i}), f_k \in \mathcal{T} \quad (23)$$

where $f_k(\cdot)$ represents the k -th decision tree, $\mathcal{T}$ denotes the set of all possible decision trees, and K is the total count of trees.

The objective function of XGBoost, as shown in Eq. (24), consists of two parts: the loss function on training data and the complexity of the model to prevent overfitting. The objective function needs to be minimized when each decision tree is added to the model. That is

$$\mathcal{L}_{XGB} = \sum_{i=1}^{\varphi} \mathcal{L}^X(y_i^X, \hat{y}_i^X) + \sum_{k=1}^K \Omega(f_k) \quad (24)$$

where $\mathcal{L}_{XGB}$ denotes the objective function, $\mathcal{L}^X(y_i^X, \hat{y}_i^X)$ is the loss function by applying XGBoost, capturing the discrepancy between the actual value y_i^X and the predicted value $\hat{y}_i^X$, the φ denotes the training number of caregivers, and the $\Omega(f_k)$ denotes the complexity penalty term for the k -th decision tree, controlling the model's complexity.

Different from the study [35], the tree model used in this paper adopts the Gini Index as the loss function. The reason is that, compared with the originally designed cross-entropy loss function, the calculation of the Gini coefficient is simpler, which can accelerate the convergence speed of the model. The values of $\mathcal{L}^X(y_i^X, \hat{y}_i^X)$ and $\Omega(f_k)$ can be derived from Eqs. (25) and (26), respectively.

$$\mathcal{L}^X(y_i^X, \hat{y}_i^X) = 1 - \left((\hat{y}_i^X)^2 + (1 - \hat{y}_i^X)^2 \right) \quad (25)$$

$$\Omega(f_k) = \beta T + \frac{1}{2} \gamma \sum_{j=1}^T \omega_j^2 \quad (26)$$

where β and γ are regularization factors, T is the number of leaf nodes, ω_j is a weight parameter of each leaf in the XGBoost model. After training, XGBoost automatically computes the importance score $I(x_{i,j}^c)$ of each feature $x_{i,j}^c$, based on its frequency and depth in the trees. Let γ denote a threshold, which is applied to the importance scores. A feature $x_{i,j}^c$ is considered important if its importance score $I(x_{i,j}^c)$ exceeds γ . Let $\mathcal{S}$ denote the resulting set of important features. The value of $\mathcal{S}$ can be defined as Eq. (27).

$$\mathcal{S} = \{x_{i,j}^c \mid I(x_{i,j}^c) \geq \gamma\} \quad (27)$$

At the end of this phase, the proposed CARES produces a refined set of important features $\mathcal{S}$. The output is forwarded to the next phase, where caregivers and postpartum women are grouped to facilitate efficient and personalized recommendations.

4.3. Feature clustering with K-Means

This phase employs K-Means to group both postpartum women and caregivers into semantically meaningful clusters. This process reduces the recommendation space and supports efficient, personalized matching, particularly in cold-start scenarios.

First, postpartum women are clustered using their normalized feature vectors. K-Means is selected for its simplicity, scalability, and suitability for high-dimensional data. Users within the same cluster share similar care needs and preferences.

Second, caregivers are grouped based on the clusters of the postpartum women they have previously served. A caregiver may belong to multiple clusters if they have served women from different groups, capturing cross-group experience and improving flexibility. When a new postpartum woman joins the system, her feature vector is compared to cluster centroids to identify the most similar group. The corresponding caregiver cluster then becomes her candidate recommendation set.

Recall that there are m postpartum women, denoted as $W = \{w_1, w_2, \dots, w_m\}$. The proposed CARES employs K-Means to separate these postpartum women into $\mathcal{R}$ clusters, denoted as $G^w = \{G_1^w, G_2^w, \dots, G_{\mathcal{R}}^w\}$. This ensures the postpartum women in the same cluster have similar features. Let $\mathcal{U}_k$ ($1 \leq k \leq \mathcal{R}$) denote the centroid of k -th cluster $G_k^w \in G^w$. Each centroid is represented as a q_w -dimensional vector corresponding to the q_w selected features. That is

$$\mathcal{U}_k = (u_{k,1}^w, u_{k,2}^w, \dots, u_{k,q_w}^w) \quad (28)$$

where $u_{k,a}^w \in \mathcal{U}_k$ ($1 \leq a \leq q_w$) denotes the average value of the a -feature across all postpartum women assigned to cluster G_k^w .

The features F_j^w of w_j are compared with each cluster centroid. Let $d(F_j^w, \mathcal{U}_k)$ denote the distance between F_j^w and $\mathcal{U}_k$. It can be calculated by Eq. (29).

$$d(F_j^w, \mathcal{U}_k) = \sum_{a=1}^{q_w} \xi_a (\mathcal{F}_{j,a}^w - u_{k,a}^w)^2 \quad (29)$$

where ξ_a denotes the weight of the a -th feature. It satisfies the following condition:

$$\sum_{a=1}^{q_w} \xi_a = 1 \quad (30)$$

The w_j is assigned to a specific cluster based on the distance $d(F_j^w, \mathcal{U}_k)$. Let G_g^w denote the cluster whose centroid is closest to w_j . This can ensure that w_j is grouped with the most similar cluster in terms of its features. The specific value of g can be derived from Eq. (31).

$$g = \arg \min_{1 \leq k \leq \mathcal{R}} d(F_j^w, \mathcal{U}_k) \quad (31)$$

Subsequently, each cluster recomputes its centroid $\mathcal{U}_k$ by updating each feature dimension based on the current members of the cluster. Specifically, the updated value $u_{k,a}^w \in \mathcal{U}_k$ for the a -feature dimension in cluster G_k^w is computed as the average of that feature among all postpartum women assigned to the cluster. That is

$$u_{k,a}^w = \frac{1}{|G_k^w|} \sum_{w_j \in G_k^w} \xi_a \mathcal{F}_{j,a}^w \quad (32)$$

where $|G_k^w|$ denotes the number of postpartum women in cluster G_k^w , and $\mathcal{F}_{j,a}^w$ denotes the a -feature value of woman w_j . Finally, the proposed CARES repeats the operations from Eqs. (29) to (32) until the centroids of the clusters stabilize, or a predetermined number of iterations is reached. These processes ensure the effective grouping of postpartum women into different clusters based on the features, facilitating a faster identification of candidate caregivers.

Caregivers $C = \{c_1, c_2, \dots, c_n\}$ who have been paired with postpartum women are grouped according to the postpartum women

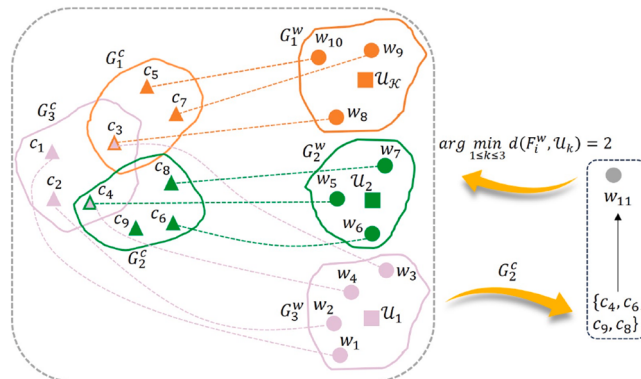


Fig. 2. The process for grouping postpartum women and caregivers applying K-Meas.

they have served. Therefore, regarding the clustering result of a postpartum woman, the caregivers who have served her should also be clustered into one group. Let $G^c = \{G_1^c, G_2^c, \dots, G_{\mathcal{R}}^c\}$ denote the set of clusters of caregivers. It is noticed that there might be overlaps among these clusters of caregivers. The reason is that one experienced caregiver may serve two or more postpartum women, but these women belong to different clusters. Therefore, the same caregiver may need to appear in different clusters.

The abovementioned clustering can cope with the cold start problem. When a new postpartum woman w_e is introduced, the proposed CARES employs Eqs. (29) and (31) to calculate the distance between w_e and each cluster and determine the closest cluster. All members in this cluster will play the role of caregiver candidates. Let G_θ^w ($1 \leq \theta \leq \mathcal{R}$) denote the closest cluster to w_e . Consequently, the candidate group of caregivers for w_e is identified as G_θ^c , ensuring an effective match between the postpartum women and their caregivers.

As shown in Fig. 2, the set of postpartum women, denoted as $W = \{w_1, w_2, \dots, w_{10}\}$, is categorized into three clusters $\{G_1^w, G_2^w, G_3^w\}$. Similarly, the caregivers, represented by $C = \{c_1, c_2, \dots, c_9\}$ who have matched with these postpartum women, are organized into corresponding clusters $\{G_1^c, G_2^c, G_3^c\}$. Notably, caregivers c_3 and c_4 are special as they belong to two clusters. The reason is that c_3 has served different postpartum women w_3 and w_8 from different clusters G_3^w and G_1^w , respectively. Furthermore, consider a new postpartum woman w_{11} who requires matching with a caregiver group. According to Eq. (31), the cluster G_2^w is identified as the most like the feature of w_{11} . Consequently, the $G_2^c = \{c_4, c_6, c_8, c_9\}$ emerges as the candidate group of caregivers to provide the necessary support for w_{11} .

At the end of this phase, the proposed CARES outputs a filtered candidate caregiver group G_θ^c , which is forwarded to the final ranking module integrating DCN and Knowledge Graphs.

4.4. Ranking with DCN and knowledge graph

In this final phase, the proposed CARES employs a DCN and a Knowledge Graph to rank candidate caregivers in G_θ^c for a given postpartum woman. DCN models both low- and high-order feature interactions between users and caregivers, generating a predicted satisfaction score for each candidate. In parallel, the Knowledge Graph enhances the recommendation process by capturing semantic relationships, such as shared experience, service context, or institutional affiliation, which is particularly valuable in cold-start scenarios where historical data is limited. To further improve recommendation quality, the proposed CARES incorporates an exploration-exploitation strategy that balances accuracy with diversity, ensuring that the results are not only highly relevant but also varied enough to support flexible and informed decision-making.

Let $\mathcal{X}_i^c = (\tilde{x}_{i,1}^c, \tilde{x}_{i,2}^c, \dots, \tilde{x}_{i,q_c}^c)$ ($\tilde{x}_{i,s}^c$ ($1 \leq s \leq q_c$) $\in \mathcal{S}$, $q_c \leq \mathcal{R}$) denote the features vector for caregiver c_i in G_θ^c . Let $\mathcal{X}_j^w = (\tilde{x}_{j,1}^w, \tilde{x}_{j,2}^w, \dots, \tilde{x}_{j,q_w}^w)$ denote the features vector for the postpartum woman w_j . As shown in Fig. 3, these features are transformed to dense features through an embedding and stacking layer, denoted by $r_{i,0}$. Subsequently, $r_{i,0}$ is fed into both cross-network and deep network. The

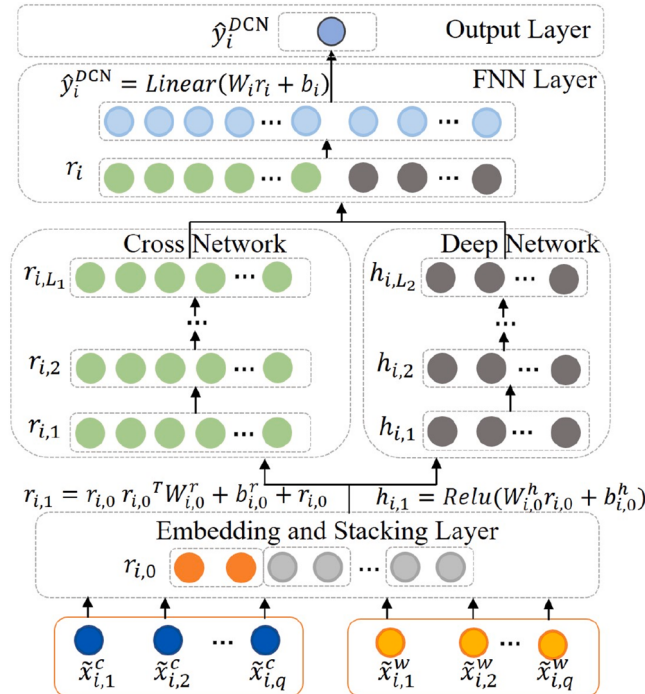


Fig. 3. The process of DCN.

cross-network facilitates explicit feature crossing, while the deep network composed of multiple fully connected layers, captures the non-linear relationships among features. Let $r_{i,k}$ and $h_{i,k}$ ($k \geq 1$) denote the output of k -th layer in the cross-network and deep network, respectively. Let L_1 and L_2 denote the total number of layers in the cross-network and deep network, respectively. The value of $r_{i,k}$ and $h_{i,k}$ can be computed using Eqs. (33) and (34), respectively.

$$r_{i,k} = r_{i,k-1} (r_{i,k-1})^T W_{i,k-1}^r + b_{i,k-1}^r + r_{i,k-1} \quad (33)$$

$$h_{i,k} = \text{Relu}(W_{i,k-1}^h r_{i,k-1} + b_{i,k-1}^h), \quad (34)$$

where $W_{i,k-1}^r$ and $W_{i,k-1}^h$ denote the weights parameters of $(k-1)$ -th layer in cross and deep networks, respectively, while $b_{i,k-1}^r$ and $b_{i,k-1}^h$ denote the bias parameters of $(k-1)$ -th layer in cross and deep networks, respectively.

The combination layer concatenates the outputs of the cross and deep networks. Let r_i denote the result of the combined layer. That is

$$r_i = \text{Concat}(r_{i,L_1}, h_{i,L_2}). \quad (35)$$

Finally, the r_i passes through the Feedforward Neural Network (FNN), which comprises a linear layer followed by a Neural cell as its output. This FNN is used to compute the satisfaction score of the caregiver c_i . Let y_i^{DCN} and $\hat{y}_i^{DCN}$ denote the labeled and prediction output, respectively. The $\hat{y}_i^{DCN}$ can be derived from Eq. (36).

$$\hat{y}_i^{DCN} = \text{Linear}(W_i r_i + b_i), \quad (36)$$

where W_i and b_i denote the weight and bias parameters. Let $L^{DCN}(y_i^{DCN}, \hat{y}_i^{DCN})$ denote the loss function by applying DCN. The value of $L^{DCN}(y_i^{DCN}, \hat{y}_i^{DCN})$ can be calculated by Eq. (37).

$$L^{DCN}(y_i^{DCN}, \hat{y}_i^{DCN}) = \sqrt{\sum_{i=1}^N \frac{(\hat{y}_i^{DCN} - y_i^{DCN})^2}{N}} \quad (37)$$

where N denotes the input number of candidate caregivers.

During the inference phase, let $\mathbb{C}_{DCN}$ denote the set of satisfactory caregivers. The initial value of $\mathbb{C}_{DCN}$ is an empty set, denoted by $\emptyset$. The set of $\mathbb{C}_{DCN}$ can be derived by Eq. (38).

$$\mathbb{C}_{DCN} = \begin{cases} \mathbb{C}_{DCN} \cup \{c_i\}, & \hat{y}_i^{DCN} > \eta, \\ \mathbb{C}_{DCN}, & \text{otherwise.} \end{cases} \quad (38)$$

where η denotes the threshold to classify the predicted output $\hat{y}_i^{DCN}$.

Although the above approach can train and predict recommendation models for postpartum women, there still exists a cold-start issue for new caregivers and postpartum women. This means that for a new caregiver, there is no record in the past, and likewise, there is no record of any caregiver serving the new postpartum women. The proposed CARES will address the cold-start problem using knowledge graphs. When a new caregiver joins, she is allocated to a specific caregiver group based on feature similarity. However, her absence of previous pairings with any postpartum woman may lead to a lower probability of being recommended when directly incorporated into the existing models. To address this cold-start issue, the proposed CARES improves the matching process by con-

Table 2
The process of constructing the knowledge graph.

Algorithm 1: Knowledge Graph Construction

Input: $w_j \in G_\theta^w$, $c_i \in G_\theta^c$, F , F_j^w , F_i^c
Output: nodes set $\mathcal{V}_\theta$, edges set $\mathcal{E}_\theta$
for each $w_j, c_i, \mathcal{F}_k \in F$
 add $w_j, c_i, \mathcal{F}_k$ in set $\mathcal{V}_\theta$
end for
for each $w_j \in G_\theta^w$
 if w_j have matched with c_i
 add $\langle w_j, c_i \rangle$ in set $\mathcal{E}_\theta$
 end if
end for
for each $\mathcal{F}_{jp}^w \in F_j^w, \mathcal{F}_{iq}^c \in F_i^c$
 add $\langle w_j, \mathcal{F}_{jp}^w \rangle$ in set $\mathcal{E}_\theta$
 add $\langle c_i, \mathcal{F}_{iq}^c \rangle$ in set $\mathcal{E}_\theta$
end for

structing a knowledge graph of postpartum women and caregivers in each cluster, which is categorized by K -Means. Let $KG_\theta = (\mathcal{V}_\theta, \mathcal{E}_\theta)$ denote the constructed knowledge graph of G_θ^w and G_θ^c , where $\mathcal{V}_\theta$ and $\mathcal{E}_\theta$ denote the sets of nodes and edges in KG_θ , respectively. Let F denote all features of the postpartum women and caregivers. The process of constructing the knowledge graph is shown in Table 2. Through the constructed knowledge graph, the proposed CARES can effectively capture potential matching opportunities, ensuring that a new caregiver is also appropriately recommended.

Assume that there is a new postpartum woman denoted as w_e . According to Eq. (31), assume that G_θ^w is identified as the cluster nearest to w_e . This implies that the feature of the new postpartum is most like the group G_θ^w of postpartum women. In the G_θ^w , assume that w_e is the nearest postpartum woman to the w_j in terms of distance according to Eq. (29). Let $Nbs(w_j)$ denote the set of caregivers neighboring to w_j in KG_θ . The value of $Nbs(w_j)$ can be derived from Eq. (39).

$$Nbs(w_j) = \{c_i \in \mathcal{V}_\theta | \langle w_j, c_i \rangle \in \mathcal{E}_\theta\} \quad (39)$$

Let $\mathcal{N}_k(w_j)$ denote the set of caregivers that are k -step neighbors of w_j in KG_θ , where $\mathcal{N}_1(w_j) = Nbs(w_j)$. The value of $\mathcal{N}_k(w_j)$ can be derived from Eq. (40).

$$\mathcal{N}_k(w_j) = \bigcup_{u \in \mathcal{N}_{k-1}(w_j)} Nbs(u) \setminus \bigcup_{z=1}^{k-1} \mathcal{N}_z(w_j) \quad (k \geq 2) \quad (40)$$

Let $\mathbb{C}_{KG}$ denote the set of satisfactory caregivers as predicted by the knowledge graph. The initial value of $\mathbb{C}_{KG}$ is an empty set, denoted by $\emptyset$. The set of $\mathbb{C}_{KG}$ can be derived by Eq. (41).

$$\mathbb{C}_{KG} = \sum_{l=1}^k \mathcal{N}_l(w_j) \quad (41)$$

Let $\hat{y}_i^{KG}$ denote the satisfactory score of recommending $c_i \in \mathbb{C}_{KG}$ to w_j using a knowledge graph. The value of $\hat{y}_i^{KG}$ can be calculated using cosine similarity as shown in Eq. (42).

$$\hat{y}_i^{KG} = \cos(\mathcal{X}_i^c, \mathcal{X}_j^w) \quad (42)$$

Fig. 4 illustrates the inference process of candidate caregiver recommendation using the knowledge graph. Specifically, when a newly registered postpartum woman w_{11} enters the system, she is classified into G_2^w and is closest to w_7 in terms of distance according to Eqs. (29) and (31). Based on this relationship, the knowledge graph retrieves the two-hop neighbors of w_7 as potential matches. Consequently, the set of candidate caregivers is obtained as $\mathbb{C}_{KG} = \mathcal{N}_2(w_7) = \{c_4, c_8, c_9\}$.

Till now, the proposed CARES has recommended the caregivers through two mechanisms, including DCN and knowledge graph. Let $\mathbb{C}_{DCN}$ and $\mathbb{C}_{KG}$ denote the two sets of caregivers recommended by applying DCN and knowledge graph, respectively. The proposed CARES further incorporates $\mathbb{C}_{DCN}$ and $\mathbb{C}_{KG}$, and utilizes exploration and exploitation strategies to balance accuracy and diversity. Let $P(c_i)$ denote the final recommendation probability, which integrates information from $\mathbb{C}_{DCN}$ and $\mathbb{C}_{KG}$. Let ϵ ($0 \leq \epsilon < 1$) denote a parameter for balancing exploration and exploitation. The exploration denotes that the proposed CARES randomly selects a caregiver c_i , who is not only recommended by $\mathbb{C}_{DCN}$ but also presented in $\mathbb{C}_{KG}$. And the exploitation denotes selecting a caregiver c_i from $\mathbb{C}_{DCN}$ based on the prediction of DCN. The value of $P(c_i)$ can be calculated by Eq. (43).

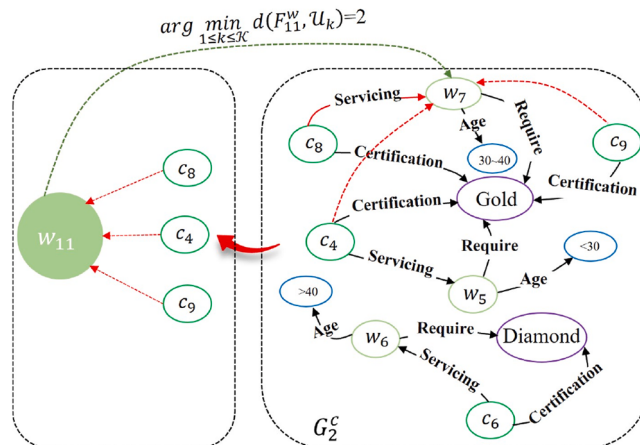


Fig. 4. Knowledge graph-based candidate caregiver inference process.

$$P(c_i) = \begin{cases} 1 - \epsilon + \frac{\epsilon}{\rho}, & \text{if } c_i \in \mathbb{C}_{DCN}, \\ \frac{\epsilon}{\rho}, & \text{if } c_i \in (\mathbb{C}_{KG} \setminus \mathbb{C}_{DCN}), \end{cases} \quad (43)$$

where ρ denotes the total number of caregivers recommended by DCN and knowledge graph.

4.5. Computational complexity analysis

The proposed CARES framework integrates four primary modules: XGBoost-based feature selection, K-Means clustering, DCN-based interaction modeling, and knowledge graph-based reasoning. The following provides a detailed analysis of the computational complexity of each component.

(1) XGBoost Feature Selection Complexity

XGBoost is used during the offline training phase to rank feature importance and eliminate redundant attributes. For $\mathfrak{N}$ training samples and ρ features, the training complexity of XGBoost is approximately:

$$O(\mathfrak{N} \cdot \rho \cdot \log \mathfrak{N}) \quad (44)$$

Since feature selection is executed only once during training and not in the online phase, its impact on runtime performance is negligible.

(2) K-Means Clustering Complexity

K-Means is used to partition users and caregivers into semantically similar groups. Let n denote the number of caregivers, k denotes the number of clusters, i denotes the maximum number of iterations, and ρ' denotes the reduced feature dimension after XGBoost. That is

$$O(k \cdot n \cdot i \cdot \rho') \quad (45)$$

As clustering is also performed offline, it does not contribute to the real-time complexity.

(3) DCN Inference Complexity

DCN is responsible for modeling high-order feature interactions during online prediction. For an input vector of dimension ρ' and l cross layers, the per-query inference complexity is computed by Eq. (46).

$$O(\rho' \cdot l) \quad (46)$$

Since DCN operates only on a small subset of selected features, the online prediction is efficient.

(4) Knowledge Graph Reasoning Complexity

During online recommendation, CARES performs reasoning over a localized knowledge graph. Let v denote the number of nodes in the subgraph and h denotes the number of hops for traversal. The worst-case complexity of Knowledge Graph reasoning is:

$$O(v^h) \quad (47)$$

However, in practice, v and h are kept small due to prior clustering, ensuring efficient subgraph access and reasoning.

(5) Overall Complexity of CARES (Online Phase)

Since XGBoost and K-Means are executed offline, the total online computational complexity of CARES is determined by DCN inference and localized KG reasoning as shown in Eq. (48). This complexity is efficient and scalable for real-time service recommendation.

$$O(\rho' \cdot l) + O(v^h) \quad (48)$$

5. Performance evaluation

This section evaluates a comparative performance analysis of the proposed CARES against the Clustering & DCN and related work DCN-V2 [31]. Clustering & DCN utilizes K-means clustering with the DCN framework to recommend caregivers. The DCN-V2 refined the DCN model with comprehensive hyper-parameter optimization, making it more practical in a large-scale environment. However, both Clustering & DCN and DCN-V2 were plagued by the cold start problem, which manifested in their inability to recommend new caregivers effectively. The proposed CARES not only leverages DCN to analyze existing feature interactions but also constructs a knowledge graph, which considers the detailed relationships between caregivers and postpartum women, including their features. It employs a hybrid strategy of exploration and exploitation to recommend caregivers effectively. This approach helps the model overcome the cold start problem and generate recommendations that are useful and specific to the needs of postpartum women, improving the accuracy and effectiveness of suggesting caregivers.

5.1. Parameters setting and datasets

The proposed CARES adopts a structured approach to partitioning each dataset, allocating 70 % for training and 30 % for testing purposes. To optimize the learning process, the batch size is set at 256, which means the model looks at 256 examples at a time during

learning. The learning rate is set to 0.0002, guiding how much the model adjusts itself with new information each time it learns.

The proposed CARES is thoroughly tested using three different datasets: the Book-Crossing and Last.FM datasets, which is a publicly available dataset, and a Caregivers' dataset that is collected from YesMa, a facility that provides care for postpartum women. The Book-Crossing and Caregivers records explicit ratings range from 0–10, and 0–5, respectively. Table 3 shows the basic statistics of Caregivers and Book-Crossing datasets, including the number of users, items, interaction records (REC), triples in the knowledge graph, and the density (DEN), where DEN represents the density of the dataset. The smaller DEN is, the sparser the dataset is.

5.2. Simulation results

Fig. 5 illustrates the impact of varying numbers of clusters on the silhouette coefficient, an indicator of cluster validity. The number of clusters ranges from 1 to 30. There is an upward trend in the silhouette coefficient as the number of clusters increases from 1 to 5, reaching a peak that signifies the best clustering result. Beyond this peak, the coefficient decreases with the number of clusters. This is evidenced by decreased differentiation between clusters and less consistency within them. Consequently, the proposed CARES adopts a five-cluster configuration, capitalizing on the highest silhouette coefficient observed, thus ensuring the most effective data segmentation for subsequent analyses.

Fig. 6 presents the results of clustering postpartum women's features after dimensionality reduction through Principal Component Analysis (PCA). The scatter plot differentiates clusters based on color, where the number of clusters is set to five. The axes represent the first, second, and third principal components, with each data point's position representing its projection value on the respective component axis.

Fig. 7 illustrates the importance of various features within different clusters (Cluster 0 to Cluster 4) for recommendation. These features include the level of caregivers, cooking skills, lactation support, emergency response capacity, Maternal and Infant Care (MIC), as well as practice. Each cluster reveals the relative importance of different features in each cluster, as well as a comprehensive analysis across all clusters. As shown in Fig. 7, different clusters have distinct aspects of caregivers' services. Universally, MIC is the most focused area, followed by emergency response capacity and levels of caregivers. This information is invaluable for enhancing the recommendation accuracy of caregivers, facilitating the DCN model to match user needs more accurately with specific skills.

Fig. 8 further compares the loss of training and validation by varying training epochs, ranging from 1 to 100, for the proposed CARES utilizing the DCN model. As shown in Fig. 8, with the progression of training epochs, both the training loss (red line) and the validation loss (blue line) decrease, indicating a gradual improvement in model performance. The training loss remains consistently lower than the validation loss in each epoch. The reason is that the model has seen all the training data, allowing it to make more accurate predictions on these data, resulting in lower training loss. Conversely, the validation set contains unseen data by the model, and its predictive ability on this new data may not be as good as on the training set, leading to a relatively higher validation loss.

To further address overfitting, underfitting, and parameter sensitivity, the proposed CARES conducted additional experiments analyzing the impact of different learning rates and batch sizes. Specifically, Fig. 8(a) compares training and validation losses under Learning Rates (LR) of 0.0001, 0.0002, and 0.0003. The results demonstrate that a LR of 0.0002 achieves the best trade-off between convergence speed and generalization, while 0.0001 leads to slow convergence (potential underfitting) and 0.0003 causes fluctuations in loss (potential overfitting). Additionally, Fig. 8(b) analyzes the effect of different batch sizes (64, 128, 256, and 512) through F1-score boxplots. The results show that a batch size of 256 yields the highest median F1-score with the lowest variance, suggesting both accuracy and stability. These experimental extensions reinforce that the proposed CARES is trained under well-chosen and justified hyperparameters and that the training process is stable without signs of overfitting or underfitting.

To address the cold start problem, presented in the DCN model, Fig. 9 constructs a knowledge graph that maps the relationships between postpartum women and caregivers along with their respective characteristics. The colored nodes within the knowledge graph represent different attributes or categories, encompassing dimensions such as postpartum women, caregivers, levels, practice, lactation support, emergency response capabilities, MIC, housework skills, and age. The links between nodes illustrate the interconnections between the postpartum women and caregivers, as well as among their features. By analyzing the knowledge graph, the proposed CARES can more accurately match the needs of postpartum women with the expertise of caregivers, particularly for those newly joining, significantly increasing their likelihood of being recommended.

Figs. 10(a), (b), and (c) compare the proposed CARES with the Clustering & DCN and DCNv2 in terms of precision, recall, and F1-score by varying ratios of training data and number of recommended caregivers, respectively. The ratio of training data varies from 0.2 to 1, while the number of recommended caregivers varies from 2 to 10. As shown in Fig. 10, all mechanisms have the common trend that the precision, recall as well as F1-score increase with the ratio of training data. The reason is that with the training data increases, the model can learn more about the pairing characteristics between caregivers and postpartum women, thereby enhancing the precision, recall as well as F1-score of the recommendations. In comparison, the proposed CARES outperforms the Clustering & DCN and DCNv2 mechanisms in all evaluations. This superior performance is due to the proposed CARES's integration of DCN and a knowledge

Table 3
Basic statistics of Caregivers, Book-Crossing, and Last.FM datasets.

Datasets	User	Item	REC	TRI	DEN
Caregivers	5560	235	8210	2580	0.67 %
Book-Crossing	17,860	14,910	139,746	19,793	0.12 %
Last.FM	1872	3864	42,346	15,518	0.59 %

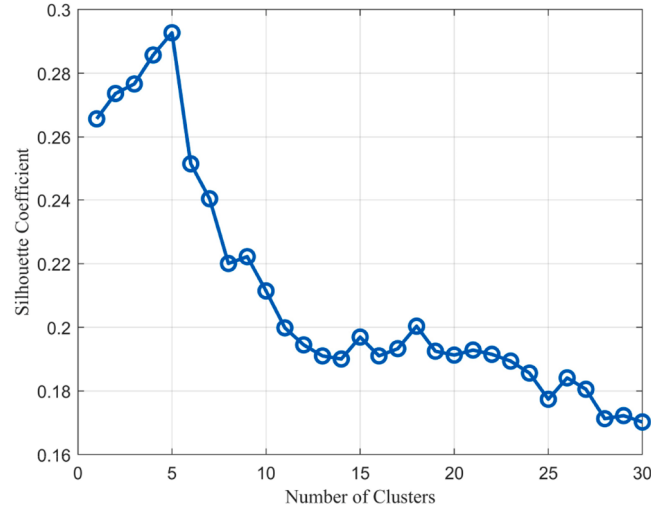


Fig. 5. Number of clusters determination using silhouette coefficient.

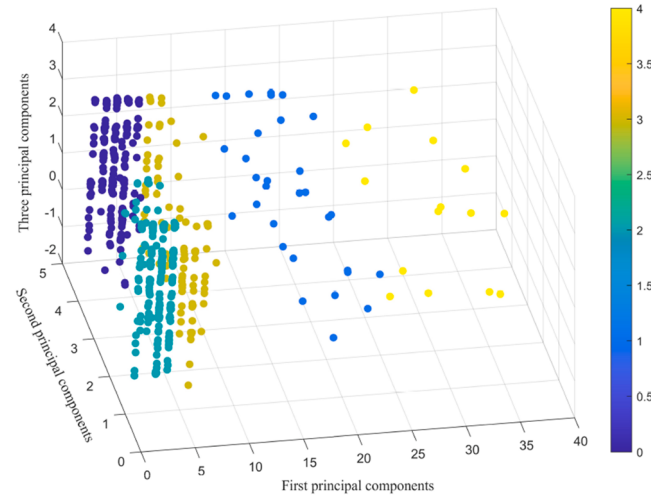


Fig. 6. Cluster distribution of postpartum women using PCA for dimensionality reduction.

graph, which effectively addresses the cold start problem and ensures reliable recommendations.

Fig. 11 compares different recommended numbers with $c_i \in \mathbb{C}_{DCN}$ and $c_i \in (\mathbb{C}_{KG} \setminus \mathbb{C}_{DCN})$ in terms of the recommended probability $P(c_i)$ by adjusting the exploration parameter ϵ , which varies from 0 to 1. The number of recommendations, denoted by r , ranges from 5 to 20. As shown in Fig. 11, $P(c_i)$ decreases with ϵ for $c_i \in \mathbb{C}_{DCN}$. This trend occurs because a rise in ϵ indicates a shift in the weight distribution from DCN and the knowledge graph. On the contrary, $P(c_i)$ is increasing with ϵ for $c_i \in (\mathbb{C}_{KG} \setminus \mathbb{C}_{DCN})$. The reason is that the proposed CARES starts to incorporate recommendations from the knowledge graph. Furthermore, as r increases, both the minimum value of $P(c_i)$ in $\mathbb{C}_{DCN}$ and the maximum value of $P(c_i)$ in $\mathbb{C}_{KG} \setminus \mathbb{C}_{DCN}$ decrease because both values are influenced by $\frac{1}{r}$.

The above experiments are conducted on a self-collected dataset. Table 4 further compares the proposed CARES with existing methods using two public datasets: Book-Crossing and Last.FM. Specifically, it reports the performance in terms of Hit Ratio (HR) and Normalized Discounted Cumulative Gain (NDCG) at ranks 10 and 20, which are HR@10, HR @20, NDCG@10, and NDCG@20, comparing the proposed CARES with baseline methods, including GMF [36], NGCF [37], LRGCN [38], RippleNet [39], KARN [40], KGAT [41] and KGARA [42]. The values of HR@10, HR @20, NDCG@10, and NDCG@20 can be calculated by Eqs. (11) and (14) of KGARA. In comparison, the proposed CARES outperforms all the other mechanisms. This occurs because the proposed CARES initially refines recommendations by clustering the dataset, enhancing its understanding and service capability for specific postpartum women groups. Second, the application of DCN enables the model to capture complex relationships within the data, significantly improving the precision of its recommendations. Furthermore, by integrating a knowledge graph, the model effectively resolves the cold start problem, ensuring reasonable recommendations even in scenarios of sparse data.

Fig. 12 illustrates the results of a 5-fold cross-validation conducted to comprehensively evaluate the generalization performance

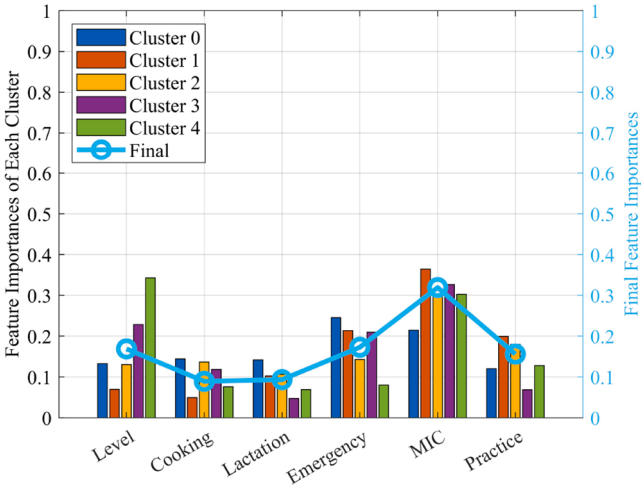


Fig. 7. The importance of various features within different clusters.

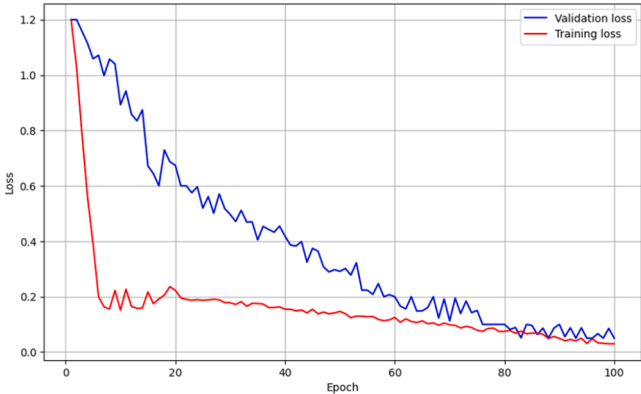


Fig. 8. Comparison with different LR and batch sizes.

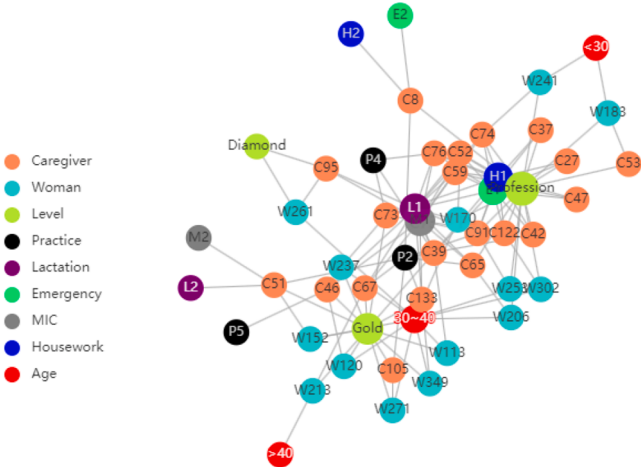


Fig. 9. Construction of knowledge graph.

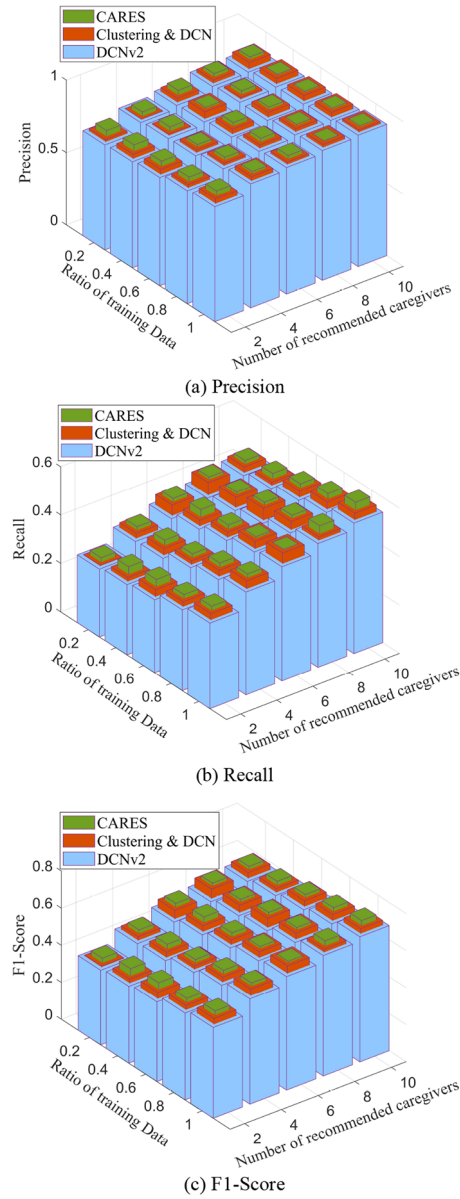


Fig. 10. Comparison of CARES with the Clustering & DCN and DCNv2 by varying ratio of training data and number of recommended caregivers.

and stability of the proposed CARES framework. In each fold, the model was trained on 80 % of the dataset and validated on the remaining 20 %, ensuring that every instance appeared exactly once in the validation set. This strategy helps mitigate the bias introduced by a single train-test split and provides a more reliable assessment of the model's robustness on unseen data.

To quantify the stability of the proposed CARE, let $\bar{M}$ and σ_M denote the mean and standard deviation of each evaluation metric across the k folds, respectively. Specifically, for each metric $M \in \{\text{F1}, \text{Precision}, \text{Recall}, \text{Accuracy}\}$, the values of $\bar{M}$ and σ_M can be computed using Eq. (49).

$$\bar{M} = \frac{1}{k} \sum_{i=1}^k M^i, \quad \sigma_M = \sqrt{\frac{1}{k} \sum_{i=1}^k (M^i - \bar{M})^2}, \quad (49)$$

where M^i denotes the value of metric M on the i -th fold, and $k = 5$ in this experiment. As shown in Fig. 12, the standard deviations for F1-score, Precision, Recall, and Accuracy are 0.012, 0.013, 0.015, and 0.010, respectively. These low variances confirm that the proposed CARES framework is not only accurate but also stable and generalizable across different data splits.

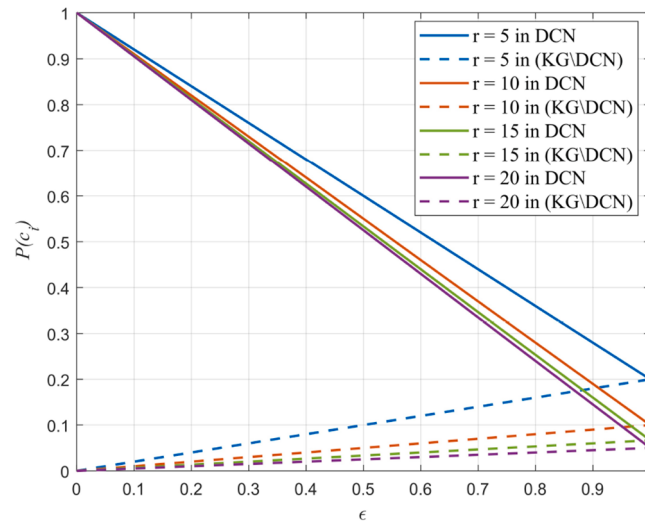


Fig. 11. The impact of $P(c_i)$ by varying exploration ϵ with different r .

6. Conclusions

This paper proposes a recommendation mechanism, called CARES. The proposed CARES employs XGBoost to identify important features and utilizes K -Means clustering alongside DCN to learn interactions of these features, aiming to achieve precise alignment between caregivers and postpartum women's requirements. Additionally, it constructs a knowledge graph to provide a robust framework for recommending new caregivers, especially in scenarios lacking historical data. The exploration and exploitation strategies further refine the recommendation process, striking a balance between diversifying caregiver options and maintaining high recommendation precision. The simulation results demonstrate that the proposed CARES outperforms the related work.

While the AI model provides objective, data-driven recommendations, it is not designed to replace human judgment. Instead, CARES should be regarded as a decision-support system that offers reliable suggestions to assist, but not override, human decision-makers such as care coordinators.

For future research, it is important to acknowledge that the proposed CARES primarily relies on user and item textual data, which may limit its ability to capture richer multimodal information and complex user-item interactions. To address these limitations, future work can explore how to integrate a wider range of data types, such as images, videos, and audio, into recommendation systems to enhance their richness and accuracy. Additionally, future research will also investigate cross-domain data learning, aiming to improve recommendation generalization capability and effectiveness by enabling knowledge transfer and utilization.

CRediT authorship contribution statement

Qiaoyun Zhang: Writing – review & editing, Writing – original draft, Validation, Methodology, Formal analysis, Conceptualization. **Sze-Han Wang:** Resources, Investigation, Conceptualization. **Chung-Chih Lin:** Writing – review & editing, Validation. **Chih-Yung Chang:** Writing – review & editing, Writing – original draft, Methodology, Formal analysis, Conceptualization. **Diptendu Sinha Roy:** Writing – review & editing.

Table 4

Comparison of H@10, H@20, N@10, and N@20 with different mechanisms on Book-Crossing and Last.FM datasets.

Models	Book-Crossing				Last.FM			
	HR@10	HR@20	NDCG@10	NDCG@20	HR@10	HR@20	NDCG@10	NDCG@20
GMF [36]	0.301	0.412	0.211	0.283	0.481	0.561	0.330	0.341
NGCF [37]	0.349	0.501	0.258	0.323	0.580	0.619	0.362	0.376
LRGCN [38]	0.415	0.638	0.318	0.379	0.591	0.648	0.378	0.395
RippleNet [39]	0.358	0.547	0.271	0.359	0.511	0.624	0.371	0.382
KARN [40]	0.369	0.521	0.279	0.346	0.561	0.604	0.354	0.367
KGAT [41]	0.380	0.582	0.284	0.364	0.638	0.678	0.396	0.418
KGAR [42]	0.402	0.603	0.338	0.382	0.625	0.663	0.392	0.432
The proposed CARES	0.425	0.65	0.352	0.393	0.640	0.681	0.408	0.435
Improv.	0.01	0.012	0.014	0.011	0.002	0.003	0.012	0.003

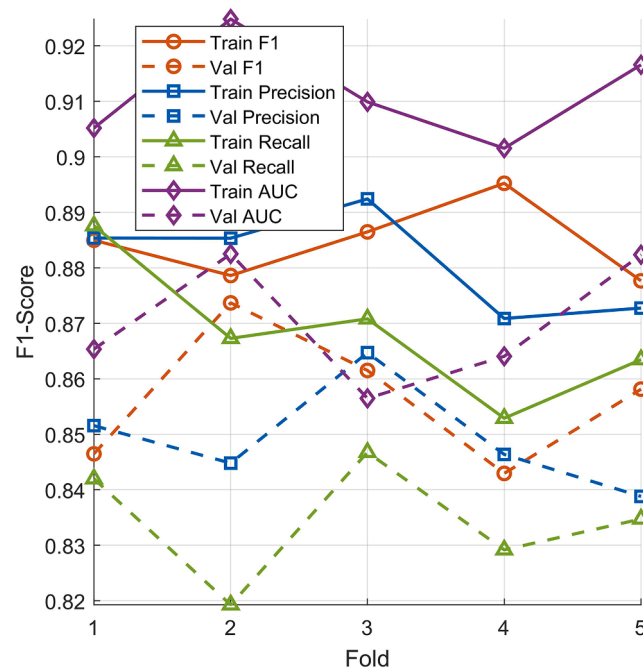


Fig. 12. Cross-Validation Results Across Folds.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

This work was supported in part by the Research Start-up Fund Project of Chuzhou University under Grant No. 2025qd12, Higher Education Research Program Project under Grant Nos. 2022AH010067 and 2024AH051412, Anhui Provincial Natural Science Foundation under Grant Nos. 2408085MF177 and 2508085MF174, and Anhui Province's University Research Project under Grant No. 2023AH040216.

Data availability

The datasets generated during the current study are not publicly available but are available from the corresponding author.

References

- [1] S.S. Khanal, P.W.C. Prasad, A. Alsadoon, A. Maag, A systematic review: machine learning based recommendation systems for e-learning, *Educ. Inform. Technol.* 25 (4) (2020) 2635–2664.
- [2] L. Wu, X. He, X. Wang, K. Zhang, M. Wang, A survey on accuracy-oriented neural recommendation: from collaborative filtering to information-rich recommendation, *IEEE Trans. Knowl. Data Eng.* 35 (5) (2022) 4425–4445.
- [3] X. Chen, W. Liang, J. Xu, C. Wang, K.C. Li, M. Qiu, An efficient service recommendation algorithm for cyber-physical-social systems, *IEEE Trans. Netw. Sci. Eng.* 9 (6) (2021) 3847–3859.
- [4] M.F. Aljunid, M.K. Hooshmand, W.A. Ali, A.M. Shetty, S.Q. Alzoubah, A collaborative filtering recommender systems: Survey, *Neurocomputing* 617 (2025) 128718.
- [5] Y. Sun, Q. Liu, Collaborative filtering recommendation based on K-nearest neighbor and non-negative matrix factorization algorithm, *J. Supercomput.* 81 (2025) 79.
- [6] S. Hwang, E. Park, Movie recommendation systems using actor-based matrix computations in South Korea, *IEEE Trans. Computat. Soc. Syst.* 9 (5) (2022) 1387–1393.
- [7] U. Javed, K. Shaukat, I.A. Hameed, F. Iqbal, T.M. Alam, S. Luo, A review of content-based and context-based recommendation systems, *Int. J. Emerg. Technol. Learn. (IJET)* 16 (3) (2021) 274–306.
- [8] P. Tenti, J. Thomas, R. Peñaloza, G. Pasi, ContReviews: A content-based recommendation system for updating Living Evidences in health care, *Knowl.-Bas. Syst.* 311 (2025) 112981.
- [9] Y. Chen, Y. Fang, R. Cheng, Y. Li, X. Chen, J. Zhang, Exploring communities in large profiled graphs, *IEEE Trans. Knowl. Data Eng.* 31 (8) (2018) 1624–1629.
- [10] X. Cai, W. Guo, M. Zhao, Z. Cui, J. Chen, A knowledge graph-based many-objective model for explainable social recommendation, *IEEE Trans. Computat. Soc. Syst.* 10 (6) (2023) 3021–3030.
- [11] J.C. Zhang, A.M. Zain, K.Q. Zhou, X. Chen, R.M. Zhang, A review of recommender systems based on knowledge graph embedding, *Expert Syst. Appl.* 250 (2024) 123876.

- [12] M.O. Kaya, M. Ozdem, R. Das, A new hybrid approach combining GCN and LSTM for real-time anomaly detection from dynamic computer network data, *Comput. Netw.* 268 (2025) 111372.
- [13] Z. Bi, T. Zhang, P. Zhou, Y. Li, Knowledge transfer for out-of-knowledge-base entities: Improving graph-neural-network-based embedding using convolutional layers, *IEEE Access* 8 (2020) 159039–159049.
- [14] Y. Wang, W. Ma, M. Zhang, Y. Liu, S. Ma, A survey on the fairness of recommender systems, *ACM Transact. Informat. Syst.* 41 (3) (2023) 1–43.
- [15] L. Wu, P. Sun, R. Hong, Y. Ge, M. Wang, Collaborative neural social recommendation, *IEEE Transact. Syst., Man, Cybernet.* 51 (1) (2021) 464–476.
- [16] B. Bai, Y. Fan, W. Tan, J. Zhang, DLTSR: a deep learning framework for recommendation of long-tail web services, *IEEE Transact. Serv. Comput.* 13 (1) (2020) 73–85.
- [17] J. Wang, Y. Shi, H. Yu, Z. Yan, H. Li, Z. Chen, A Novel KG-based Recommendation Model Via Relation-Aware Attentional GCN, 275, *Knowledge-Based Systems*, 2023 110702.
- [18] Q. Guo, F. Zhuang, C. Qin, H. Zhu, X. Xie, H. Xiong, A survey on knowledge graph-based recommender systems, *IEEE Trans. Knowl. Data Eng.* 34 (8) (2022) 3549–3568.
- [19] G. Chen, Y. Zheng, N. Li, Y. Li, Y. Qin, J. Piao, Y. Quan, J. Chang, D. Jin, X. He, Y. Li, A survey of graph neural networks for recommender systems: challenges, methods, and directions, *ACM Transact. Recommend. Syst.* 1 (1) (2023) 1–51.
- [20] Y. Liu, S. Yang, Y. Xu, C. Miao, M. Wu, J. Zhang, Contextualized graph attention network for recommendation with item knowledge graph, *IEEE Trans. Knowl. Data Eng.* 35 (1) (2023) 181–195.
- [21] Y. Wang, X. Huang, J. Ma, Q. Jin, LASGRec: a personalized recommender based on learnable attribute sampling and graph neural network, *IEEE Trans. Comput. Soc. Syst.* 11 (2) (2024) 2930–2939.
- [22] J. Liu, W. Huang, T. Li, S. Ji, J. Zhang, Cross-domain knowledge graph chiasmal embedding for multi-domain item-item recommendation, *IEEE Trans. Knowl. Data Eng.* 35 (5) (2023) 4621–4633.
- [23] Z. Huang, Y. Liu, C. Zhan, C. Lin, W. Cai, Y. Chen, A novel group recommendation model with two-stage deep learning, *IEEE Transac. Syst., Man, Cybernet.* 52 (9) (2022) 5853–5864.
- [24] C. Liu, Y. Li, H. Lin, C. Zhang, GNNRec: gated graph neural network for session-based social recommendation model, *J. Intell. Inf. Syst.* 60 (1) (2023) 137–156.
- [25] Q. Wang, Y. Wei, J. Yin, J. Wu, X. Song, L. Nie, Dualgnn: dual graph neural network for multimedia recommendation, *IEEE Trans. Multimedia* 25 (2021) 1074–1084.
- [26] S. Kim, S. Yun, J. Lee, G. Chang, W. Roh, D. Sohn, J.T. Lee, H. Park, S. Kim, Self-supervised multimodal graph convolutional network for collaborative filtering, *Inf. Sci. (Ny)* 653 (2024) 119760.
- [27] Y. Wang, D. Ma, J. Ma, Q. Jin, HGCR: a heterogeneous graph-enhanced interactive course recommendation scheme for online learning, *IEEE Transac. Learn. Technol.* 17 (2023) 364–374.
- [28] K. Yao, J. Liang, J. Liang, M. Li, F. Cao, Multi-view graph convolutional networks with attention mechanism, *Artif. Intell.* 307 (2022) 103708.
- [29] X. Li, Y. Tian, B. Dong, S. Ji, MD-GCCF: multi-view deep graph contrastive learning for collaborative filtering, *Neurocomputing* 590 (2024) 127756.
- [30] Y. Han, J. Kim, D. Enke, A machine learning trading system for the stock market based on N-period Min-max labeling using XGBoost, *Expert Syst. Appl.* 211 (2023) 118581.
- [31] R. Wang, R. Shivanna, D. Cheng, S. Jain, D. Lin, L. Hong, E. Chi, Dcn v2: improved deep & cross network and practical lessons for web-scale learning to rank systems, in: *The Web Conference*, Ljubljana, Slovenia, 2021.
- [32] J. Castro, J. Lu, G. Zhang, Y. Dong, L. Martínez, Opinion dynamics-based group recommender systems, *IEEE Transac. Syst., Man, Cybernet.* 48 (12) (2018) 2394–2406.
- [33] D. Qin, X. Zhou, L. Chen, G. Huang, Y. Zhang, Dynamic connection-based social group recommendation, *IEEE Trans. Knowl. Data Eng.* 32 (3) (2020) 453–467.
- [34] B. Wang, S. Song, S. Liu, X. Deng, MultiFDF: multi-community clustering for fairness-aware recommendation, *IEEE Transac. Computat. Soc. Syst.* 10 (6) (2023) 2959–2970.
- [35] T. Chen, C. Guestrin, Xgboost: a scalable tree boosting system, in: *The 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, California, USA, 2016.
- [36] X. He, L. Liao, H. Zhang, L. Nie, X. Hu, Neural collaborative filtering, in: *The 26th International Conference on World Wide Web*, Perth, Australia, 2017.
- [37] X. Wang, X. He, M. Wang, F. Feng, T.S. Chua, Neural graph collaborative filtering, in: *The 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, Paris, France, 2019.
- [38] L. Chen, L. Wu, R. Hong, K. Zhang, M. Wang, Revisiting graph based collaborative filtering: a linear residual graph convolutional network approach, in: *The 34th AAAI Conference on Artificial Intelligence*, New York, USA, 2020.
- [39] H. Wang, F. Zhang, J. Wang, M. Zhao, W. Li, X. Xie, Ripplenet: propagating user preferences on the knowledge graph for recommender systems, in: *The 27th ACM International Conference on Information and Knowledge Management*, New York, USA, 2018.
- [40] Q. Zhu, X. Zhou, J. Wu, J. Tan, L. Guo, A knowledge-aware attentional reasoning network for recommendation, in: *The 34th AAAI Conference on Artificial Intelligence*, New York, USA, 2020.
- [41] X. Wang, X. He, Y. Cao, M. Liu, T. Chua, Kgat: knowledge graph attention network for recommendation, in: *The 25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, New York, USA, 2019.
- [42] Y. Zhang, M. Yuan, C. Zhao, M. Chen, X. Liu, Aggregating knowledge-aware graph neural network and adaptive relational attention for recommendation, *Appl. Intellig.* 52 (15) (2022) 17941–17953.