

MetaEDL: Meta Evidential Learning For Uncertainty-Aware Cold-Start Recommendations

Krishna Prasad Neupane, Ervine Zheng, Qi Yu
Rochester Institute of Technology
 {kpn3569, mxz5733, qi.yu}@rit.edu

Abstract—Recommender systems have been widely used to predict users’ interests and filter information from a large number of candidate items. However, accurately capturing the interests of users having limited interactions with a system remains a long-lasting challenge. Furthermore, existing recommender systems primarily focus on predicting user preferences without quantifying the prediction uncertainty. Uncertainty can help to quantify the model confidence when making a recommendation where low model confidence could serve as a more accurate indicator of a user’s cold-start level than simply using the number of interactions. We present a novel recommendation model that seamlessly integrates a meta-learning module with an evidential learning approach. The former module generalizes meta knowledge to tackle cold-start recommendations by exploiting fast adaptation. The latter quantifies both aleatoric and epistemic uncertainty without performing expensive posterior inference. Evidential learning achieves this by placing evidential priors and treating the output of the meta-learning module as evidence-based pseudo counts and learns a function to directly predict the evidence of a target interaction. Experiments on four benchmark datasets justify that our proposed model captures the uncertainty of users and demonstrates its superior performance over the state-of-the-art recommendation models.

I. INTRODUCTION

Recommender systems exploit data mining techniques and prediction algorithms to predict users’ interest in products, services, and information among a large number of available items [1]. Commonly used approaches can be generally categorized as collaborative filtering-based, content-based, and hybrid systems. Collaborative filtering methods recommend items to the target users based on the similar taste of existing users [2]. These methods mostly suffer from data sparsity that leads to the cold-start problems (*i.e.*, inability to handle new users and/or items with limited interactions). Content-based methods [3] address this issue by utilizing users’ demographic information (*e.g.*, age, gender, and location) and item content (*e.g.*, genres, directors, and actors of movies). While various extra information is available for items, acquiring users’ personalized information is usually difficult due to privacy issues. Hybrid models combine the benefits of both collaborative and content-based systems but remain less effective to the cold-start users/items.

Several recent works have attempted to address the cold-start problem in recommender systems through meta-learning [4], [5]. In particular, meta-learning models the cold-start recommendation as a few-shot learning problem. By arranging existing users’ item interaction history as the training tasks, it learns a global meta-model that can adapt to users/items with

Table I: Epistemic uncertainty and RMSE loss for two users from Movielens 1M based on the number of interactions.

UserID	Interactions	Epistemic	RMSE
4515	24	0.5133	1.0345
4575	164	0.6812	1.7038

limited interactions with improved recommendation accuracy. Most existing methods, including meta-learning models, use the number of interactions as the primary factor to identify cold-start users. However, they ignore the nature of the interactions as not all the interactions are equally important for a recommender system to construct an accurate (latent) profile for users to provide effective recommendations. As certain interactions can bring much higher value to the system than others, it is essential to consider the *value* of the interactions to most effectively tackle cold-start recommendations and uncertainty provides an effective means to quantify such value (as evidenced by our experimental results in Section V-D).

Table I shows two example users from the Movielens dataset with significantly different numbers of interactions. As can be seen, the second user is much more active than the first one, who has much fewer interactions and may be regarded as cold-start. However, more interactions may not necessarily lead to a more accurate recommendation result, which is evidenced by a higher root mean squared error (RMSE) for the second user. In fact, the larger recommendation error is also reflected by a higher model (or epistemic) uncertainty. This example, along with more illustrative examples provided in our experiment section, helps confirm the distinct values of different interactions further. It also implies the important role of using *uncertainty* to quantify the model confidence when making a recommendation that could indicate the cold-start level of a user (*i.e.*, how well the system knows the user).

In general, a recommender system’s prediction is very sensitive to the observed user-item interactions, especially when they are limited. Hence, a precise and calibrated uncertainty estimation is useful for interpreting the model confidence in cold-start recommendations. There are two common types of uncertainty: *aleatoric* that captures the uncertainty introduced by the noises in the data and *epistemic* that captures the model uncertainty due to lack of understanding of the data [6]. Aleatoric uncertainty can be directly estimated from data and Bayesian models offer a natural way to capture model uncertainty. Hence, Bayesian neural networks have been commonly used to estimate the epistemic uncertainty of

deep learning (DL) models. However, Bayesian DL models usually conduct posterior inference through Monte Carlo (MC) sampling, which poses a very high computational cost due to a large number of parameters in a DL model and their complex dependencies. Consequently, directly extending the current DL-based recommender systems through Bayesian modeling will prevent them from scaling to a large user-item space.

To address the above key challenges, we propose a meta evidential learning model, referred to as *MetaEDL*, to provide uncertainty-aware cold-start recommendations. By integrating a meta-learning module with evidential learning, MetaEDL is able to leverage all existing users' historical interactions to learn a global model that can easily and accurately adapt to cold-start users with limited interactions. Furthermore, we construct a hierarchical Bayesian model that provides a generative process to model the likelihood of the user-item interactions. Instead of performing an expensive posterior inference, evidential learning is adopted to directly predict the hyper-parameters of the posterior distributions of the parameters in the likelihood function, using a non-Bayesian deep neural network. These predicted hyperparameters have a natural interpretation as *pseudo counts*, which can serve as evidence to quantify the model confidence for its recommendations. The main contribution of this paper is fourfold:

- A novel recommendation model that integrates meta-learning and evidential learning to provide uncertainty-aware cold-start recommendations.
- Bayesian posterior inference through evidential learning to ensure good efficiency that allows a recommender system to scale to a large user-item space.
- Using pseudo-count-based evidence that provides a deeper insight to understand the value of different interactions that is instrumental to provide effective uncertainty-aware recommendations to cold-start users and accurately capture their latent profile using limited but highly 'informative' interactions.
- An integrated end-to-end training process that optimizes the embeddings and meta evidential learning modules.

We conduct extensive experiments over four real-world datasets and compare with state-of-the-art models to demonstrate the effectiveness of the proposed MetaEDL model.

II. RELATED WORK

Matrix Factorization. Matrix factorization (MF) based models utilize user and item latent factors [7] to make predictions. Incorporating implicit information of users and items, the basic MF method is extended into SVD++ [8]. Variations of these models have also been applied to dynamic settings, including timeSVD++ [9], dynamic Poisson Factorization (DPF) [10], and collaborative Kalman Filter (CKF) [11].

Deep Learning Models. In recent years, deep learning-based recommender systems [12], [13] are successfully proposed due to their effective improvement over traditional methods. DeepFM [13] which adapts deep learning with traditional FM, integrates the power of deep learning and factorization

machines to learn low- and high-order feature interactions simultaneously from the input. Cheng et al. [12] propose to jointly train wide linear models and deep neural networks to combine the benefits of memorization and generalization.

Graph-Based Models. Another popular line of recommendation systems is graph-based models. A graph captures high-order user-item interactions through an iterative process to provide effective recommendations [14]. Users and items are represented as a bipartite graph in [15] and links are predicted to provide recommendations.

Meta-learning. Meta-learning [16] is a few-shot learning approach that learns from similar tasks and can generalize quickly and efficiently for the unseen new tasks. The meta-learning strategy introduced in [4] addresses the cold-start problem in item recommendation. Similarly, recent works take advantage of users' and items' side-information to generate user and item embeddings and fed those embeddings to the meta-learning to estimate user preferences [5].

Uncertainty in Recommender Systems. All the above methods primarily focus on personalized recommendation but lack in handling uncertainty. Recently, Gaussian embedding-based recommendation [17] attempts to capture user and item uncertainty but does not measure model uncertainty. Our proposed method not only provides an effective recommendation but also measures both data and model uncertainty utilizing the evidential learning approach [18].

III. PROBLEM FORMULATION

For a recommendation model, input data is represented as $\{U, I\}$, where U is the user set and I is the item set. We perform recommendation and uncertainty quantification for each user with a recommendation function as:

$$f_{\theta_u, E_u, E_i}(u, i) = \{\gamma_{(u,i)}, \nu_{(u,i)}, \alpha_{(u,i)}, \beta_{(u,i)}\} \quad \forall u \in U, i \in I \quad (1)$$

where $\gamma_{(u,i)}$ is the recommended score for item i assigned by user u , $\nu_{(u,i)}$, and $\alpha_{(u,i)}$ are the model evidence, and $\beta_{(u,i)}$ is a total uncertainty coming from both pseudo and actual data samples (more details are provided along with the meta evidential learning module) for user u on item i and all parameters are scalar quantities, θ_u is user specific model parameter; E_u , and E_i are user and item embedding module parameters. The goal of a recommender system is to predict the scores with confidence so that it can accurately capture a user's true preference on items in belief that the recommended items are likely to be adopted by the user.

We formulate recommendations as a few-shot regression problem in the meta-learning setting. Users are *dynamically* partitioned into *meta-train* and *meta-test* sets. The meta-train user set includes users with sufficient interactions, while the meta-test user set includes cold-start users who have only a few interactions. We consider a distribution over tasks $P(\mathcal{T})$, and each user is represented as a few-shot regression task $\mathcal{T}_u$ sampled from the given task distribution. In general, a task includes a *support set* $\mathcal{S}_u$ and a *query set* $\mathcal{Q}_u$. The support set includes a user's K interactions where K is interpreted as

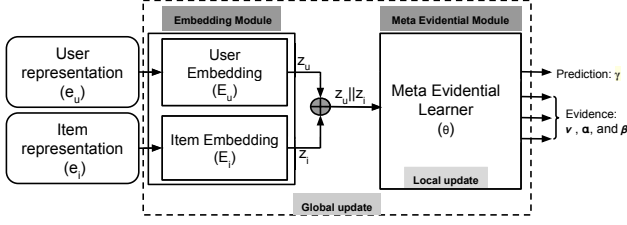


Figure 1: Overview of the proposed model.

the number of shots (i.e., interactions). The query set includes the rest interactions of this user.

$$\begin{aligned} \mathcal{T}_u \sim P(\mathcal{T}) : \quad S_u &= \{(u, i_j), r_{(u, i_j)}\}_{j=1}^K \\ Q_u &= \{(u, i_j), r_{(u, i_j)}\}_{j=K+1}^N \end{aligned} \quad (2)$$

where N is the number of items a user interacted, and $r_{(u, i_j)}$ represents label (i.e., rating or count) from user u to item i_j .

We adopt episodic training [19], where the training task mimics the test task for efficient meta-learning. The support set $\mathcal{S}$ in each episode works as the labeled training set on which the model is trained to minimize the loss over the query set $\mathcal{Q}$. This training process is iteratively carried out episode by episode until convergence.

IV. META EVIDENTIAL LEARNING

The proposed model consists of two major components: an embedding module and a meta evidential learning module, as shown in Figure 1. The embedding module generates user and item embeddings and is forwarded to the evidential meta-learning module, where prediction and model evidence are produced as a final output.

A. Embedding Module

We represent user u , and item i in one hot encoding considering unique user and item IDs: $e_u \in \mathcal{R}^n$ where n is the total number of users and $e_i \in \mathcal{R}^m$ where m is the total number of items respectively. This one hot vector is then transformed using the embedding matrix E_u for user and E_i for item in d -dimension: $z_u = E_u e_u$ and $z_i = E_i e_i$. The embedding matrix is optimized along with the model training process. We use gradient descent to update both users and item embedding matrices:

$$\begin{aligned} E_u &= E_u - \xi \nabla_{E_u} (\mathcal{L}_{\mathcal{T}_u} [f_{\theta_u, E_u, E_i}]) \\ E_i &= E_i - \xi \nabla_{E_i} (\mathcal{L}_{\mathcal{T}_u} [f_{\theta_u, E_u, E_i}]) \end{aligned} \quad (3)$$

where ξ is the step size, and $\mathcal{L}_{\mathcal{T}_u} [f_{\theta_u, E_u, E_i}]$ is an end-to-end meta evidential user-specific loss and detail is given in Equation (20).

B. Meta Evidential Learning Module

We formulate the meta-learning module as a non-Bayesian neural network to estimate a target interaction score and its associated evidence to learn both aleatoric and epistemic uncertainty. We accomplish this by placing evidential priors over the original Gaussian likelihood function and training the

neural network to infer the hyperparameters of the evidential distribution similar to [18]. The key intuition of employing evidential learning in recommender systems is that it allows us to assign evidence to the predicted interaction, where the evidence can be used to formulate the prediction score while capturing the model confidence.

A hierarchical Bayesian model. The recommendation problem is set up in such a way that the target (e.g., rating or count), y_n , is drawn i.i.d. from a Gaussian distribution with unknown mean and variance (μ, σ^2) . Model evidence can be introduced by further placing a prior distribution on (μ, σ^2) , leading to a hierarchical Bayesian model. To ensure conjugacy, we choose a Gaussian prior on the unknown mean and an Inverse-Gamma prior on the unknown variance:

$$p(y_n | \mu, \sigma^2) = \mathcal{N}(\mu, \sigma^2) \quad (4)$$

$$p(\mu | \gamma, \sigma^2 \nu^{-1}) = \mathcal{N}(\gamma, \sigma^2 \nu^{-1}) \quad (5)$$

$$p(\sigma^2 | \alpha, \beta) = \text{Inv-Gamma}(\alpha, \beta) \quad (6)$$

where $\text{Inv-Gamma}(z | \alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} \left(\frac{1}{z}\right)^{\alpha+1} \exp(-\frac{\beta}{z})$ with $\Gamma(\cdot)$ being a gamma function; γ, ν, α , and β are parameters of the corresponding prior distributions.

Interpreting hyper-parameters. Besides serving as the parameters of the corresponding prior distributions in the hierarchical Bayesian model, the hyper-parameters $(\gamma, \nu, \alpha, \beta)$ offer very intuitive meanings, which set the stage to use them in the proposed evidential learning model. The best way to show this is to couple these prior distributions with a set of actual observations, i.e., $\mathbf{y} = (y_1, \dots, y_N)^\top$. Given the Gaussian likelihood in (4), we compute joint posterior distribution $p(\mu, \sigma^2 | \mathbf{y})$ factorized as $p(\mu | \mathbf{y}, \sigma^2) p(\sigma^2 | \mathbf{y})$. We first derive the conditional posterior of μ :

$$p(\mu | \mathbf{y}, \sigma^2) = \mathcal{N}(\gamma_N, \sigma_N^2) \quad (7)$$

$$\gamma_N = \frac{\nu}{\nu + N} \gamma + \frac{1}{N + \nu} \sum_{n=1}^N y_n \quad (8)$$

$$\sigma_N^2 = \frac{\sigma^2}{\nu + N} = \frac{\sigma^2}{\nu_N} \quad (9)$$

where $\nu_N = \nu + N$. From (8), we can see that the posterior mean is the convex combination of the prior mean γ and the maximum likelihood estimation of the mean, given by $\frac{1}{N} \sum_{n=1}^N y_n$. Similarly, the variance in the posterior distribution is ν_N times smaller than the prior variance. As a result, ν can be interpreted as the ‘effective’ prior observations for the prior mean γ . We continue to derive the posterior of σ^2 :

$$p(\sigma^2 | \mathbf{y}) = \text{Inv-Gamma}(\alpha_N, \beta_N) \quad (10)$$

$$\alpha_N = \alpha + \frac{N}{2} \quad (11)$$

$$\beta_N = \beta + \frac{1}{2} \sum_{n=1}^N (y_n - \mu)^2 \quad (12)$$

First, (11) shows that after observing N data samples, the prior parameter α is increased by $\frac{N}{2}$ to reach the posterior parameter α_N . This has the effect of treating the prior hyper-parameter

as 2α ‘effective’ prior observations of ‘pseudo’ data samples. Similarly, by multiplying both sides of (12) by 2, we have

$$2\beta_N = 2\beta + \sum_{n=1}^N (y_n - \mu)^2 = 2\beta + N\sigma_{ML}^2 \quad (13)$$

where σ_{ML}^2 denotes the maximum likelihood estimator of the variance arising from data samples $(y_1, \dots, y_N)$. From this, hyper-parameter β can be interpreted as the 2β total ‘prior’ variance arising from the corresponding 2α ‘effective’ prior observations of ‘pseudo’ data samples.

Mapping hyper-parameters to evidence-based uncertainty.

The above analysis provides an intuitive interpretation of key hyper-parameters introduced along with the prior distributions in the hierarchical Bayesian model. This will help to understand their key roles in defining different types of uncertainties introduced next. In particular, since both ν and α are essentially the ‘effective’ prior observations, it is natural to treat their posterior counterpart ν_N and α_N as the *evidence* to support (or suspect) a prediction given training samples $(y_1, \dots, y_N)$. Furthermore, β_N can be treated as the total uncertainty that combines two sources of uncertainty: the prior variance β from the pseudo samples and the variance σ_{ML}^2 of the actually observed data samples.

We start defining the model prediction and uncertainty from the data (referred to as aleatoric uncertainty) as

$$\text{Prediction: } \mathbb{E}[\mu] = \gamma_N, \quad \text{Aleatoric: } \mathbb{E}[\sigma^2] = \frac{\beta_N}{\alpha_N - 1} \quad (14)$$

where both can be directly obtained as the mean from the corresponding Gaussian and Inv-Gamma posteriors defined in (7) and (10), respectively. It is interesting to see that the uncertainty from the data is proportion to the total uncertainty β_N and decreases with (both pseudo and actual) observations. Next, we quantify the uncertainty of the model prediction (referred to as epistemic uncertainty) by showing an important relationship with the aleatoric uncertainty through the following theorem.

Theorem 1. *Given a hierarchical Bayesian model as specified by (4)-(6) and a set of observed (training) data samples $(y_1, \dots, y_N)$, the epistemic uncertainty that quantifies the variance of the posterior mean (as the model prediction), given by $\text{Var}[\mu]$, is $\frac{1}{\nu_N}$ times of the aleatoric uncertainty:*

$$\text{Var}[\mu] = \frac{E[\sigma^2]}{\nu_N} = \frac{\beta_N}{\nu_N(\alpha_N - 1)} \quad (15)$$

where ν_N is defined in (9).

Proof. First, note that we cannot directly use the variance given by the posterior distribution in (7) as it is still conditioned on σ^2 . Since $\text{Var}[\mu]$ is defined on the marginal posterior $p(\mu|y)$, we need to further marginalize σ^2 , which gives

$$\begin{aligned} \text{Var}[\mu] &= \int \int [\mu^2 p(\mu|\sigma^2) - (\mathbb{E}[\mu])^2] p(\sigma^2) d\mu d\sigma^2 \\ &= \gamma_N^2 - (\mathbb{E}[\mu])^2 + \int \frac{\sigma^2}{\nu} p(\sigma^2) d\sigma^2 = \frac{\beta_N}{\nu_N(\alpha_N - 1)} \end{aligned} \quad (16)$$

where we omit the dependency on y to keep the notation uncluttered. $\square$

Now we define a loss function that is formed through the evidence and total uncertainty parameters. Given an observed score $r_{(u,i)}$ resulted from an interaction between user u and item i , we marginalize the likelihood parameters (μ, σ^2) , which gives the marginal likelihood function

$$\begin{aligned} &p(r_{(u,i)} | \gamma_{(u,i)}, \nu_{(u,i)}, \alpha_{(u,i)}, \beta_{(u,i)}) \\ &= \int \int \mathcal{N}(\mu, \sigma^2) \mathcal{N}(\gamma_{(u,i)}, \sigma^2 \nu^{-1}) \text{IG}(\alpha_{(u,i)}, \beta_{(u,i)}) d\mu d\sigma^2 \\ &= \text{St} \left(r_{(u,i)}; \gamma_{(u,i)}, \frac{\beta_{(u,i)}(1 + \nu_{(u,i)})}{\nu_{(u,i)} \alpha_{(u,i)}}, 2\alpha_{(u,i)} \right) \end{aligned} \quad (17)$$

where IG is short for Inv-Gamma and St(.) is a student-t distribution on target variable $r_{(u,i)}$ with respective location and scale parameters.

We adopt an evidential loss, which utilizes the above marginal likelihood while computing the predicted loss. This includes the negative log-likelihood ($\mathcal{L}^{NLL}[f_{\theta_u, E_u, E_i}]$) to maximize the marginal likelihood and an evidential regularizer ($\mathcal{L}^R[f_{\theta_u, E_u, E_i}]$) to impose a high penalty on the predicted error with a low uncertainty (or a large confidence). We first formulate the negative log-likelihood, given by

$$\mathcal{L}^{NLL}[f_{\theta_u, E_u, E_i}] = -\log(p(r_{(u,i)} | \gamma_{(u,i)}, \nu_{(u,i)}, \alpha_{(u,i)}, \beta_{(u,i)})) \quad (18)$$

We formalize our own evidence regularizer, which considers epistemic uncertainty to penalize confidently predicted errors. We multiply the predicted error with the inverse epistemic uncertainty that scales up the error when the predicted evidence is high causing high inverse epistemic uncertainty and vice-versa. Conversely, it will be less penalized if the prediction is close to the target score:

$$\mathcal{L}^R[f_{\theta_u, E_u, E_i}] = |r_{(u,i)} - \gamma_{(u,i)}| \cdot \left(\frac{\nu_{(u,i)}(\alpha_{(u,i)} - 1)}{\beta_{(u,i)}} \right) \quad (19)$$

In the meta evidential setting, we compute the loss for a specific user u , which can be formulated with user evidential loss as:

$$\begin{aligned} \mathcal{L}_{\mathcal{T}_u}[f_{\theta_u, E_u, E_i}] &= \sum_{u, i \sim \mathcal{T}_u} \mathcal{L}[f_{\theta_u, E_u, E_i}(u, i)], \\ \mathcal{L}[f_{\theta_u, E_u, E_i}(u, i)] &= \mathcal{L}^{NLL}[f_{\theta_u, E_u, E_i}(u, i)] + \\ &\quad \lambda_1 \mathcal{L}^R[f_{\theta_u, E_u, E_i}(u, i)] \end{aligned} \quad (20)$$

where λ_1 is a regularization parameter.

The total loss is formed by aggregating all users in the meta-train set, regularized by the L_2 norm of key model parameters. Let θ_u and θ denote the local (i.e., user-specific) and global parameters of the meta evidential learner. Training the meta evidential learning as a recommendation model can be formulated as the following optimization problem:

$$\begin{aligned} \min_{\theta} \quad &\sum_{\mathcal{T}_u \sim p(\mathcal{T})} \mathcal{L}_{\mathcal{T}_u}[f_{\theta_u, E_u, E_i}] + \frac{\lambda_2}{2} \|\theta\|_2^2, \\ &\theta_u = \theta - \eta \nabla_{\theta} \mathcal{L}_{\mathcal{T}_u}(f_{\theta_u, E_u, E_i}) \end{aligned} \quad (21)$$

where θ_u is one (or a few) gradient step updates from global parameter θ of the meta evidential learner with η being the step size, and λ_2 is the regularization parameter.

We apply an optimization-based meta-learning approach [20] to learn user specific factors, as shown in Figure 1. The meta evidential learning consists of three fully connected linear layers with ReLU activation in the first two, while the last layer predicts ratings or count and its evidence. The input to the meta evidential learning model is the concatenation of user embedding (z_u) and item embedding (z_i) for each user, i.e., $(z_u || z_i)$. For the meta evidential learning module, the local update is done for the user-specific parameter, which is achieved by one or more gradients from the global parameter:

$$\theta_u = \theta - \eta \nabla_{\theta} \mathcal{L}_{\mathcal{T}_u} [f_{\theta, E_u, E_i}] \quad (22)$$

In this update, the loss function is computed with the support set. Global update is done with the new item interactions of each user from the query set:

$$\theta = \theta - \xi \nabla_{\theta} \sum_{\mathcal{T}_u \sim p(\mathcal{T})} \mathcal{L}_{\mathcal{T}_u} [f_{\theta, E_u, E_i}] \quad (23)$$

This process will continue until it converges to a good global parameter shared by all users.

V. EXPERIMENTS

A. Dataset Description

We evaluated our model on four public benchmark datasets. Three are explicit datasets where users provide explicit ratings: MovieLens 1M (1M explicit ratings made by 6,040 users on 3,900 distinct movies from 04/2000 to 02/2003), Netflix (6,042 users with their interaction history from 01/2002 to 12/2005), and Book Crossing (751 users and their interacted books in a 4-week span during August-September of 2004); and one implicit dataset, where users have implicit interactions, (captured by count): Last.fm-1K dataset (listening history for nearly 1,000 users).

B. Baselines

For comparison, we include two matrix factorization based deep learning models: *DeepFM* [13] and *Wide & Deep* [12], one graph based model: *GC-MC* [15], and a meta-learning based recommendation model: *MeLU* [5].

C. Results and Discussion

The experimental results for the proposed model and baselines are summarized in Table II. We compute the average RMSE considering all users with the range of deviation for all datasets: MovieLens 1M, Book Crossing, Netflix, and Last.fm, respectively. The proposed model benefits from the meta-learning module, and hence it can effectively handle cold-start users who have few interactions like those in Book Crossing datasets. We also observe from Table II that deep learning and graph based models have poor performance on the Book Crossing datasets than meta-learning models like MeLU and the proposed model achieves significant improvements. For the last.fm dataset, the meta-learning models have shown a

clear indication of improvement again over deep learning and graph based models is not applicable due to implicit datasets. For the MovieLens 1M and Netflix datasets, most users have enough interactions, and hence all models achieve comparable performances. We further provide top N NDCG performance ranging from top 5 to 25 and their respective values for each model. For this, we chose those test users with 30 interactions so that we can use 25 interactions for query set to compute NDCG. The result is consistent with the RMSE results.

D. Uncertainty-Aware Recommendations

In this set of experiments, we show how the model effectively leverages predicted uncertainty to recommend the *most informative items* rather than solely based on the predicted ratings. For this, we randomly chose a test user (*ID: 41*) from the MovieLens-1M dataset. This user has a total of 25 interactions, and we randomly choose 20 interactions that serve as the candidate pool to form the support set. The remaining 5 interactions are used for a query set. We perform uncertainty-based recommendation to tackle cold-start problem where we recommend a few items from the pool according to their epistemic uncertainty (instead of predicted ratings). By collecting only limited interaction results, we expect the model to learn the most from the cold-start user (by reducing the epistemic uncertainty) to provide more accurate recommendation in the future. To demonstrate that the uncertainty-based recommendation can lead to better future predictions, we also employ the classical rating based recommendation to select same number of highest rated items. After the adaptation using the selected support set, both methods will be evaluated on the same query set for comparison.

We first show total counts of genres and ratings by the left and middle plots of Figure 2. From those plots, we can clearly see that the epistemic method selects more diverse genres with more count in *others* genres. It also selects items with relatively lower ratings than the rating-based method. This suggests that rating based recommendation seems more specific to the adventure movies, whereas epistemic method selects more diverse genres, including *drama*, *adventure*, and a higher number of *others* genres.

We further investigate how interactions selected based on epistemic uncertainty help to provide a better future recommendation. For this, we make fast adaptation of our meta-train model with those few interactions resulted from the recommended items and then perform testing on the query set. We start by adding 5 interactions and continue to add 5 in each round until all the items in the candidate pools are used. As we can see from the right plot of Figure 2, after adding 10 interactions based on the recommended items, the epistemic method achieves almost optimal performance on the query set. In contrast, the rating based method requires more than 15 interactions to achieve similar performance on the query set.

VI. CONCLUSIONS

This paper presents a novel meta evidential learning recommendation framework that integrates evidential learning with

Table II: Performance of Recommendation (average RMSE and NDCG)

Model	MovieLens-1M		Book Crossing		Netflix		Last.fm	
	RMSE	NDCG	RMSE	NDCG	RMSE	NDCG	RMSE	NDCG
deepFM	1.0254±0.03	0.2913	4.0889±0.06	0.2733	0.9699±0.02	0.2915	1.1939±0.05	0.2807
Wide & Deep	1.0218±0.03	0.2932	4.1341±0.08	0.2745	0.9686±0.02	0.2944	1.1847±0.05	0.2812
GC-MC	1.0313±0.03	0.2872	4.1405±0.10	0.2712	0.9816±0.03	0.2814	N/A	N/A
MeLU	1.0195±0.02	0.3308	3.7388±0.05	0.2811	0.9613±0.02	0.3265	1.0711±0.03	0.3102
MetaEDL	1.0114±0.02	0.3493	3.7026±0.04	0.3046	0.9525 ±0.02	0.3488	1.0183±0.03	0.3233

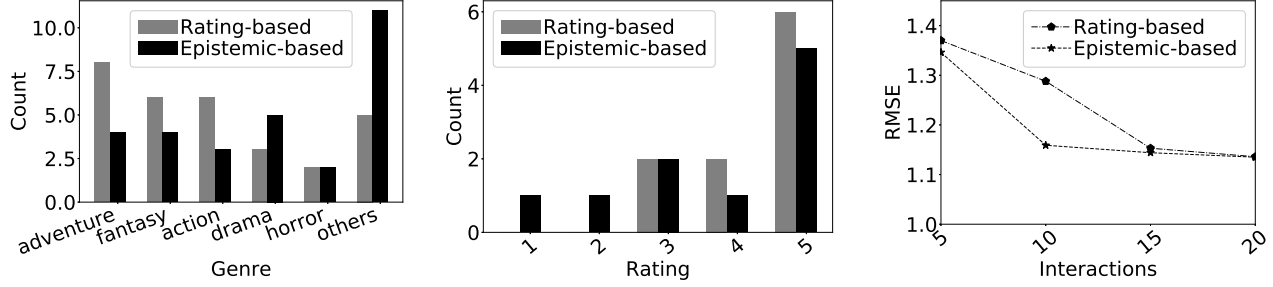


Figure 2: Genre count and rating count for the items selected in the support set with size 10 and RMSE for the query set

meta-learning to provide uncertainty-aware cold-start recommendations. The proposed framework handles the user cold-start problem by adopting global knowledge of similar users from their interaction information and leveraging evidential learning for efficient posterior inference to quantify the model confidence. Experimental results on four real-world datasets and comparison with the state-of-the-art competitive models clearly demonstrate the effectiveness of the proposed model.

ACKNOWLEDGEMENT

This research was supported in part by an NSF IIS award IIS-1814450 and an ONR award N00014-18-1-2875. The views and conclusions contained in this paper are those of the authors and should not be interpreted as representing any funding agency.

REFERENCES

- [1] Dhoha Almazro, Ghadeer Shahatah, Lamia Alabdulkarim, Mona Kharees, Romy Martinez, and William Nzoukou. A survey paper on recommender systems. *arXiv preprint arXiv:1006.5278*, 2010.
- [2] David Goldberg, David Nichols, Brian M Oki, and Douglas Terry. Using collaborative filtering to weave an information tapestry. *Communications of the ACM*, 35(12):61–70, 1992.
- [3] Pasquale Lops, Marco De Gemmis, and Giovanni Semeraro. Content-based recommender systems: State of the art and trends. In *Recommender systems handbook*, pages 73–105. Springer, 2011.
- [4] Manasi Vartak, Arvind Thiagarajan, Conrado Miranda, Jeshua Bratman, and Hugo Larochelle. A meta-learning perspective on cold-start recommendations for items. In *Advances in neural information processing systems*, pages 6904–6914, 2017.
- [5] Hoyeop Lee, Jinbae Im, Seongwon Jang, Hyunsouk Cho, and Sehee Chung. Melu: Meta-learned user preference estimator for cold-start recommendation. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1073–1082, 2019.
- [6] Alex Kendall and Yarin Gal. What uncertainties do we need in bayesian deep learning for computer vision? In *Advances in neural information processing systems*, pages 5574–5584, 2017.
- [7] Yehuda Koren, Robert Bell, and Chris Volinsky. Matrix factorization techniques for recommender systems. *Computer*, 42(8):30–37, 2009.
- [8] Yehuda Koren. Factorization meets the neighborhood: a multifaceted collaborative filtering model. In *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 426–434, 2008.
- [9] Yehuda Koren. Collaborative filtering with temporal dynamics. In *Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 447–456, 2009.
- [10] Laurent Charlin, Rajesh Ranganath, James McInerney, and David M Blei. Dynamic poisson factorization. In *Proceedings of the 9th ACM Conference on Recommender Systems*, pages 155–162, 2015.
- [11] San Gultekin and John Paisley. A collaborative kalman filter for time-evolving dyadic processes. In *2014 IEEE International Conference on Data Mining*, pages 140–149. IEEE, 2014.
- [12] Heng-Tze Cheng, Levent Koc, Jeremiah Harmsen, Tal Shaked, Tushar Chandra, Hrishikesh Aradhye, Glen Anderson, Greg Corrado, Wei Chai, Mustafa Ipsir, et al. Wide & deep learning for recommender systems. In *Proceedings of the 1st workshop on deep learning for recommender systems*, pages 7–10, 2016.
- [13] Hui Feng Guo, Ruiming Tang, Yunming Ye, Zhenguo Li, and Xiuqiang He. Deepfm: a factorization-machine based neural network for ctr prediction. *arXiv preprint arXiv:1703.04247*, 2017.
- [14] Qingyu Guo, Fuzhen Zhuang, Chuan Qin, Hengshu Zhu, Xing Xie, Hui Xiong, and Qing He. A survey on knowledge graph-based recommender systems. *IEEE Transactions on Knowledge and Data Engineering*, 2020.
- [15] Rianne van den Berg, Thomas N Kipf, and Max Welling. Graph convolutional matrix completion. *arXiv preprint arXiv:1706.02263*, 2017.
- [16] Jürgen Schmidhuber. *Evolutionary principles in self-referential learning, or on learning how to learn: the meta-meta-... hook*. PhD thesis, Technische Universität München, 1987.
- [17] Junyang Jiang, Deqing Yang, Yanghua Xiao, and Chenlu Shen. Convolutional gaussian embeddings for personalized recommendation with uncertainty. *arXiv preprint arXiv:2006.10932*, 2020.
- [18] Alexander Amini, Wilko Schwarting, Ava Soleimany, and Daniela Rus. Deep evidential regression. *Advances in Neural Information Processing Systems*, 33, 2020.
- [19] Oriol Vinyals, Charles Blundell, Timothy Lillicrap, Daan Wierstra, et al. Matching networks for one shot learning. In *Advances in neural information processing systems*, pages 3630–3638, 2016.
- [20] Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. *arXiv preprint arXiv:1703.03400*, 2017.