

Micro- and Macro-Level Churn Analysis of Large-Scale Mobile Games

Xi Liu^{1*}, Muhe Xie², Xidao Wen^{3*}, Rui Chen², Yong Ge⁴, Nick Duffield¹ and Na Wang²

¹Texas A&M University, College Station, TX, USA

²Samsung Research America, Mountain View, CA, USA

³University of Pittsburgh, Pittsburgh, PA, USA

⁴The University of Arizona, Tucson, AZ, USA

Abstract. As mobile devices become more and more popular, mobile gaming has emerged as a promising market with billion-dollar revenue. A variety of mobile game platforms and services have been developed around the world. A critical challenge for these platforms and services is to understand the churn behavior in mobile games, which usually involves churn at micro level (between an app and a specific user) and macro level (between an app and all its users). Accurate micro-level churn prediction and macro-level churn ranking will benefit many stakeholders such as game developers, advertisers, and platform operators. In this paper, we present the first large-scale churn analysis for mobile games that supports both micro-level churn prediction and macro-level churn ranking. For micro-level churn prediction, in view of the common limitations of the state-of-the-art methods built upon traditional machine learning models, we devise a novel semi-supervised and inductive embedding model that jointly learns the prediction function and the embedding function for user-app relationships. We model these two functions by deep neural networks with a unique edge embedding technique that is able to capture both contextual information and relationship dynamics. We also design a novel attributed random walk technique that takes into consideration both topological adjacency and attribute similarities. To address macro-level churn ranking, we propose to construct a relationship graph with estimated micro-level churn probabilities as edge weights and adapt link analysis algorithms on the graph. We devise a simple algorithm *SimSum* and adapt two more advanced algorithms *PageRank* and *HITS*. The performance of our solutions to the two-level churn analysis problem is evaluated on real-world data collected from the Samsung Game Launcher platform. The data includes tens of thousands of mobile games and hundreds of millions of user-

* This work was done during the authors' internships at Samsung Research America

Received 11 January 2019

Revised 26 March 2019

Accepted 10 June 2019

app interactions. The experimental results with this data demonstrate the superiority of our proposed models against existing state-of-the-art methods.

Keywords: Churn prediction; Representation learning; Graph embedding; Inductive learning; Semi-supervised learning; Mobile games

1. Introduction

With the wide adoption of mobile devices (e.g., smartphones and tablets), the past decade has seen a rapid increase of the mobile gaming industry into a billion-dollar market around the globe. \$70.3 billion revenue was generated by mobile games in 2018 according to Newzoo’s Global Games Market Report [1]. It is expected that mobile games will generate \$106.3 billion revenue in 2021, accounting for more than half of the overall game market. Driven by this increasingly vital market, software and hardware providers of mobile devices (e.g., Apple, Google and Samsung) have provided integrated mobile game platforms and services for end users, game developers and other stakeholders. Some examples of such platforms and services are Apple App Store, Google Play and Samsung Game Launcher [2].

One particularly crucial task, within these mobile game platforms and services, is understanding the churn behavior in mobile games, which usually involves two levels of analysis: the micro level (between an app and a specific user) and the macro level (between an app and all its users). In many traditional business applications [3], accurately predicting micro-level churn has been a long-standing and important task, the objective of which is to predict the likelihood that a user will stop using a service or product. In our problem setting, the aim is to predict the likelihood that a user will stop using a particular game app in the future. This task is vital for the following reasons. First, the churn rate of a mobile game is an important business metric to measure its success. With accurately predicted churn probabilities of player-game pairs, a game platform is able to prioritize its resources for better operation and management. Second, predicting individual churn probabilities will enable a game platform to design better marketing strategies to improve user retention. Examples include sending push notifications and providing free items in games to users who are likely to churn. Since, as well known, the acquisition cost for new users is much higher than the retention cost for existing users, successful micro-level churn prediction could largely reduce costs for game developers and platform operators to increase the number of active users, which plays a critical role in the success of a mobile game. Third, the micro-level churn prediction provides direct input to determine the right timing of app recommendations for a game platform. The results of this research will enable testing of the hypothesis that a user is more likely to act on an app recommendation when he/she is about to stop playing other games.

The objective of macro-level churn ranking is to provide a list of games ranked by their total number of users who will churn in the near future. It is an equally important problem, and its solution can be applicable to various mobile app services such as trend prediction [4], rating and review spam detection, ranking fraud detection [5], and especially trend-sensitive recommendation [6]. There has been a general consensus that the mobile app market is highly dynamic: new apps continuously enter the market and distract users’ attention from existing

ones. A recommender can use the results of macro-level churn ranking as prior knowledge, e.g., an attribute in the representation of games. A recommender can also slightly adjust the positions of recommended games before delivering to users in preference to those that have smaller total numbers of near-future churn. Macro-level churn ranking in general can be very challenging because the order of the games on the list (1) can change very fast with the dynamics of the market, and (2) is based on future events, which can only be estimated. The macro-level churn ranking problem is intimately related to the micro-level churn prediction problem. Consider the extreme case that all micro-level churn predictions are 100% correct, then the number of users to churn in the future can be directly computed.

There have been several previous studies [7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17] on micro-level mobile game churn prediction by using traditional machine learning models (e.g., logistic regression, random forests, Cox regression). However, we observe several major limitations of these studies. First, they were developed for predicting churn of a single or a few mobile games. None is capable of handling churn prediction of large-scale mobile apps and users. For real-world applications, a solution needs to handle tens of thousands of mobile games and hundreds of millions of user-app interactions on a daily basis. Second, user-app interaction data often comes with rich contextual information (e.g., WiFi connection status, screen brightness, and audio volume), which has never been considered in the existing studies. Third, the existing methods rely on handcrafted features that usually cannot scale well in practice.

To overcome all these limitations, we propose a novel inductive semi-supervised embedding model that *jointly* learns the prediction function and the embedding function for user-game interactions. The user-app interaction data includes detailed information of opens, closes, installs and uninstalls of game apps for each individual user. This data is collected from the Samsung Game Launcher platform [2], which is pre-loaded in most smartphones manufactured by Samsung. We model the interplay between users and games by an attributed bipartite graph and then learn these two functions by deep neural networks with a unique embedding technique that is able to capture both contextual information and dynamic user-game interactions. Our method is fully automatic and can be easily integrated into existing mobile game platforms.

With the predicted micro-level churn probabilities, we devise a simple method *SimSum* to estimate the macro-level churn ranking. Under the presupposition that the churn probabilities are properly estimated, we show that *SimSum* is based on an unbiased estimation. Furthermore, as the churn behavior may be related among users and games, we propose to construct a user-game graph with edges weighted by the estimated churn probabilities and adapt link analysis algorithms to infer the macro-level churn ranking of games. Among the family of link analysis algorithms, *PageRank* [18] and *HITS* (Hyperlink-Induced Topic Search) [19] are chosen because they have been widely applied in abundant real-world applications and shown to have stable performance.

Contributions. Our research contributions are summarized as follows:

1. To the best of our knowledge, this paper is the first to develop solutions for both micro-level churn prediction and macro-level churn ranking of mobile games at scale. Although the paper mainly applies the proposed solutions to

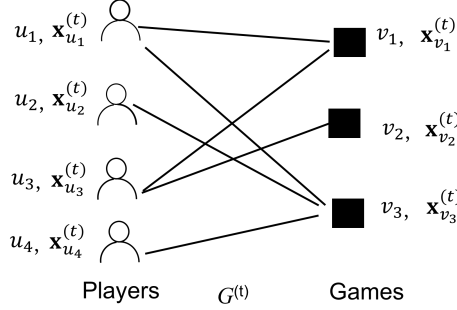


Fig. 1. An example of an attributed bipartite graph for mobile game churn analysis

mobile game churn analysis, they are also applicable to churn analysis in other contexts.

2. We propose a novel semi-supervised and inductive model based on embedding. Our model can capture the dynamics between users and mobile games based on the introduced temporal loss in the formulated objective function. The model is able to embed new users or games not used in training. This is critical for mobile game churn prediction because new games and users continually enter the market.
3. We develop an attributed random walk technique that enables us to sample the contexts of edges in an attributed bipartite graph and that takes into account both topological adjacency and attribute similarities.
4. We propose a simple method *SimSum* and adapt two link analysis algorithms *PageRank* and *HITS* to solve the macro-level churn ranking problem.
5. We conduct a comprehensive experimental evaluation with large-scale real-world data consisting of hundreds of millions of user-app interaction records collected from Samsung Game Launcher, one of the largest commercial mobile game platforms. The experimental results demonstrate that our model outperforms all state-of-the-art methods with respect to different evaluation metrics for micro-level churn prediction and macro-level churn ranking.

The rest of the paper is organized as follows. Section 2 formulates the two-level mobile game churn analysis problem. Section 3 discusses our solution in detail. Section 4 presents our experimental results on large-scale real-world data. Section 5 reviews the related literature. Finally, Section 6 concludes our work.

2. Problem Formulation

In the context of mobile games, *churn* is defined as a player stopping using a game within a given period (i.e., there is no app usage in the period). The duration T of the period may vary from application to application depending on different business goals. $T = 14$ days and $T = 30$ days are some typical settings used in industry [7, 8, 12]. In this paper, we consider the generic micro-level game churn prediction and macro-level churn ranking problems without assuming any particular value of T . We note that uninstall is different from churn. Regarding

Table 1. Notations and definitions

Notations	Descriptions or Definitions
$\mathcal{G}^{(t)}$	Attributed graph at time t
$\mathcal{H}^{(t)}$	Historical attributed graphs $\{\mathcal{G}_i\}_{i=t_0}^{i=t}$
$\mathcal{U}^{(t)}$	Set of all player nodes in $\mathcal{G}^{(t)}$
$\mathcal{V}^{(t)}$	Set of all game nodes in $\mathcal{G}^{(t)}$
$\mathcal{E}^{(t)}$	Set of all edges in $\mathcal{G}^{(t)}$
$e_{uv}^{(t)}$	Indicator of the existence of edge (u, v) in $\mathcal{G}^{(t)}$
$\mathbf{x}_u^{(t)}$	Feature vector of user $u \in \mathcal{U}^{(t)}$
$\mathbf{x}_v^{(t)}$	Feature vector of game $v \in \mathcal{V}^{(t)}$
$\mathbf{z}_{uv}^{(t)}$	Aggregated feature vector of edge $(u, v) \in \mathcal{E}^{(t)}$
n_u, n_v	Numbers of attributes in $\mathbf{x}_u^{(t)}$ and $\mathbf{x}_v^{(t)}$, resp.
d	Number of attributes in $\mathbf{z}_{uv}^{(t)}$
m	Embedding dimension
g	Edge embedding function $g : \mathcal{R}^d \rightarrow \mathcal{R}^m$
f	Churn prediction function $f : \mathcal{R}^m \rightarrow [0, 1]$
l_p	Number of embedding hidden layers in Part I
l_n	Number of prediction hidden layers in Part III
$ \cdot $	Cardinality of a set
$A \setminus B$	Set of elements in A but not in B
$\mathbb{1}(c)$	Indicator function for condition c
$\mathcal{S}^{(t)}(v)$	Ranking score of game v at time t

only uninstall as churn would be problematic since there may be a large time gap between cessation of playing and uninstall, if any.

The relationships between players and games can be represented by an attributed bipartite graph as illustrated in Fig.1, whose two parts correspond to players and games. In the sequel, we use the terms *player* and *user*, and *game* and *app* interchangeably. Let $\mathcal{G}^{(t)}$ be the attributed bipartite graph at time t , the vertex set $\mathcal{U}^{(t)}$ denote the set of users and the vertex set $\mathcal{V}^{(t)}$ denote the set of games. A player is represented by a node $u \in \mathcal{U}^{(t)}$ and a game is represented by a node $v \in \mathcal{V}^{(t)}$. Each user u is associated with a feature vector $\mathbf{x}_u^{(t)} \in \mathcal{R}^{n_u}$, where n_u is the size of $\mathbf{x}_u^{(t)}$; each game v is associated with a feature vector $\mathbf{x}_v^{(t)} \in \mathcal{R}^{n_v}$, where n_v is the size of $\mathbf{x}_v^{(t)}$. There is an edge between nodes u and v in $\mathcal{G}^{(t)}$ if player u played v in the time window $[t+1, t+T]$. The set of neighbor nodes that connect to game v at time t is denoted by $\mathcal{N}_v^{(t)}$. The set of edges at time t is denoted by $\mathcal{E}^{(t)}$.

Now we are ready to define the two-level mobile game churn analysis problem below.

Definition 2.1 (Mobile game churn analysis). Consider a collection of attributed bipartite graphs observed from time t_0 to time t ($t > t_0$), which is denoted by $\mathcal{H}^{(t)} = \{\mathcal{G}^{(i)}\}_{i=t_0}^{i=t}$. Let $\mathcal{N}_v^{(i)}$ be the set of neighbors of game v at time i and $e_{uv}^{(i)}$ be the indicator of the existence of edge (u, v) in $\mathcal{G}^{(i)}$. The dual objectives are:

1. *Micro-level churn prediction* that estimates the probability $Pr(e_{uv}^{(t+1)} = 0 | e_{uv}^{(t)} = 1, \mathcal{H}^{(t)})$ for an edge (u, v) in $\mathcal{G}^{(t)}$, which is the probability that (u, v) disappears in $\mathcal{G}^{(t+1)}$;
2. *Macro-level churn ranking* that ranks the games in a descending order of the total number of users to churn at time $t + 1$.

The notations used in this paper and their descriptions are listed in Table 1.

3. Methods

3.1. Overview of Our Solution

Most existing works [7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17] rely on traditional machine learning models to solve the micro-level churn prediction problem, which suffer from several major limitations as mentioned in Section 1. In view of the recent progress in graph embedding, a promising way is to adopt graph embedding frameworks for churn prediction. However, we face several key technical challenges that have not been addressed in prior studies: (1) Most existing methods are transductive, and thus cannot produce embeddings for new player-game pairs. (2) Existing methods are either purely supervised or unsupervised, and thus do not take full advantage of relevance between embedding and a task. (3) Existing methods are node-centric, and thus are not directly applicable to edge-centric tasks. (4) Existing methods mainly handle a static graph and do not incorporate graph dynamics in embedding. In the mobile game industry, players and games, however, change very quickly. There are many new players, new games, and new player-game relationships almost every day.

In addressing these challenges, we propose a novel inductive semi-supervised embedding model in dynamic graphs that jointly learns the prediction function f and the embedding function g . The prediction function f and the embedding function g are learned by deep neural networks (DNNs). The architecture of the proposed DNN is presented in Fig. 2, which consists of three parts. Part I is responsible for producing embedding feature vectors $g(\mathbf{z}_{uv}^{(t)})$ from raw edge feature vectors $\mathbf{z}_{uv}^{(t)} \in \mathcal{R}^d$, where d is the size of raw edge feature vectors. To learn the probability of churn, we need to construct a feature vector for each (u, v) with $e_{uv}^{(t)} = 1$. However, it is impractical to calculate features for all possible edges that may appear in the prediction period because the number of possible edges is huge, which is $\mathcal{O}(|\mathcal{U}^{(t)}| \cdot |\mathcal{V}^{(t)}|)$. Instead, we construct the feature vector $\mathbf{z}_{uv}^{(t)}$ of (u, v) from attribute-wise cosine similarity aggregation of $\mathbf{x}_u(t)$ and $\mathbf{x}_v(t)$. As such, all $\mathbf{z}_{uv}^{(t)}$ can be readily computed from the features of $|\mathcal{U}^{(t)}| + |\mathcal{V}^{(t)}|$ nodes, where $|\cdot|$ denotes the cardinality of a set.

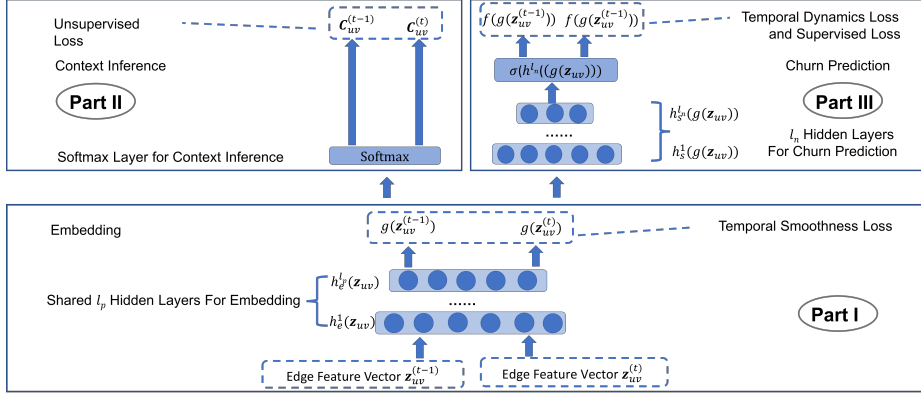


Fig. 2. Deep neural network architecture of the inductive semi-supervised embedding model in training

Part II is for inferring contexts from embedding feature vectors. Here the context of an edge refers to the edges that are similar to and co-occur with the edge under some graph sampling strategy, for example, random walk. Part I and Part II form the *unsupervised* component of our model. They are jointly trained by minimizing the error of incorrect context inference and inconsistency with respect to temporal smoothness (see Section 3.3 for an explanation of temporal smoothness). Part I and Part II are trained in an inductive and edge-centric way. In contrast to transductive node embedding that learns a distinct embedding vector for each node, our model is to learn an embedding function that is generalized to any *unseen* edges as long as their feature vectors are available. We propose a novel attributed random walk to sample similar edges as contexts (see Section 3.4 for details).

Part III fulfills the supervised churn prediction task from embedding feature vectors. Part III forms the *supervised* component of the proposed model, which is trained by minimizing the error of incorrect churn predictions. The supervised component and unsupervised component are simultaneously trained as combined into a single objective function. Traditional unsupervised embedding techniques are not designed in a task-specific way and hence are not able to incorporate task-specific information to improve performance. In contrast, in our model Part III and Part II share the common hidden layers in Part I, and therefore they are implicitly coupled with each other. This helps the embedding align with the supervised prediction task.

Part I and Part III both consider graph dynamics in training. Part I handles graph dynamics by requiring the embeddings of the same edge at two consecutive timestamps to stay close. Part III handles graph dynamics by requiring the churn probabilities of the same edge at two consecutive timestamps to follow a decaying pattern.

The overall objective function of our model consists of four parts:

$$\mathcal{L} := \mathcal{L}_S + \alpha \mathcal{L}_U + \beta \mathcal{L}_T + \gamma \mathcal{L}_R. \quad (1)$$

$\mathcal{L}_S$ denotes the supervised loss due to incorrect predictions and will be discussed in Section 3.2.1. $\mathcal{L}_U$ denotes the unsupervised loss, which comes from failures of context inference and will be addressed in Section 3.2.2. $\mathcal{L}_T$ is the temporal loss

that consists of two parts: temporal smoothness and temporal dynamics, and will be explained in Section 3.3. $\mathcal{L}_R$ presented in Section 3.2.3 is the regularization term, and (α, β, γ) are trade-off weights.

After the prediction function f and embedding function g are learned through minimizing the loss function in Equation (1), the churn probability $Pr(e_{uv}^{(t+1)} = 0 | e_{uv}^{(t)} = 1, \mathcal{H}^{(t)})$ of any individual user-game pair (u, v) can be estimated by $f(g(\mathbf{z}_{uv}^{(t)}))$. Under the assumption that the micro-level churn probability is properly estimated, we are able to show that for a specific game v , the sum of its churn probabilities of all its users is an *unbiased* estimator of the ground truth $|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}|$. This inspires us to derive a simple solution for macro-level churn ranking: first estimate the total number of users to churn $|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}|$ for each game $v \in \mathcal{V}^{(t)}$ by summing all its users' churn probabilities and then rank all games based on the estimations. This method is referred to as *SimSum*. Furthermore, we observe that there may be correlations in the churn behavior between different games and different users. For instance, for the same user, the probability of his/her churning a game is likely to be influenced by the probability of churning other games. On the other hand, similar users may exhibit similar patterns of churning the same game. Therefore, we propose to construct a user-game graph with edges weighted by the estimated churn probabilities and adapt link analysis algorithms to estimate macro-level churn ranking. Specifically, *PageRank* [18] and *HITS* [19] are chosen due to their wide applications and stable performance. Details will be discussed in Section 3.5.

3.2. Static Loss Functions

3.2.1. Supervised Loss Function $\mathcal{L}_S$

The supervised loss function $\mathcal{L}_S$ is designed for Part III. Let $h_s^k(g(\mathbf{z}_{uv}^{(t)})) = \phi(W_s^k h_s^{k-1}(g(\mathbf{z}_{uv}^{(t)})) + b^k)$ be the k -th hidden layer for churn prediction (referred to as *prediction hidden layer* in the sequel), where W_s^k and b_s^k are the weights and biases in the k -th prediction hidden layer, and $\phi(\cdot)$ is a non-linear activation function. We model the churn prediction function f by l_n such layers in Part III. Then the prediction output layer can be represented by:

$$\begin{aligned} f(g(\mathbf{z}_{uv}^{(t)})) &:= \widehat{Pr}(e_{uv}^{(t+1)} = 0 | e_{uv}^{(t)} = 1, \mathcal{H}^{(t)}) \\ &:= \sigma(h_s^{l_n}(g(\mathbf{z}_{uv}^{(t)}))) \\ &:= \frac{\exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(t)}))^T w_s)}{1 + \exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(t)}))^T w_s)} \end{aligned} \quad (2)$$

where $\widehat{Pr}(\cdot)$ is the estimate of $Pr(\cdot)$, $\sigma(\cdot)$ is the sigmoid function, and w_s is the sigmoid weight vector that combines the outputs from the last hidden layer to predict churn. Now we can define the supervised loss $\mathcal{L}_S$ as follows:

$$\mathcal{L}_S = \frac{1}{L} \sum_{i=0}^{t-1} \sum_{(u,v) \in \mathcal{E}^{(i)}} \delta_{uv}^{(i+1)} \left(1 - e_{uv}^{(i+1)} - \frac{\exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(i)}))^T w_s)}{1 + \exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(i)}))^T w_s)} \right)^2, \quad (3)$$

where L is the number of training examples and $\delta_{uv}^{(i+1)}$ is a censoring indicator that will be discussed in Section 3.3.

3.2.2. Unsupervised Loss Function $\mathcal{L}_U$

The unsupervised loss function $\mathcal{L}_U$ is devised for Part II, which guides to embed the handcrafted features $\mathbf{z}_{uv}^{(t)} \in \mathcal{R}^d$ into a latent space $g(\mathbf{z}_{uv}^{(t)}) \in \mathcal{R}^m$, where m is the size of the latent space. Denote the k -th hidden layer for embedding (referred to as *embedding hidden layer* in the sequel) by $h_e^k(\mathbf{z}_{uv}^{(t)}) = \phi(W_e^k h_e^{k-1}(\mathbf{z}_{uv}^{(t)}) + b_e^k)$, where W_e^k and b_e^k are the weights and biases in the k -th embedding hidden layer. We use the l_p layers in Part I to represent the embedding function g , and the embedding output layer can be represented by:

$$g(\mathbf{z}_{uv}^{(t)}) := h_e^{l_p}(\mathbf{z}_{uv}^{(t)}). \quad (4)$$

We can define the unsupervised loss function as follows:

$$\mathcal{L}_U = - \sum_{i=t_0}^t \sum_{(u,v) \in \mathcal{E}^{(i)}} \sum_{(u',v') \in C_{uv}^{(i)}} \delta_{uv}^{(i)} \log \left(\text{Pr}((u', v') | g(\mathbf{z}_{uv}^{(i)})) \right), \quad (5)$$

where $C_{uv}^{(i)}$ denotes the context (i.e., contextual edges) of (u, v) in $\mathcal{G}^{(i)}$. The contextual edges $C_{uv}^{(i)}$ are obtained by attributed random walk on the bipartite graph, which will be discussed in Section 3.4.

The likelihood of having a contextual edge (u', v') of (u, v) conditional on the embedding of (u, v) is:

$$\text{Pr}((u', v') | g(\mathbf{z}_{uv}^{(i)})) = \frac{\exp \left(g(\mathbf{z}_{uv}^{(i)})^T w_{u'v'} \right)}{\sum_{(u^*, v^*)} \exp \left(g(\mathbf{z}_{uv}^{(i)})^T w_{u^*v^*} \right)}, \quad (6)$$

where $w_{u^*v^*}$ and $w_{u'v'}$ are the vectors of weights for edges (u^*, v^*) and (u', v') in the softmax layer, respectively. The denominator is computed by negative sampling [20]. Although the embedding function is learned by training a context inference task, it is still considered as “unsupervised” because the contexts are calculated by sampling on the attributed graph, which is independent of any supervised learning task [21].

The objective function in Equation (1) contains both the supervised loss function and the unsupervised loss function. Thus the embedding hidden layers are jointly trained with prediction hidden layers. Compared to traditional embedding methods, the semi-supervised approach makes the embedding more suitable for the prediction task.

Note that the target of the embedding process is to learn a mapping function from the feature space to the embedding space, instead of directly learning the embeddings. Thus, the input for the embedding hidden layers only contains attributes. A new edge can be embedded for churn prediction as long as we can observe its attributes. This indicates that the proposed approach is inductive. It is able to produce the embedding for an edge that is not used in training.

3.2.3. Regularization Loss $\mathcal{L}_R$

Regularization loss is introduced mainly to avoid overfitting. The weights for regularization consist of $\{\{W_e^i, b_e^i\}_{i=1}^{l_p}, \{W_s^i, b_s^i\}_{i=1}^{l_n}, w_s\}$. Therefore the regularization part can be expressed as:

$$\begin{aligned} \mathcal{L}_R := & \lambda_0 \sum_{i=1}^{l_p} \|W_e^i\|_2^2 + \lambda_1 \sum_{i=1}^{l_p} \|b_e^i\|_2^2 + \lambda_2 \sum_{i=1}^{l_n} \|W_s^i\|_2^2 \\ & + \lambda_3 \sum_{i=1}^{l_n} \|b_s^i\|_2^2 + \lambda_4 \|w_s\|_2^2, \end{aligned} \quad (7)$$

where $\{\lambda_i\}_{i=0}^4$ are trade-off weights on different regularization terms.

3.3. Temporal Loss Function

Temporal loss refers to the loss related to graph dynamics. Different from existing works [20, 21, 22], the proposed embedding model takes into account graph dynamics. Understanding the temporal dynamics of attributed bipartite graphs is crucial to precisely model the churn behavior. We make several key observations of the temporal dynamics, which help to achieve good prediction performance in practical settings.

Observation 1. For a given user-game play relationship, the longer the relationship exists, the more likely the user is to churn the game.

This is because the content of a mobile game is usually somewhat fixed. Players can easily lose interests after going through all contents and passing all levels in the game, let alone many players churn before passing all levels. With more days of play, their initial interests in the game gradually efface. Indeed, 71% of all mobile app users churn within 90 days [23]. The churn rate of mobile games is even higher. We plot the average retention rate of mobile games as a function of time, based on Game Launcher data in Fig. 3. It shows that 95% of user-game play relationships end after 40 days, which well justifies our observation. We formally state Observation 1 below.

$$f(g(\mathbf{z}_{uv}^{(i)})) \lesssim f(g(\mathbf{z}_{uv}^{(i+1)})) \quad \forall 0 \leq i \leq t-1, (u, v) \in D^{(i)}, \quad (8)$$

where $D^{(i)} = \{(u, v) : (u, v) \in \mathcal{E}^{(i)} \cap \mathcal{E}^{(i+1)}\}$ and $\lesssim$ represents “almost always smaller than”. We refer to this observation as *temporal dynamics* in the following discussion.

The second observation we make is stated in Observation 2.

Observation 2. For a given user-game play relationship denoted by an edge in the attributed bipartite graph, its context usually evolves slowly at two consecutive timestamps.

This observation is also derived from the real-world data. Due to space limitation, we omit the figure here. It follows that the topology and attribute values of the attributed bipartite graph mostly evolve smoothly at two consecutive timestamps, resulting in similar contexts for a given edge at consecutive timestamps.

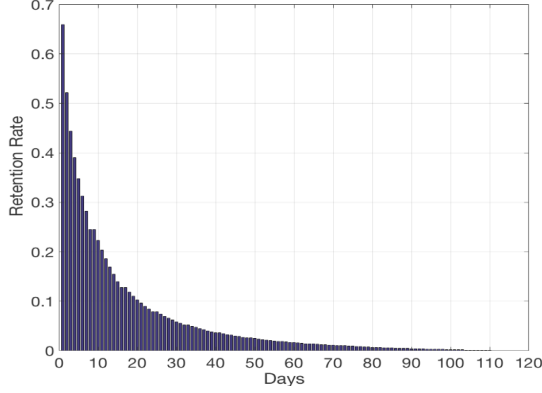


Fig. 3. Average mobile game retention rate as a function of time

Therefore, its embeddings at these two timestamps should also be close, that is,

$$g(\mathbf{z}_{uv}^{(i+1)}) \approx g(\mathbf{z}_{uv}^{(i)}) \quad \forall 0 \leq i \leq t-1, (u, v) \in D^{(i)}, \quad (9)$$

where $\approx$ represents “almost always equal to”. We refer to this observation as *temporal smoothness* in the later discussion.

The final observation is:

Observation 3. By definition, churn in nature introduces right censoring to the training dataset.

The problem of right censoring has been widely studied [24] and is illustrated in Fig. 4. The observation period refers to some time duration in history. Suppose we are at time t and the observation period is from time t_0 to time t . Data for training and testing all comes from the observation period. Since the label of a player-game pair at a specific timestamp is determined by their interaction in the next T time duration, the labels of some player-game pairs in the last T duration could be *unknown*. For instance, the last observation of the pair of player 2 and game 1 (denoted by p2-g1) was in the last T time duration, and therefore the labels after that time are unknown. This is known as right censoring. In contrast, the pair p3-g3 has play records every day in the last T days, and it is not censored during the observation period.

Considering the existence of censored instances, we introduce a binary indicator $\delta_{uv}^{(i)}$ to indicate whether an edge (u, v) is censored at timestamp i . $\delta_{uv}^{(i)} = 0$ if (u, v) is censored; $\delta_{uv}^{(i)} = 1$ otherwise. Then Inequality (8) needs to be updated by

$$\begin{aligned} f(g(\mathbf{z}_{uv}^{(i)})) &\lesssim f(g(\mathbf{z}_{uv}^{(i+1)})) \quad \forall 0 \leq i \leq t_{uv} - 1, (u, v) \in D^{(i)} \\ f(g(\mathbf{z}_{uv}^{(t_{uv})})) &\lesssim f(g(\mathbf{z}_{uv}^{(i)})) \quad \forall t_{uv} \leq i \leq t, (u, v) \in \mathcal{E}^{(t_{uv})}, \end{aligned} \quad (10)$$

where t_{uv} denotes the timestamp when the edge (u, v) was observed, i.e., $t_{uv} = \max\{i : \delta_{uv}^{(i)} = 1\}$. The update reflects the fact that after timestamp t_{uv} the label of (u, v) becomes unknown. Since the existence of edge (u, v) after t_{uv} is unknown, it is more reasonable to just require that Inequality (10) holds *pairwise* between time point t_{uv} and all time points after t_{uv} . Therefore the temporal loss

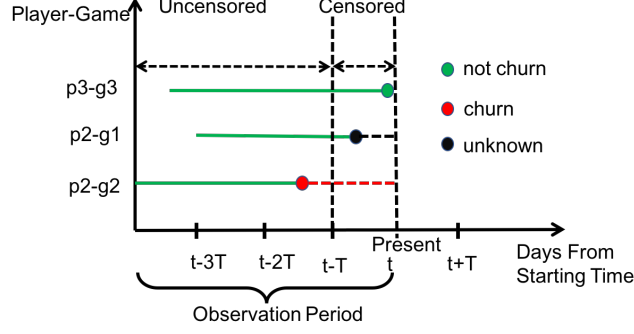


Fig. 4. An illustration of censored data by three randomly sampled player-game pairs

can be expressed as:

$$\mathcal{L}_T := \sum_{i=t_0}^{t-1} \sum_{(u,v) \in D^{(i)}} \left\{ \|g(\mathbf{z}_{uv}^{(i+1)}) - g(\mathbf{z}_{uv}^{(i)})\|_2 + \right. \quad (11)$$

$$\left. [\mathbb{1}(\delta_{uv}^{(i+1)} = 1)f(g(\mathbf{z}_{uv}^{(i)})) + \mathbb{1}(\delta_{uv}^{(i+1)} = 0)f(g(\mathbf{z}_{uv}^{(t_{uv})})) - f(g(\mathbf{z}_{uv}^{(i+1)}))]_+ \right\},$$

where $\mathbb{1}(A)$ denotes the indicator function for event A and $[x]_+ = \max\{x, 0\}$. The first term corresponds to the temporal smoothness and the second term corresponds to the temporal dynamics. Taking Equations (2) and (4) into Equation (11), we have the temporal loss $\mathcal{L}_T$ as follows.

$$\mathcal{L}_T := \sum_{i=0}^{t-1} \sum_{(u,v) \in D^{(i)}} \left\{ \|h_e^{l_p}(\mathbf{z}_{uv}^{(i+1)}) - h_e^{l_p}(\mathbf{z}_{uv}^{(i)})\|_2 + \right. \quad (12)$$

$$\left[\mathbb{1}(\delta_{uv}^{(i+1)} = 1) \frac{\exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(i)}))^T w_s)}{1 + \exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(i)}))^T w_s)} + \right.$$

$$\mathbb{1}(\delta_{uv}^{(i+1)} = 0) \frac{\exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(t_{uv})}))^T w_s)}{1 + \exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(t_{uv})}))^T w_s)} -$$

$$\left. \frac{\exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(i+1)}))^T w_s)}{1 + \exp(h_s^{l_n}(g(\mathbf{z}_{uv}^{(i+1)}))^T w_s)} \right]_+ \Big\}.$$

3.4. Context Generation by Attributed Random Walk

All above discussion assumes the availability of contexts of edges. There have been a series of node-centric research on graph embedding proposing to apply random walk to sample contexts [25]. These methods are normally topology-based. In our problem setting, our goal is to embed an *edge* (relationship), not a node (identity). In this case, a simple topology-based random walk may return two adjacent edges having the same player or the same game while totally ignor-

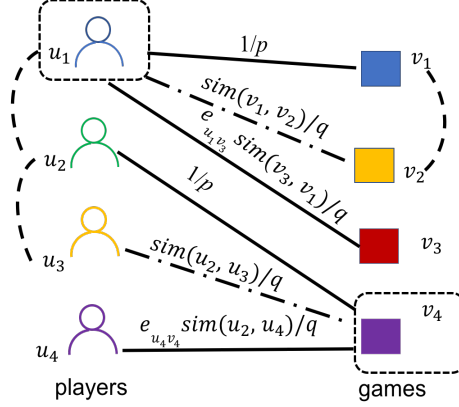


Fig. 5. An illustrative example of attributed random walk in an attributed bipartite graph

ing the similarity of the other end. This is undesirable. In contrast, attributed random walk measures such similarities by attributes and allows to transit to similar nodes even if they are not connected.

To this end, we propose a novel attributed random walk technique that takes into account both topological adjacency and attribute similarities to make the transition decision of the walk. Fig. 5 gives an illustrative example of attributed random walk in an attributed bipartite graph. For clarity, we omit the time index in the following discussion. The *solid line* indicates that there exists an edge in the attributed bipartite graph. We denote the type of node o by $type(o)$, which can be of value either player or game. The *dash-dotted* lines between users and games do not exist in the original attributed bipartite graph but may be considered as transitions by our attributed random walk due to attribute similarities.

The *dashed* lines in Fig. 5 represent the added augmented edges between the nodes and their similar same-typed nodes. Consider a random walker that just traversed edge (v_1, u_1) in Fig. 5 and now resides at node u_1 . The walker needs to decide which node to transit to. The goal is to traverse a new edge that is as similar as possible to the edge (v_1, u_1) . For two edges to be similar, their both ends should be similar. Since the next edge to traverse shares the same end u_1 , we only need to consider the attribute similarity of the other end. As a result, the walker should evaluate not only those nodes that are neighbors of u_1 as suggested in [25], but also the non-neighbor nodes that are similar to v_1 . It is time-consuming to calculate pairwise similarities between v_1 and all other nodes with the same type at each step. For this reason, we add augmented edges for a proportion of the same-typed nodes. For example, the augmented edges of node v_1 are $\{(v_1, o) : sim(v_1, o) > 1 - \epsilon, type(v_1) = type(o)\}$, where $0 < \epsilon < 1$ is a filtering parameter and:

$$sim(v_1, o) := \frac{\mathbf{x}_{v_1} \cdot \mathbf{x}_o}{\|\mathbf{x}_{v_1}\| \|\mathbf{x}_o\|}. \quad (13)$$

After adding the augmented edges, the nodes connecting v_1 through augmented edges are one hop from v_1 . We denote such nodes by $N_1(v_1) = \{v : d(v_1, v) =$

$1, \text{type}(v) = \text{type}(v_1)\}$, where $d(v_1, v)$ is the length of shortest path between v_1 and v . For example, $v_2 \in N_1(v_1)$ (due to the augmented edge between v_2 and v_1). Similarly, we denote the set of two-hop same-typed neighbor nodes of v_1 by $N_2(v_1) = \{v : d(v_1, v) = 2, \text{type}(v) = \text{type}(v_1)\}$. It contains those that share users with v_1 , e.g., v_3 . Then the transition probability in attributed random walk can be calculated as follows:

$$p(o|u_1) = \begin{cases} 0, & \text{if } \text{type}(o) \neq \text{type}(v_1) \\ \frac{1}{p}, & \text{if } o = v_1 \\ \frac{\text{sim}(v_1, o)}{q}, & \text{if } o \in N_1(v_1) \\ \frac{\text{sim}(v_1, o)e_{u_1 o}}{q}, & \text{if } o \in N_2(v_1) \end{cases} \quad (14)$$

where p and q are normalization constants used to control the walk strategy. Recall that $e_{u_1 o} = 1$ if there is an edge between u_1 and o . Similar equations can be defined when the random walker transits to a game node from user node, e.g., v_4 as shown in Fig. 5. The attributed walk in an attributed bipartite graph walks through different types of nodes repeatedly. It heuristically enlarges the probability of making any two consecutive edges in a path similar, and thus the probability of making the entire set of edges in the attributed random walk similar. As an example, suppose the walker transited from v_1 to u_1 as shown in Fig. 5. Since v_1 and v_2 are highly similar, the walker may transit to v_2 and produce a path with edges (v_1, u_1) and (u_1, v_2) even though u_1 is not connected to v_2 in the bipartite graph.

Our proposed method shares some similarities with the idea of [26]. However, in their work attribute similarities have no influence on which nodes a walker goes through. In our proposed method, when choosing the next node to visit, there is a non-zero probability to choose those that are not connected to the present node but share similar attributes to the previous node. For those that are connected to the present node, the probability to choose from them is weighted by the similarity between the descendant and the ancestor nodes. In this way, an edge can be the context of another even when they are not adjacent but similar in both ends.

3.5. Macro-Level Churn Ranking

The objective of macro-level churn ranking is to provide a ranked list of games based on their total numbers of users to churn in the near future. The ground truth of the number for game $v \in \mathcal{V}^{(t)}$ is $|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}|$. Let $\mathcal{S}^{(t)}(v)$ denote the score of ranking for game v , where a larger $\mathcal{S}^{(t)}(v)$ means a higher position. One solution is to first estimate the total number of users to churn $|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}|$ as the score $\mathcal{S}^{(t)}(v)$ for a game $v \in \mathcal{V}^{(t)}$ and then rank all games based on the estimations. We refer to it as *SimSum*. After the prediction function f and the embedding function g are learned, the churn probability $\hat{Pr}(e_{uv}^{(t+1)} = 0 | e_{uv}^{(t)} = 1, \mathcal{H}^{(t)})$ of any individual user-game pair (u, v) can be estimated by $f(g(\mathbf{z}_{uv}^{(t)}))$. We sum up the churn probabilities across all users of game v and consider the sum as the estimation of the total number of users to churn:

$$\mathcal{S}_{ss}^{(t)}(v) = \sum_{u \in \mathcal{U}^{(t)}} e_{uv}^{(t)} f(g(\mathbf{z}_{uv}^{(t)})). \quad (15)$$

where the subscript “ss” marks that the score is computed by the *SimSum* method.

Now we show that *SimSum* provides an unbiased estimation of the ground truth $|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}|$ under the assumption that the micro-level churn probability is properly estimated. This is because we have:

$$\begin{aligned} \mathbb{E}[|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}| - \mathcal{S}_{ss}^{(t)}(v) | \mathcal{H}^{(t)}] &= \mathbb{E}[|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}| \mathcal{H}^{(t)}] - \mathbb{E}[\mathcal{S}_{ss}^{(t)}(v) | \mathcal{H}^{(t)}] \\ &= \mathbb{E}\left[\sum_{u \in \mathcal{U}^{(t)}} \mathbb{1}(e_{uv}^{(t)} = 1) \mathbb{1}(e_{uv}^{(t+1)} = 0) | \mathcal{H}^{(t)}\right] - \sum_{u \in \mathcal{U}^{(t)}} e_{uv}^{(t)} f(g(\mathbf{z}_{uv}^{(t)})), \end{aligned}$$

and $\mathcal{H}^{(t)} = \{\mathcal{G}^{(i)}\}_{i=t_0}^t$, then

$$\begin{aligned} &\mathbb{E}[|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}| - \mathcal{S}_{ss}^{(t)}(v) | \mathcal{H}^{(t)}] \\ &= \sum_{u \in \mathcal{U}^{(t)}} e_{uv}^{(t)} \mathbb{E}[\mathbb{1}(e_{uv}^{(t+1)} = 0) | \mathcal{H}^{(t)}, e_{uv}^{(t+1)} = 0] - \sum_{u \in \mathcal{U}^{(t)}} e_{uv}^{(t)} f(g(\mathbf{z}_{uv}^{(t)})) \\ &= \sum_{u \in \mathcal{U}^{(t)}} e_{uv}^{(t)} (Pr((e_{uv}^{(t+1)} = 0 | e_{uv}^{(t)} = 1, \mathcal{H}^{(t)})) - \hat{Pr}(e_{uv}^{(t+1)} = 0 | e_{uv}^{(t)} = 1, \mathcal{H}^{(t)})). \end{aligned} \tag{16}$$

Since the right hand side of Equation 16 is 0, we complete the proof.

Furthermore, we observe that there may exist correlation in the churn behavior between different games and users. As such, we employ link analysis algorithms, which are able to take into consideration the global information and mutual reinforcing effects among different games and users, for inferring the macro-level churn of games. In particular, we choose *PageRank* [18] and *HITS* [19] algorithms due to their wide applications and stable performance. To apply both algorithms, we first construct a user-game relation graph $\mathcal{G}_r^{(t)} = \{\mathcal{V}^{(t)}, \mathcal{E}^{(t)}, \mathcal{W}_r^{(t)}\}$, where $\mathcal{V}^{(t)}$ and $\mathcal{E}^{(t)}$ are copied from the attributed graph $\mathcal{G}^{(t)}$, and $\mathcal{W}_r^{(t)}$ is an adjacency matrix with the element at u -th row and v -th column satisfying

$$\mathcal{W}_r^{(t)}(u, v) = f(g(\mathbf{z}_{uv}^{(t)})). \tag{17}$$

This relation graph will be the input to both algorithms. The ranking scores computed by *PageRank* and *HITS* are referred to as $\mathcal{S}_{pg}^{(t)}(\cdot)$ and $\mathcal{S}_{hits}^{(t)}(\cdot)$, respectively.

The procedure to compute $\mathcal{S}_{pg}^{(t)}(\cdot)$ is given in Algorithm 1. The algorithm assigns a larger ranking score to a game if the game has more about-to-churn players with larger scores. Moreover, the score will be propagated along with a damping factor from one game to another game. For a specific game, the algorithm handles different players differently. In the algorithm, players with distinct propensities of churn contribute differently to the computation of the game’s ranking score. Considering that $\mathcal{G}_r^{(t)}$ is a bipartite graph, in each iteration, the algorithm computes a game’s ranking score by summing up its players’ scores. The sum is weighted by the churn probabilities, where the players with larger churn probabilities are more likely to contribute greater proportions to the game’s ranking score. Meanwhile, this preponderance is normalized by the total number of games the player is playing. The intuition is that the algorithm reduces the proportion contributed by a game if the player is more likely to churn other games at this time point. This is because a player is rarely observed to churn all games at a single time point. In this way, the churn correlation between different games for a specific player is addressed.

Algorithm 1: *PageRank* Based Macro-Level Churn Ranking

Input: Graph $\mathcal{G}_r^{(t)}$, maximal number of iterations M , damping factor d_α
Output: The ranking score $\mathcal{S}_{pg}^{(t)}(v)$ of $v \in \mathcal{V}^{(t)}$

```

1  $k = 0$ , initialize  $\mathcal{T}_k^{(t)}(i) = 1/|\mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}|$  for  $i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}$ ;
2 while  $k < M$  do
3   for  $i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}$  do
4      $\mathcal{T}_{k+1}^{(t)}(i) = \frac{1 - d_\alpha}{|\mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}|} + d_\alpha \times \sum_{j \in \mathcal{N}_i^{(t)}} \mathcal{T}_k^{(t)}(j) \times \frac{\mathcal{W}_r^{(t)}(i, j)}{\sum_{\ell \in \mathcal{N}_i^{(t)}} \mathcal{W}_r^{(t)}(i, \ell)}$ ;
5   end
6 end
7  $\mathcal{S}_{pg}^{(t)}(v) = \mathcal{T}_{k+1}^{(t)}(v)$  for  $v \in \mathcal{V}^{(t)}$ ;
8 return  $\mathcal{S}_{pg}^{(t)}(v)$  for  $v \in \mathcal{V}^{(t)}$ 

```

Algorithm 2: *HITS* Based Macro-Level Churn Ranking

Input: Graph $\mathcal{G}_r^{(t)}$, maximal number of iterations M
Output: The ranking score $\mathcal{S}_{hits}^{(t)}(v)$ of $v \in \mathcal{V}^{(t)}$

```

1  $k = 0$ , initialize  $\mathcal{A}_{hits}^{(t)}(i) = \mathcal{H}_{hits}^{(t)}(i) = 1$  for  $i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}$ ;
2 while  $k \leq M$  do
3   Copy last iteration scores to  $\mathcal{T}^{(t)}$ :  $\mathcal{T}^{(t)}(i) = \mathcal{A}_{hits}^{(t)}(i)$  for  $i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}$ ;
4   for  $i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}$  do
5      $\mathcal{A}_{hits}^{(t)}(i) = \sum_{j \in \mathcal{N}_i^{(t)}} \mathcal{W}_r^{(t)}(i, j) \mathcal{H}_{hits}^{(t)}(j)$ ;
6      $\mathcal{H}_{hits}^{(t)}(i) = \sum_{j \in \mathcal{N}_i^{(t)}} \mathcal{W}_r^{(t)}(i, j) \mathcal{T}^{(t)}(j)$ ;
7   end
8   Normalize  $\{\mathcal{A}_{hits}^{(t)}(i)\}_{i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}}$  and  $\{\mathcal{H}_{hits}^{(t)}(i)\}_{i \in \mathcal{V}^{(t)} \cup \mathcal{U}^{(t)}}$ ;
9 end
10  $\mathcal{S}_{hits}^{(t)}(v) = \mathcal{A}_{hits}^{(t)}(v)$  for  $v \in \mathcal{V}^{(t)}$ ;
11 return  $\mathcal{S}_{hits}^{(t)}(v)$  for  $v \in \mathcal{V}^{(t)}$ 

```

The procedure to compute $\mathcal{S}_{hits}^{(t)}(\cdot)$ is presented in Algorithm 2. The algorithm uses the final authority score $\mathcal{A}_{hits}^{(t)}(v)$ as the ranking score $\mathcal{S}_{hits}^{(t)}(v)$ for game v . Since the graph $\mathcal{G}_r^{(t)}$ is a bipartite graph composed of players and games, a larger authority score is assigned to a game if the game has more players with large hub scores. Similarly, a player has a large hub score when the player is playing many games with large authority scores. In other words, the hub score for a player in some sense describes the chance that a churn will happen with the player. And the authority score of a game to some extent captures the number of its players with high churn propensity. Meanwhile, for a specific game, the algorithm handles different players differently, and for a specific player, different games contribute different proportions. This is reflected by Lines 5 – 6 of Algorithm 2, where

the computation of the authority score and the hub score are weighted by the estimated probability of churn.

4. Experimental Evaluation

In this section, we conduct a comprehensive experimental evaluation over the large-scale real data collected from the Samsung Game Launcher platform. We compare our semi-supervised model (referred to as **SS** in the sequel) with the following state-of-the-art models for mobile game churn prediction:

- **LR**: the logistic regression based solution used in [9, 13, 17]
- **RS**: the supervised variant of our model, in which the loss function contains only the supervised component and the regularization term
- **DT**: a decision tree based solution
- **RF**: the random forests based solution used in [9, 17]
- **SVM**: the SVM based solution used in [9, 13].

In the experiments, we consider the churn duration $T = 14$, but again the proposed solution is not restricted to any particular choice of T .

4.1. Dataset and Feature/Label Construction

Two anonymous datasets were collected *independently* from the Samsung Game Launcher platform within a 4-month period (from August 1st, 2017 to November 30th, 2017) with users’ consent, one from users in *USA* and the other from users in *Korea*. We summarize the key statistics of these two datasets in Table 2. After some investigation of the datasets, we find that the players from US and Korea have very different tastes of mobile games. In our datasets, there are only 18 common games between the top 100 games played by US users and the top 100 games played by Korean users. One major reason is that the two countries have different cultures, different life styles, and different native languages. Games interesting Korean users do not necessarily interest US users. Representing users and games from the two countries in one relation graph faces the risk that the resulting graph is highly disconnected, which in turn increases the complexity of graph embedding. Moreover, as shown in Table 2, the statistics of the datasets from the two countries differ, e.g., the numbers of records and users from Korea are larger than those from US. Training a single model for both datasets may make the parameter update of one of them dominate the other in gradient descent. Therefore, we use the same DNN prediction architecture for both datasets but train them separately. We would also like to note that training multiple vertical prediction models for different regions or countries is also a common practice in industry.

The collected data contains three major types of information: (1) play history, (2) game profiles, and (3) user information. Each play record in the play history contains an anonymous user id, a game package name, and the timestamp of play. It is also accompanied with rich contextual information, such as WiFi connection status, screen brightness, audio volume, etc. Game profiles are collected from different game stores, which include features like genre, developer, number of downloads, rating values, number of ratings, etc. User information contains the

Table 2. Dataset statistics

Dataset	# of users	# of games	# of play records
<i>USA</i>	15,000	19,705	76,468,301
<i>Korea</i>	25,000	18,470	106,544,313

device model, region, OS version, etc. However, data within games (e.g., levels) is not available due to privacy concerns.

Features and labels need to be generated with care in order to avoid data leakage. In general, labels and features are taken from disjoint periods to avoid data leakage. The semi-supervised model is trained based on features and labels within the observation period. Labels indicate whether a player churns a game on that day. Features are constructed from historical data before the day of prediction. The training set and test set are split by label days in chronological order [27], for example, taking the labeled data in the first 2/3 of the observation period as the training set and the remaining 1/3 as the test set. This is also to ensure that there is a time difference between the test set and the training set. The semi-supervised model along with the *SimSum*, *PageRank* and *HITS* methods are evaluated on the test datasets.

4.2. Experimental Settings

We tune the hyperparameters, including learning rate, batch size, regularization terms, number of layers and number of neurons per layer, based on the model performance on the test datasets. To determine the parameters for micro-level churn prediction, we follow what we did in our previous work [28]. A grid search on these parameters is performed and the combination yielding the best performance is chosen. The regularization parameters $\{\lambda_i\}_{i=0}^4$ are all set to 1. α , β , and γ in Equation (1) are set to 0.02, 0.01, 1e-5, respectively. ϵ , p , and q discussed in Section 3.4 are set to 1, 1, and 0.05, respectively. The parameters used for training the deep neural network are summarized below. The maximal number of iterations M for the *PageRank* and *HITS* methods is set to 100. The damping factor for *PageRank* is set to 0.85. We summarize the parameter settings of the two datasets below.

Parameter settings of the *USA* dataset:

- Player feature dimension: 10,042
- Game feature dimension: 10,042
- Player-game feature dimension: 30
- Learning rate: the initial value η_0 is 0.017 and decay by $\eta = \eta_0/(1 + k/2)$, where k is the number of epochs
- Number of neurons: input layers 30, embedding layers 50, output layers 380K
- Number of epochs: 6-8
- Batch size: 1,024
- Context number per user-game pair: 4
- Optimizer: Adam method [29]
- Activation function: rectified linear unit (ReLU)

Table 3. Performance of different micro-level churn prediction models on *USA*

Model	AUC	Recall	Precision
SS	<u>0.82</u>	<u>0.78</u>	0.32
RS	0.77	0.75	0.27
LR	0.66	0.38	0.26
DT	0.59	0.28	0.32
RF	0.75	0.31	<u>0.41</u>
SVM	0.61	0.78	0.18

Table 4. Performance of different micro-level churn models on *Korea*

Model	AUC	Recall	Precision
SS	<u>0.82</u>	<u>0.70</u>	0.34
RS	0.76	0.70	0.25
LR	0.67	0.59	0.21
DT	0.58	0.26	0.30
RF	0.73	0.27	<u>0.42</u>
SVM	0.63	0.67	0.18

Parameter settings of the *Korea* dataset:

- Player feature dimension: 10,042
- Game feature dimension: 10,042
- Player-game feature dimension: 30
- Learning rate: the initial value η_0 is 0.019 and decay by $\eta = \eta_0/(1 + k/2)$, where k is the number of epochs
- Number of neurons: input layers 30, embedding layers 50, output layers 632K
- Number of epochs: 8-12
- Batch size: 4,096
- Context number per user-game pair: 4
- Optimizer: Adam method [29]
- Activation function: rectified linear unit (ReLU)

4.3. Experimental Results*4.3.1. Micro-Level Churn Prediction*

We use three widely-used evaluation metrics to compare the performance of different micro-level churn prediction models. The most important metric with respect to the business goals is *the area under the ROC curve* (AUC). Following previous studies [9, 13, 17], we also consider *precision* and *recall*. Accuracy is not

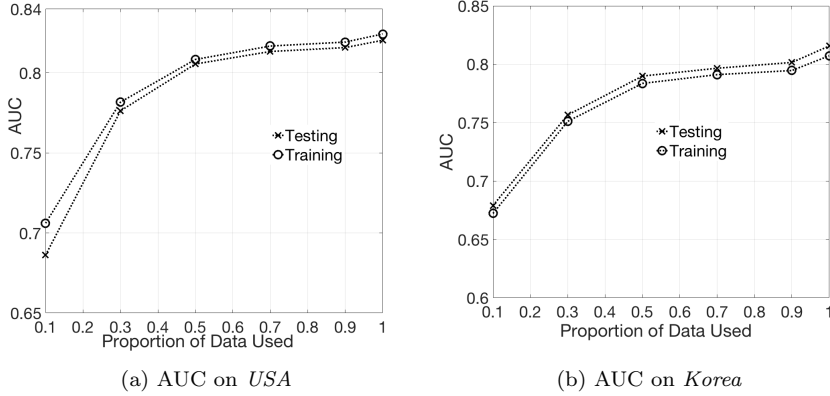


Fig. 6. Comparison of AUC for SS in training and testing

used because our data is imbalanced with around 85% negative instances in the *USA* dataset and 86% negative instances in the *Korea* dataset.

We report the main experimental results in Table 3 and Table 4. It can be observed that in general our model achieves the best AUC and recall on both datasets. Our model outperforms all single models (i.e., LR, DT and SVM) in terms of all the three metrics. In particular, it is worth mentioning that SVM achieves a high recall at the cost of a very low precision. This is because it makes a large number of false positives. This fact makes it less useful for business decision making. Compared with the ensemble method RF, which has been considered so far the best method in the field, our model still achieves 34% AUC improvement and 250% recall improvement. RF achieves the highest precision on the datasets. However, we point out that this number is actually misleading because it can easily overfit and only recognize a small proportion of churn labels. This is partially evidenced by its poor recall, which makes it difficult to meet business requirements of churn prediction (e.g., targeted promotion campaigns). We also carefully compare the AUC of SS in training and testing in Fig. 6. Since the curves are very close, it can be learned that our model is neither overfitting nor underfitting.

The performance difference between SS and RS justifies the benefits of incorporating unsupervised loss and temporal loss in the objective function. We provide a further comparison between SS and RS with respect to the number of epochs in Fig. 7. Both models take 5-7 epochs to reach relatively stable performance. We observe that SS outperforms RS in general under different numbers of epochs and for both Korean users and USA users.

Since the architecture of our DNN is novel and unique (e.g., contain both supervised outputs and unsupervised outputs), we expose more details on how we choose the parameters and train the model. A comparison of AUC in training under different learning rates is given in Fig. 8. It can be observed that the choice of the learning rate greatly influences the model performance after the initial epoch. We experimentally find that 0.1 is too large for the learning rate, which makes the step in gradient descent too large to find a good minimum and that 0.001 is too small, making it converge very slowly to the optimal point. Therefore we experimentally test learning rates between 0.1 and

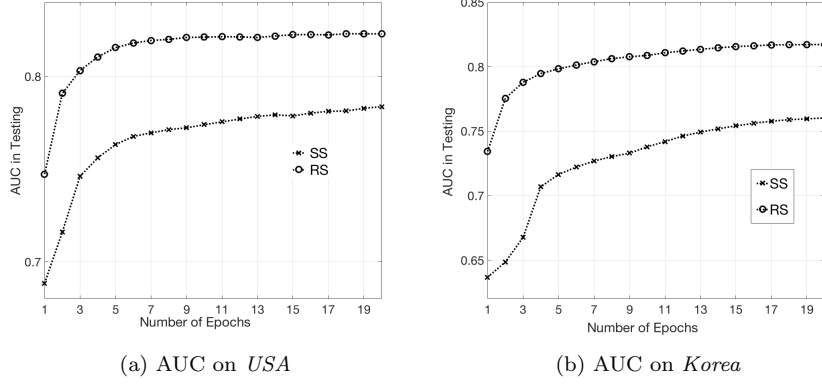


Fig. 7. AUC comparison between SS and RS under different numbers of epochs

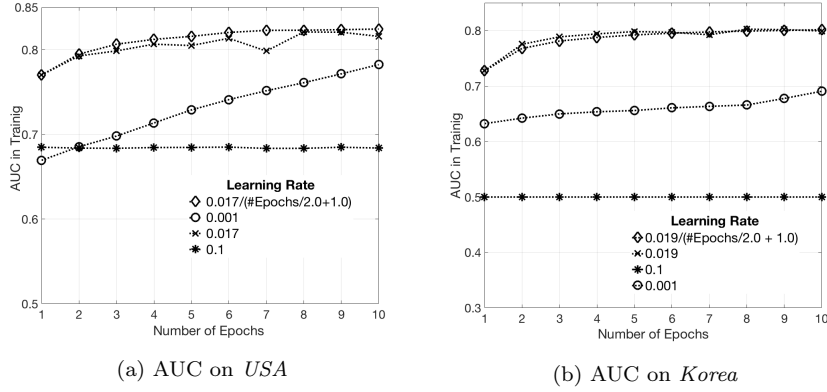


Fig. 8. Comparison of AUC for SS in training under different learning rates

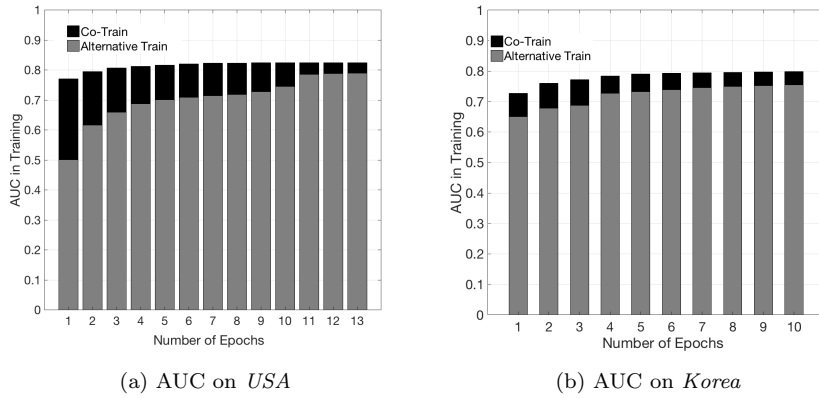


Fig. 9. Comparison of AUC for SS in training under different training methods

0.001 and find that $0.017/(\#Epochs/2+1)$ works best for training on *USA* while $0.019/(\#Epochs/2+1)$ works best for training on *Korea*.

For the supervised and unsupervised components, we try two different training methods: *co-train* and *alternative train*. Co-train means that we simultaneously train the supervised loss function and the unsupervised loss function; alternative train means that we alternately train the unsupervised component with the unsupervised loss function and the supervised component with the supervised loss function. Alternative train is a widely-used training method for similar structures [20, 21]. It is interesting to observe that, however, co-train outperforms alternative train in terms of AUC under different numbers of epochs as shown in Fig. 9. Therefore, we choose co-train as the final training method in our experiments.

4.3.2. Macro-Level Churn Ranking

We use five widely-used evaluation metrics to compare the performance of different macro-level churn ranking methods including *Kendall’s Tau correlation coefficient*, *weighted Kendall’s Tau correlation coefficient*, *Spearman correlation coefficient*, *Average Precision at K* and *Mean Average Precision (MAP)*.

In statistics, the Kendall’s Tau coefficient, also known as Kendall rank correlation coefficient, the weighted Kendall’s tau coefficient and the Spearman coefficient are all measures of rank correlation: the similarity of two orderings on the same data. For instance, Kendall’s Tau coefficient [30] is defined as:

$$\text{Kendall’s Tau} := \frac{P - Q}{n(n - 1)/2},$$

where P is the number of concordant pairs in two orderings, Q is the number of discordant pairs, and n is the number of elements in the ranked list. In our problem setting, one ordering is obtained by ranking the games based on one of the ranking scores, namely $\mathcal{S}_{ss}^{(t)}(v)$, $\mathcal{S}_{hits}^{(t)}(v)$, or $\mathcal{S}_{pg}^{(t)}(v)$, and the other ordering is obtained by ranking the games based on the ground truth $|\mathcal{N}_v^{(t)} \setminus \mathcal{N}_v^{(t+1)}|$. The weighted Kendall’s Tau is a weighted version of Kendall’s Tau, in which exchanges of high rankings are more influential than exchanges of low rankings, e.g., in [31] an exchange to break the ties between elements with rank r and s has weight $1/(r + 1) + 1/(s + 1)$. Spearman coefficient, on the other hand, quantifies how well the relationship between two ranked lists can be described using a monotonic function. Unlike Pearson correlation that can only assess linear relationships, Spearman correlation considers both linear and nonlinear monotonic relationships. The range of all three evaluation metrics is between -1 and 1 . Having a value closer to 1 indicates that the ranking of the games is more similar to that based on the ground truth.

Since macro-level churn ranking is based on the results of micro-level churn prediction, we compare the performance of different combinations of the two parts. Each combination is referred to in the “macro-micro” format. For example, *SimSum-SS* represents the combination of using *SimSum* as the macro-level churn ranking method and our semi-supervised method as the micro-level churn prediction model. Fig. 10, Fig. 11, and Fig. 12 present the performance of different combinations in terms of Kendall’s Tau, weighted Kendall’s Tau and Spearman coefficient, respectively. Each performance metric is experimentally evaluated on the same test datasets described in Section 4.1. Since micro-level

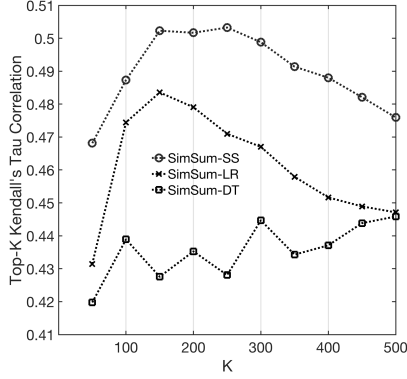
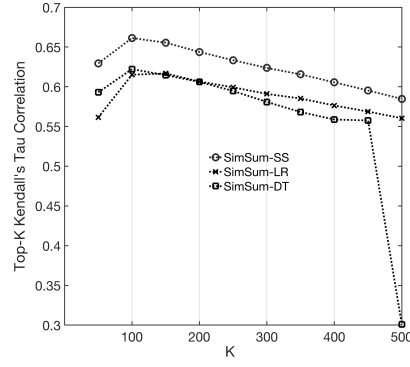
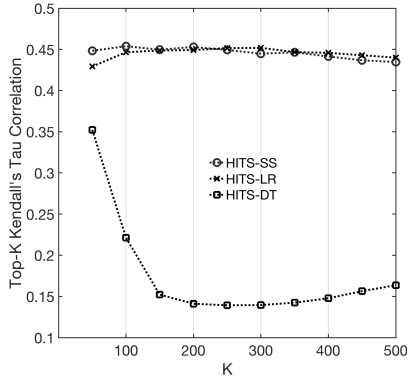
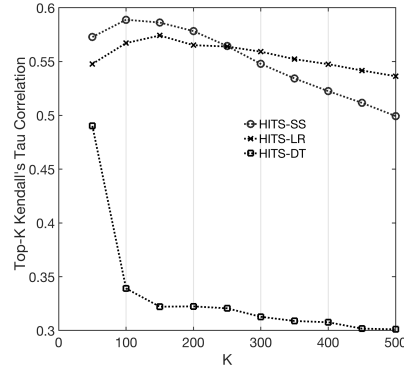
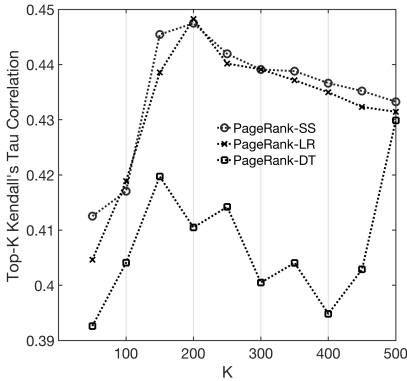
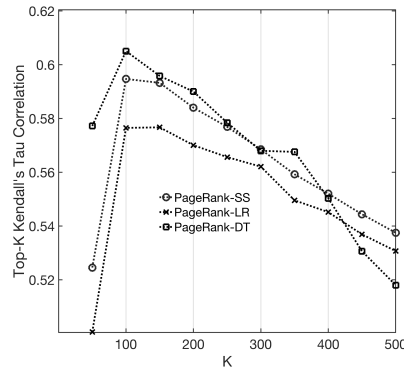
(a) *SimSum* Kendall's Tau on *USA*(b) *SimSum* Kendall's Tau on *Korea*(c) *HITS* Kendall's Tau on *USA*(d) *HITS* Kendall's Tau on *Korea*(e) *PageRank* Kendall's Tau on *USA*(f) *PageRank* Kendall's Tau on *Korea*

Fig. 10. Comparison of Kendall's Tau correlation coefficients in testing under different macro-level churn ranking and micro-level churn prediction methods

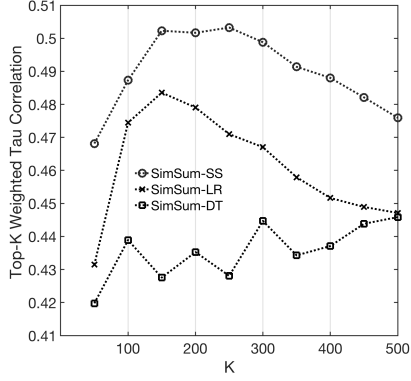
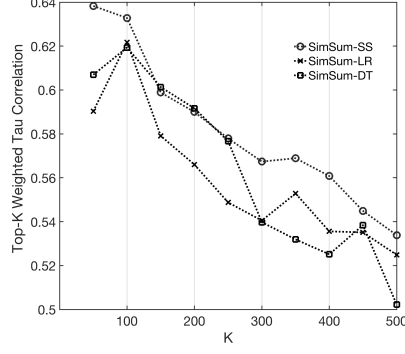
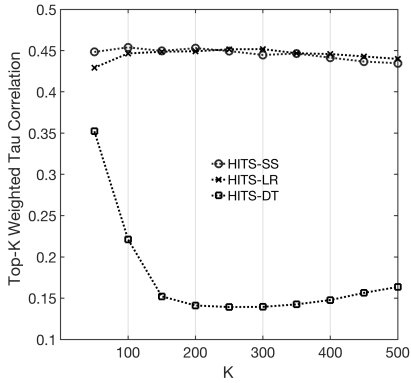
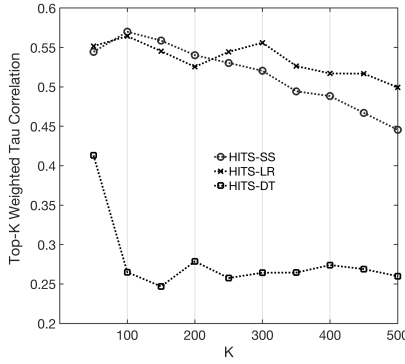
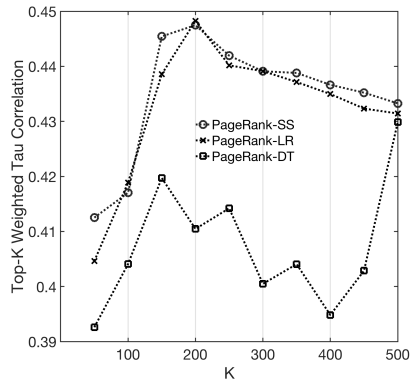
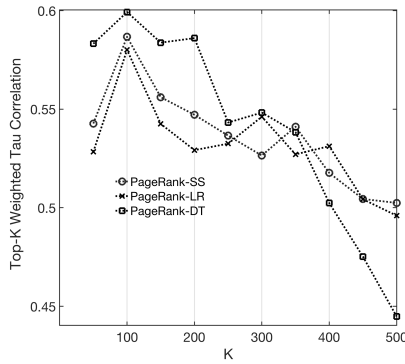
(a) *SimSum* Weighted Kendall's Tau on USA(b) *SimSum* Weighted Kendall's Tau on Korea(c) *HITS* Weighted Kendall's Tau on USA(d) *HITS* Weighted Kendall's Tau on Korea(e) *PageRank* Weighted Kendall's Tau on USA(f) *PageRank* Weighted Kendall's Tau on Korea

Fig. 11. Comparison of Weighted Kendall's Tau correlation coefficients in testing under different macro-level churn ranking and micro-level churn prediction methods

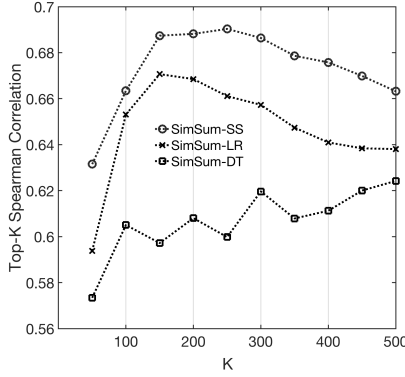
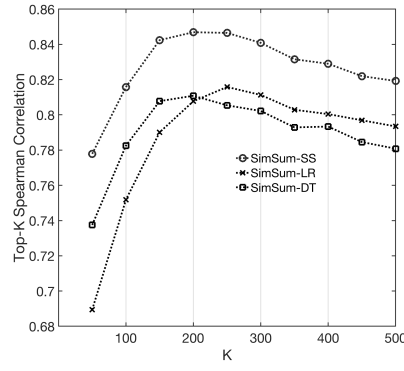
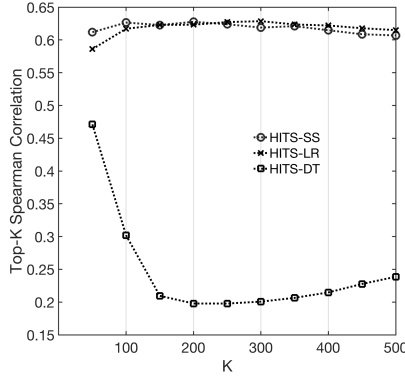
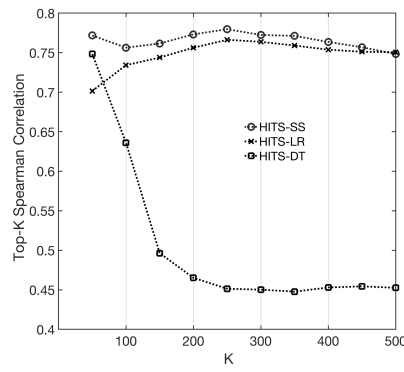
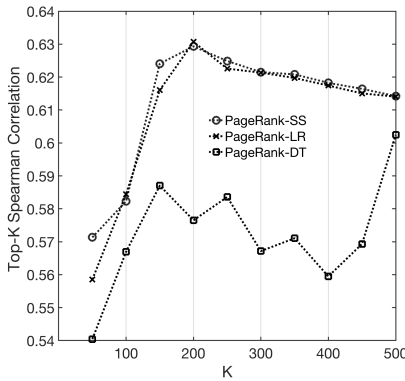
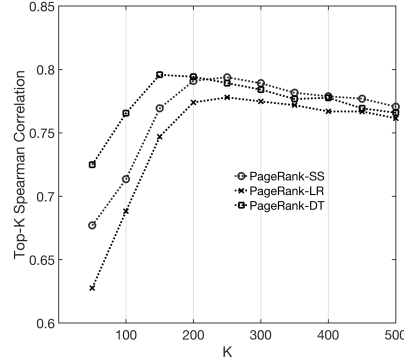
(a) *SimSum* Spearman Coefficient on *USA*(b) *SimSum* Spearman Coefficient on *Korea*(c) *HITS* Spearman Coefficient on *USA*(d) *HITS* Spearman Coefficient on *Korea*(e) *PageRank* Spearman Coefficient on *USA*(f) *PageRank* Spearman Coefficient on *Korea*

Fig. 12. Comparison of Spearman correlation coefficients in testing under different macro-level churn ranking and micro-level churn prediction methods

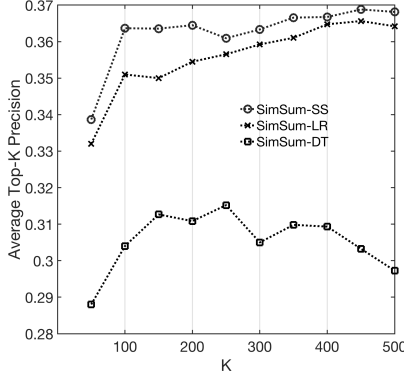
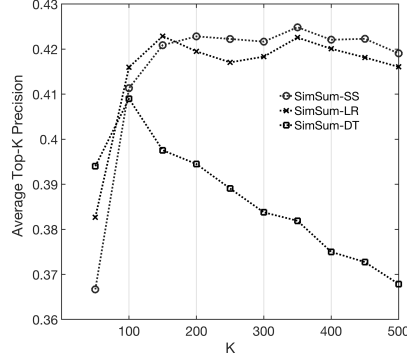
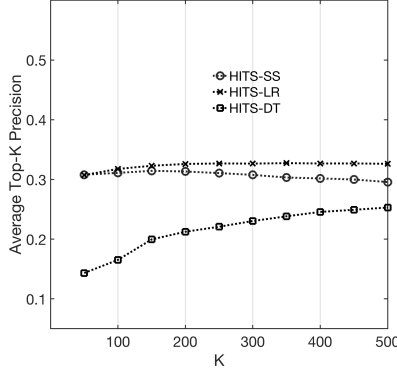
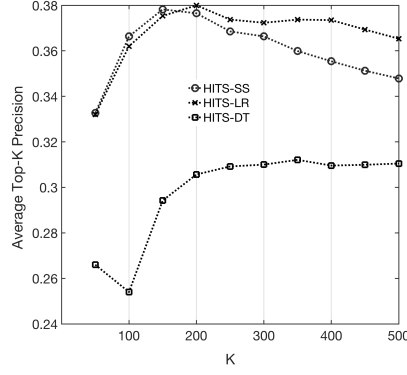
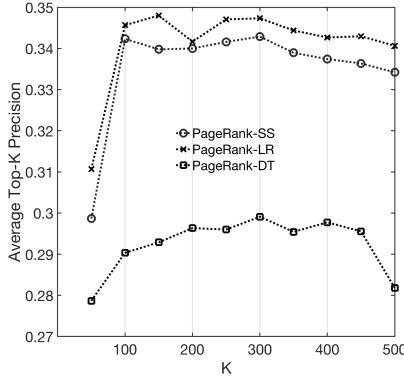
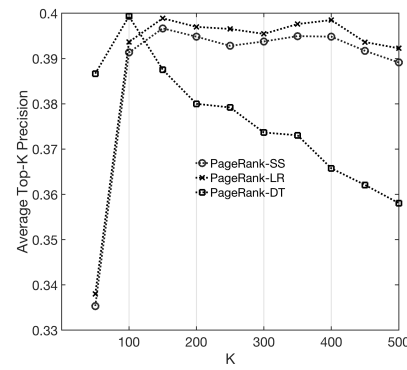
(a) *SimSum* Average Precision at K on *USA*(b) *SimSum* Average Precision at K on *Korea*(c) *HITS* Average Precision at K on *USA*(d) *HITS* Average Precision at K on *Korea*(e) *PageRank* Average Precision at K on *USA*(f) *PageRank* Average Precision at K on *Korea*

Fig. 13. Comparison of Average Precision at K in testing under different macro-level churn ranking and micro-level churn prediction methods

Table 5. Mean Average Precision under different micro-level churn prediction and macro-level churn ranking methods on *Korea* and *USA*

Area-Model	<i>SimSum</i>	<i>HITS</i>	<i>PageRank</i>
<i>Korea</i> -SS	<u>0.41</u>	<u>0.40</u>	<u>0.40</u>
<i>Korea</i> -LR	0.40	0.39	0.39
<i>Korea</i> -DT	0.40	0.40	0.39
<i>USA</i> -SS	<u>0.41</u>	<u>0.40</u>	<u>0.38</u>
<i>USA</i> -LR	0.38	0.38	0.37
<i>USA</i> -DT	0.38	0.38	0.38

churn prediction is performed daily, macro-level churn ranking is also calculated every day. The performance reported in the figures is the average over the testing period. As the ranked list of games can be as long as 20,000, we only report the performance on the top K games where the maximal K value is set to 500. It can be observed that in most cases the macro-level churn ranking methods in combination with our semi-supervised method outperform those in combination with different baselines (i.e., LR, DT), especially when the value of K is small. This is because SS achieves better performance in churn prediction than other baselines. This indicates that the quality of the churn prediction results is important for subsequent churn ranking, which takes the predicted churn probabilities as inputs. Moreover, we find that based on the Kendall’s Tau, Weighted Kendall’s Tau and Spearman coefficient, the performance of *HITS* is more sensitive to the quality of churn prediction results. This may be because in the iteration, *HITS* normalizes the scores over the sum of scores across all nodes, which can be large and lead to a skewed score distribution. In contrast, *SimSum* and *PageRank* normalize the scores based on only the scores of neighbors, which is more robust.

We also evaluate the performance of different churn ranking methods by two widely-adopted metrics in the recommendation domain: Average Precision at K and Mean Average Precision (MAP). Precision at K corresponds to the percentage of relevant results in the top K games of the ranked list. A game at position i in the ranked list based on our ranking scores is considered relevant if it is in the top i games of the ranked list based on the ground truth. The comparison is reported in Fig. 13. Similarly, the scores reported are the average of all days in the testing period. In Fig. 13, we observe similar patterns to those observed under the previous statistic metrics. In most cases, different macro-level churn ranking methods in combination with the proposed semi-supervised mode SS achieve better performance than the combinations with other baselines. This again indicates the importance of the positive interaction between micro-level churn prediction and macro-level churn ranking. One shortcoming of this metric is that it fails to take into account the positions of relevant games. To overcome this limitation, we further adopt Mean Average Precision [32], in which the relevance score is weighted by the position before being averaged. We report the comparison of the methods in terms of MAP in Table 5. Similar conclusions can be drawn from the results. It is interesting to note that after putting larger weights on relevant games with higher ranking, the difference of the sensitivity to churn prediction quality between *HITS* and the other two methods is reduced.

5. Related Work

In this section, we review two categories of existing works that are relevant to this paper. The first category includes the existing works for (game) churn analysis. The early works [7, 9, 8, 11, 13, 12] are based on more traditional machine learning models, such as logistic regression, random forests, SVM, naive Bayes, etc., and are experimentally evaluated on extremely small numbers of games (i.e., less than five). As shown in our experiments, their performance on large-scale real data with tens of thousands of mobile games and hundreds of millions of user-app interactions is generally not satisfactory. In addition, we also observe scalability issues when they are applied to large-scale data. Some recent research has started to use more advanced models. [10, 14, 16] propose to use survival model for churn prediction, in which churn probability is modeled as a function of playtime. Kim et al. [17] achieve good performance by using convolutional neural networks (CNN) and long short-term memory networks (LSTM). There are also several recent deep-learning-based studies [33, 34, 35] for non-game churn prediction problems, which report better performance. This motivates us to employ deep neural network models. While being a generic solution, our model is able to accommodate the unique characteristics of mobile games. We provide a comparison between all existing works and ours in Table 6. Much of the research relevant to macro-level churn ranking pays particular attention to popularity prediction, the orthogonal to the churn ranking problem. However, they solve the problem from scratch and thus require efforts on feature engineering [4] or a complex model [5]. In contrast, we propose to reuse estimated churn probabilities of the micro-level churn prediction task to reduce the overhead.

The second category contains recent works on graph embedding. The availability of temporal player-game relation graphs $\{\mathcal{G}_r^{(i)}\}_{i=t_0}^t$ motivates us to make use of the graph embedding technique in data representation. Graph embedding has been recently applied to many large-scale machine learning applications, demonstrating strong performance. It automates the entire process of feature engineering by casting feature extraction as a representation learning problem. It frees models from human bottlenecks due to handcrafted features and is able to utilize the full richness of data. The churn prediction task is trained not with the original, maybe noisy, handcrafted features but with more informative latent features generated by the embedding technique. Moreover, it reduces the dimensions of the feature vectors used in training as well as the corresponding computational complexity. For example, the dimension of the original feature vectors in our problem is 10,042, while in the latent space, the dimension is reduced to 30, which is much smaller and more efficient for the downstream prediction task. In addition, it is an inductive framework and is able to embed a new player-game relationship into a latent vector close to the embedding vectors of similar player-game relationships. This enables the framework to capture the churn tendency of a player-game pair even when the churn of the pair was never observed before. However, most existing works such as Node2Vec [25], Deep Walk [36] and LINE [37] are *node-centric* embedding. When it comes to game churn prediction, the entities to be embedded are edges that represent the relationships between players and games. Very limited work has studied edge embedding. Abu-El-Haija et al. [38] propose to model edge embedding as a function of node embedding, in which the two ends of an edge are first embedded and then passed into a deep neural network with edge embedding as output. This method embeds edges in

Table 6. Comparison between this paper and existing works in data size and key techniques

Paper	Data size	Key techniques
[7]	5 games 50 thousand users	Decision tree, naive Bayes
[8]	2 games 10 thousand users	Hidden Markov Model combined with a single layer neural network
[10, 14]	1 game 3 thousand users	Survival ensembles
[16]	1 game 1 thousand users	Survival model
[11]	1 game 10 thousand users	Hidden Markov Model
[9]	2 games 1 thousand users	SVM, decision tree, logistic regression
[12]	1 game 130 thousand users	Heuristic decision tree
[13]	3 games 60 thousand users	Logistic regression, decision tree, SVM
[17]	3 games 200 thousand users	Logistic regression, random forests, CNN, LSTM
Ours	40,000 games 40 thousand users	Deep attributed edge embedding

an indirect way and does not take into account any attribute information. In addition, these models [25, 36, 37] all belong to the transductive framework, in which embedding cannot be generated if an object has never appeared in training. However, in our problem new users and new games appear continuously; new relationships between existing users and games may form anytime in the future. Therefore, for game churn prediction the capability of handling new edges is indispensable.

Several very recent works have proposed the idea of inductive graph embedding [20, 22, 21, 39], inspired by which we propose our inductive edge embedding model for micro-level churn prediction. Our model improves these works in several major ways. First, our embedded features are learned in a *semi-supervised* manner, where the supervised component is for churn prediction while the unsupervised component is for context recovery. Compared to unsupervised methods, embedding features learned in a semi-supervised way have been shown to achieve better performance [21]. Second, our model captures graph dynamics by imposing temporal loss to the embeddings of the same edge in consecutive graph snapshots. Unlike all existing works, where embeddings represent structures or attribute information, embeddings in our model are designed to simultaneously capture contexts and graph dynamics. To our best knowledge, our model is the first to achieve such benefits.

6. Conclusion

Churn analysis of mobile games is a vital research and business problem that is backed up by a billion-dollar market. In this paper, we consider two important aspects of churn analysis of mobile games: micro-level churn prediction and macro-level churn ranking. We proposed a novel inductive semi-supervised embedding model for large-scale micro-level game churn prediction. It jointly learns a prediction function and an edge embedding function, the latter of which can automatically map high-dimensional raw feature vectors to more informative latent feature vectors. We modeled the prediction function and the embedding function by deep neural networks, where the embedding component is designed to capture both contextual information and relationship dynamics. The contexts of an edge are sampled by a novel attributed random walk technique, which considers both topological adjacency and attribute similarities. To solve the macro-level churn ranking problem without much overhead, we proposed a simple method *SimSum* based on the outputs of the micro-level churn prediction model and adapted the link analysis algorithms *PageRank* and *HITS* to further take into consideration the global information and mutual reinforcing effects among different games and users. We compared our approaches with several state-of-the-art baseline methods on large-scale real-world data collected from the Samsung Game Launcher platform. The extensive experimental results clearly demonstrate the effectiveness of our proposed solutions.

Although the paper focuses on mobile game churn analysis, the proposed methods are not restricted to this specific problem and also work for more general problems. Since the inputs of our model only contain the attributes of objects and their relationships at different timestamps, the proposed model can be generalized to fulfill any churn analysis task where the underlying data can be modeled in a similar way, for instance, customer disengagement prediction in membership business (e.g., Apple Music, Costco, and insurance companies) and interest group unsubscribe prediction in social networks (e.g., Facebook, Meetup, and Google+).

References

- [1] NewZoo. Mobile revenues account for more than 50% of the global games market as it reaches 137.9 billion in 2018. <https://newzoo.com/insights/articles/global-games-market-reaches-137-9-billion-in-2018-mobile-games-take-half/>, 2018. Accessed: 2019-01-09.
- [2] Cameron Summerson. What is samsung’s game launcher, and should you use it? <https://www.howtogeek.com/262157/what-is-samsungs-game-launcher-and-should-you-use-it/>, 2016. Accessed: 2019-01-09.
- [3] Mohamed Zaki, Dalia Kandeil, Andy Neely, and Janet R. McColl-Kennedy. The fallacy of the net promoter score: Customer loyalty predictive model. *Cambridge Service Alliance*, pages 1–25, 2016.
- [4] Eric Malmi. Quality matters: Usage-based app popularity prediction. In *Proceedings of ACM International Joint Conference on Pervasive and Ubiquitous Computing: Adjunct Publication (UbiComp Adjunct)*, pages 391–396, 2014.
- [5] Hengshu Zhu, Chuanren Liu, Yong Ge, Hui Xiong, and Enhong Chen. Popularity modeling for mobile apps: A sequential approach. *IEEE Transactions on Cybernetics*, 45(7):1303–1314, 2015.

- [6] Da Cao, Liqiang Nie, Xiangnan He, Xiaochi Wei, Jialie Shen, Shunxiang Wu, and Tat-Seng Chua. Version-sensitive mobile app recommendation. *Information Sciences*, 381:161–175, 2017.
- [7] Fabian Hadiji, Rafet Sifa, Anders Drachen, Christian Thureau, Kristian Kersting, and Christian Bauckhage. Predicting player churn in the wild. In *Proceedings of IEEE Conference on Computational Intelligence and Games (CIG)*, 2014.
- [8] Julian Runge, Peng Gao, Florent Garcin, and Boi Faltings. Churn prediction for high-value players in casual social games. In *Proceedings of IEEE Conference on Computational Intelligence and Games (CIG)*, 2014.
- [9] Hanling Xie, Sam Devlin, Daniel Kudenko, and Peter Cowling. Predicting player disengagement and first purchase with event-frequency based data representation. In *Proceedings of IEEE Conference on Computational Intelligence and Games (CIG)*, 2015.
- [10] África Periañez, Alain Saas, Anna Guitart, and Colin Magne. Churn prediction in mobile social games: Towards a complete assessment using survival ensembles. In *Proceedings of IEEE International Conference on Data Science and Advanced Analytics (DSAA)*, pages 564–573, 2016.
- [11] Marco Tamassia, William Raffe, Rafet Sifa, Anders Drachen, Fabio Zambetta, and Michael Hitchens. Predicting player churn in destiny: A hidden markov models approach to predicting player departure in a major online game. In *Proceedings of IEEE Conference on Computational Intelligence and Games (CIG)*, 2016.
- [12] Anders Drachen, Eric Thurston Lundquist, Yungjen Kung, Pranav Rao, Rafet Sifa, Julian Runge, and Diego Klabjan. Rapid prediction of player retention in free-to-play mobile games. In *Proceedings of AAAI Conference on Artificial Intelligence and Interactive Digital Entertainment (AIIDE)*, pages 23–29, 2016.
- [13] Hanling Xie, Sam Devlin, and Daniel Kudenko. Predicting disengagement in free-to-play games with highly biased data. In *Proceedings of AAAI Conference on Artificial Intelligence and Interactive Digital Entertainment (AIIDE)*, pages 143–150, 2016.
- [14] Paul Bertens, Anna Guitart, and África Periañez. Games and big data: A scalable multi-dimensional churn prediction model. In *Proceedings of IEEE Conference on Computational Intelligence and Games (CIG)*, 2017.
- [15] Markus Viljanen, Antti Airola, Anne-Maarit Majanoja, Jukka Heikkonen, and Tapio Pahikkala. Measuring player retention and monetization using the mean cumulative function. *arXiv preprint arXiv:1709.06737*, 2017.
- [16] Markus Viljanen, Antti Airola, Jukka Heikkonen, and Tapio Pahikkala. Playtime measurement with survival analysis. *IEEE Transactions on Games*, 10(2):128–138, 2017.
- [17] Seungwook Kim, Daeyoung Choi, Eunjung Lee, and Wonjong Rhee. Churn prediction of mobile and online casual games using play log data. *PloS one*, 12(7):1–19, 2017.
- [18] Lawrence Page, Sergey Brin, Rajeev Motwani, and Terry Winograd. The pagerank citation ranking: Bringing order to the web. Technical report, Stanford InfoLab, 1999.
- [19] Jon M. Kleinberg. Authoritative sources in a hyperlinked environment. *Journal of the ACM (JACM)*, 46(5):604–632, 1999.
- [20] Jiongqian Liang, Peter Jacobs, and Srinivasan Parthasarathy. Seano: Semi-supervised embedding in attributed networks with outliers. In *Proceedings of SIAM International Conference on Data Mining (SDM)*, 2018.
- [21] Zhilin Yang, William W. Cohen, and Ruslan Salakhutdinov. Revisiting semi-supervised learning with graph embeddings. In *Proceedings of International Conference on Machine Learning (ICML)*, pages 40–48, 2016.
- [22] Will Hamilton, Zhitao Ying, and Jure Leskovec. Inductive representation learning on large graphs. In *Proceedings of Conference on Neural Information Processing Systems (NIPS)*, 2017.
- [23] Localytics. Mobile apps: What’s a good retention rate? <http://info.localytics.com/blog/mobile-apps-whats-a-good-retention-rate>, 2018. Accessed: 2019-01-09.
- [24] Margaret C. Wu and Raymond J. Carroll. Estimation and comparison of changes in the presence of informative right censoring by modeling the censoring process. *Biometrics*, 44(1):175–188, 1988.

- [25] Aditya Grover and Jure Leskovec. node2vec: Scalable feature learning for networks. In *Proceedings of ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, pages 855–864, 2016.
- [26] Nesreen K. Ahmed, Ryan A. Rossi, Rong Zhou, John Boaz Lee, Xiangnan Kong, Theodore L. Willke, and Hoda Eldardiry. A framework for generalizing graph-based representation learning methods. *arXiv preprint arXiv:1709.04596*, 2017.
- [27] Benjamin Paul Chamberlain, Angelo Cardoso, Chak H. Liu, Roberto Pagliari, and Marc Peter Deisenroth. Customer lifetime value prediction using embeddings. In *Proceedings of ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, pages 1753–1762, 2017.
- [28] Xi Liu, Muhe Xie, Xidao Wen, Rui Chen, Yong Ge, Nick Duffield, and Na Wang. A semi-supervised and inductive embedding model for churn prediction of large-scale mobile games. In *Proceedings of IEEE International Conference on Data Mining (ICDM)*, 2018.
- [29] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *Proceedings of International Conference on Learning Representations (ICLR)*, 2015.
- [30] Maurice G. Kendall. The treatment of ties in ranking problems. *Biometrika*, 33(3):239–251, 1945.
- [31] Sebastiano Vigna. A weighted correlation index for rankings with ties. In *Proceedings of International World Wide Web Conference (WWW)*, pages 1166–1176, 2015.
- [32] Wanhua Su, Yan Yuan, and Mu Zhu. A relationship between the average precision and the area under the roc curve. In *Proceedings of International Conference on The Theory of Information Retrieval (ICTIR)*, pages 349–352, 2015.
- [33] Artit Wangperawong, Cyrille Brun, Olav Laudy, and Rujikorn Pavasuthipaisit. Churn analysis using deep convolutional neural networks and autoencoders. *arXiv preprint arXiv:1604.05377*, 2016.
- [34] V. Umayaparvathi and K. Iyakutti. Automated feature selection and churn prediction using deep learning models. *International Research Journal of Engineering and Technology (IRJET)*, 4(3):1846–1854, 2017.
- [35] Helder Martins. Predicting user churn on streaming services using recurrent neural networks. Master’s thesis, KTH Royal Institute of Technology, 2017.
- [36] Bryan Perozzi, Rami Al-Rfou, and Steven Skiena. Deepwalk: Online learning of social representations. In *Proceedings of ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, pages 701–710, 2014.
- [37] Jian Tang, Meng Qu, Mingzhe Wang, Ming Zhang, Jun Yan, and Qiaozhu Mei. Line: Large-scale information network embedding. In *Proceedings of International World Wide Web Conference (WWW)*, pages 1067–1077, 2015.
- [38] Sami Abu-El-Haija, Bryan Perozzi, and Rami Al-Rfou. Learning edge representations via low-rank asymmetric projections. In *Proceedings of ACM International Conference on Information and Knowledge Management (CIKM)*, pages 1787–1796, 2017.
- [39] Xi Liu, Ping-Chun Hsieh, Nick Duffield, Rui Chen, Muhe Xie, and Xidao Wen. Real-time streaming graph embedding through local actions. In *Companion Proceedings of The 2019 World Wide Web Conference*, pages 285–293. ACM, 2019.

Author Biographies



Xi Liu is a Computer Engineering PhD candidate at Texas A&M University and received B.E. at Harbin Institute of Technology. He has published papers in ICML, ICDM, and WWW. His research interests include bandit learning, representation learning, deep reinforcement learning and their application in recommendation systems.



Muhe Xie received his M.S. degree in data science from New York University. He is currently a software engineer at Samsung Research America. His research interests include machine learning and deep learning. He published a paper in ICDM.



Xidao Wen is a Ph.D candidate at PITT Computational Social Science Lab (Picso), from the School of Computing and Information, University of Pittsburgh. His research interests include data mining, machine learning and Information Visualization.



Rui Chen received the Ph.D. degree in Computer Science from Concordia University. He is currently the Head of Data Science at Samsung Research America, leading the R&D of production-scale machine/deep learning solutions for multiple Samsung services. Before that, he was a Research Assistant Professor in the Department of Computer Science at Hong Kong Baptist University and a Postdoctoral Fellow at the University of British Columbia. His research interests include machine learning, data privacy and databases. He has published more than 40 technical papers in top venues, including CSUR, VLDBJ, PVLDB, IEEE TKDE, IEEE TDSC, ACM KDD, ACM CCS and IEEE ICDE, and won CIKM 2015 Best Paper Runner Up, the Best Papers of ICDE 2016 and the Best Papers of ICDM 2018. He has served as program committee member for leading conferences, including ACM KDD, IEEE ICDM, and ACM CIKM, and as reviewer for numerous flagship journals, including VLDBJ, PVLDB, IEEE TKDE, IEEE TDSC, IEEE TIFS, and IEEE TOPS.



Yong Ge received his Ph.D. in Information Technology from Rutgers, The State University of New Jersey in 2013, the M.S. degree in Signal and Information Processing from the University of Science and Technology of China (USTC) in 2008, and the B.E. degree in Information Engineering from Xi'an Jiao Tong University in 2005. He is currently an Assistant Professor at University of Arizona. His research interests include data mining and business analytics. He received the ICDM-2011 Best Research Paper Award, Excellence in Academic Research (one per school) at Rutgers Business School in 2013, and the Dissertation Fellowship at Rutgers University in 2012. He has published prolifically in refereed journals and conference proceedings, such as IEEE TKDE, ACM TOIS, ACM TKDD, ACM TIST, ACM SIGKDD, SIAM SDM, IEEE ICDM, and ACM RecSys. He has served as Program Committee members at premier conferences such as ACM SIGKDD conferences and IEEE ICDM conferences. Also he has served as a reviewer for numerous journals, including IEEE TKDE, ACM TKDD, ACM TIST, KAIS, Information Science, and TSMC-B.



Nick Duffield is a Professor in the Department of Electrical and Computer Engineering at Texas A&M University, a Professor by Courtesy, in the Department of Computer Science and Engineering, and Director of the Texas A&M Institute of Data Science. From 1995 until 2013, he worked at AT&T Labs-Research, Florham Park, NJ, where he was a Distinguished Member of Technical Staff and an AT&T Fellow. He previously held post-doctoral and faculty positions in Dublin, Ireland, and Heidelberg, Germany. He received a BA in Natural Sciences in 1982 and an MMath in 1983 from the University of Cambridge, UK, and a PhD in Mathematical Physics from the University of London, U.K., in 1987. His research concerns data science and computer networking, with current projects involving algorithms for data streaming and machine learning, computer network measurement and resilience, and applications of data science to transportation, agriculture and hydrology. He is an IEEE Fellow, an IET Fellow, and has twice received the ACM SIGMETRICS Test-of-Time award. He is an elected member of the governing board of ACM SIGMETRICS. He was an Associate Editor for the IEEE/ACM Transactions on Networking from 2007-2011, and has been an Editor-at-Large since 2014. He is Chief Editor for Big Data Network at Frontiers.



Na Wang received her M.S. and Ph.D. degrees in Information Sciences and Technology from the Pennsylvania State University. She is currently a Staff Software Engineer at Samsung Research America. Her research interests include machine learning and human-computer interaction. She has published papers in top conference proceedings, including IEEE ICDM, ACM SIGCHI, ACM IUI, ACM MobileHCI, and ACM CSCW, and won Best Paper Award of CHI 2015.

Correspondence and offprint requests to: Rui Chen, Samsung Research America, Mountain View, CA, USA. Email: rui.chen1@samsung.com