

Twitterにおけるセレンディピティなユーザ推薦手法の評価

徐 哲林¹ 周 娟² 高田 秀志²

概要: Twitter のお薦めユーザシステムでは、ユーザ自身がフォローしているユーザと同じジャンルのユーザが多く推薦される傾向がある。しかし、このようなお薦めユーザは本人が既に知っている可能性が高く、推薦しなくても自身が見つけられるため、ユーザに高い満足度を与えることが難しい。そのため、本研究では、興味のあるものはユーザの Timeline から抽出した名詞に現れると想定した上で、「個人を惹き付ける興味」を推定することによりセレンディピティなユーザを推薦する手法を構築し、評価実験を行った。その結果、ユーザの Timeline から抽出した名詞に基づいて、興味を推定できるだけではなく、提案手法によりセレンディピティなユーザを発見できることとともに、Twitter のお薦めユーザシステムよりセレンディピティに関する性能も向上することが示された。

Evaluation on a Serendipitous User Recommendation Method for Twitter

XU ZHELIN¹ ZHOU JUAN² TAKADA HIDEYUKI²

Abstract: Twitter's user recommendation system is likely to focus on prediction accuracy, which tends to recommend users matching with people's profiles. However, such a recommendation system can hardly provide users with high satisfaction because they can likely already know such a recommended user or easily find him or her by themselves. In this paper, assuming that some interests appear in nouns extracted from the user's Timeline, we propose a serendipity-oriented user recommendation scheme to predict serendipitous users that are both unexpected and useful to the user by estimating "user's attracting interests". The evaluation experiment shows not only that user's interests can be estimated based on nouns extracted from user's tweets, but also that serendipitous users can be found by the proposed scheme and the performance measure on serendipity can be improved compared to the Twitter's user recommendation system.

1. はじめに

近年、インターネットの発展に伴い、Web 上のデータ量は膨大になっている。それによって、ユーザが自分にとって価値のあるものを探すのが難しくなるという問題が発生している。代表的なマイクロブログである Twitter を一例に挙げると、現在、世界で月間 3 億 2600 万人、日本国内では 4500 万人のアクティブユーザに利用されており、国内外に広く普及している。しかしながら、Twitter を利用しているユーザが多くなるとともに、ユーザがフォローしたい人を探すのも難しくなる。そのため、このような問題を解決するために、データマイニング技術を応用してユー

ザの手助けをする推薦システムが登場している [1]。

現在に至るまで、多くの推薦システムでは、一般的に予測の正確さ (*Prediction Accuracy*) という特性 (*Property*) を重視し、ユーザのプロファイル (たとえば、行動履歴や好きなアイテムなど) と類似しているアイテムしか推薦しないという傾向がある [2]。

Twitter のお薦めユーザシステム (「おすすめユーザ」という機能) では、自分がフォローしているユーザと同じジャンルのユーザが多く推薦される傾向がある。たとえば、ユーザが多くの歌手をフォローしていれば、有名な歌手が多く推薦される。これは、上で述べた予測の正確さを重視している可能性が高いからではないかと考える。しかしながら、このようなお薦めユーザは本人が既に知っている可能性が高いので、推薦しなくても自身が見つけられ、しか

¹ 立命館大学院情報理工学研究科

² 立命館大学情報理工学部

も、同じジャンルのユーザを多く推薦すれば、推薦ジャンルの範囲も制限される。このように、単に予測の正確さを重視すると、フォローしたい、すなわち関心があるユーザが多くは推薦されず、ユーザに高い満足度を与えることができない。このようなことから、推薦システムに対して予測の正確さ以外の特性、特にセレンディピティ (*Serendipity*) が注目されている [3,4]。

セレンディピティとはどれだけ魅力的かだけではなく、意外性という要素が加わった特性であると考えられる [2,3]。図1はTwitterにおけるセレンディピティなユーザが推薦されるシナリオの例である。ユーザAと何らかの関係があるユーザBは火鍋が好きで、火鍋に関する情報を発信した。また、ユーザAはパンケーキが好きで、普段はパンケーキに関するtweetを見ているが、パンケーキと違うものである火鍋に関する情報を受け取るのは予想していなかった。そのため、意外に感じるとともに関心を持ち、ユーザBはセレンディピティなお薦めユーザになった。このように、Twitterにおいて、ユーザAにとってセレンディピティのある情報を発信するユーザはセレンディピティなユーザであると考えられる。

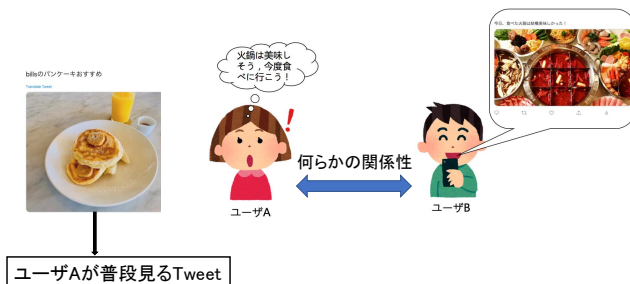


図1 セレンディピティなお薦めユーザの例

本研究では、Twitterにおいて、以前に我々が提案したセレンディピティなお薦めユーザを推定する手法 [5] を改善した上で、実際のTwitterを利用している人を被験者にユーザを推薦する実験を行い、提案手法により推定されるユーザ自身の興味、かつ、推薦されるユーザに対して、セレンディピティに関する性能について評価する。実施した実験の結果、提案手法によりユーザ自身の興味を推定できるだけでなく、セレンディピティなお薦めユーザを発見できることとともに、Twitterのお薦めユーザシステムよりセレンディピティに関する性能も向上することが示された。

本論文の構成は以下に示す通りである。第2章ではTwitterにおける提案するセレンディピティのあるお薦めユーザを推定する手法について述べる。第3章では提案手法を評価する実験について述べる、また、実験結果に基づいて考察を行う。最後に、第4章はまとめと今後の展望について述べる。

2. 提案手法

2.1 手法の全体像

図2は提案手法の全体像である。ユーザの興味はtweetとretweetに現れる [6]。そのため、まず、ユーザベース協調フィルタリングに基づいて、ユーザ p と、ユーザ p がフォローしているユーザのtweetとretweetに現れる興味から、ユーザ p と類似しているユーザを抽出する。次に、ユーザ p と、ユーザ p と類似しているユーザの興味によって、「 p の惹かれた興味」を推定する。また、ユーザ p と類似しているユーザの興味によって、「 p の関心がありそうな興味」を推定する。その上で、これらの2つの興味を合わせ、それぞれの興味の重要度に基づいて「 p を惹き付ける興味」を推定する。

ユーザ p と、ユーザ p と類似している人がもつ興味の中で、 p 自身の興味を除いた上で推定される惹き付ける興味は、「どれだけ魅力的か」と「意外性」を同時に満足し、ユーザ p のセレンディピティな興味になると想定する。そのため、ユーザ p と類似している人がフォローしているユーザの中で、 p を惹き付ける興味を常に発信しているユーザ、すなわちこのような興味をもつユーザをセレンディピティのあるユーザとしてユーザ p に推薦する。

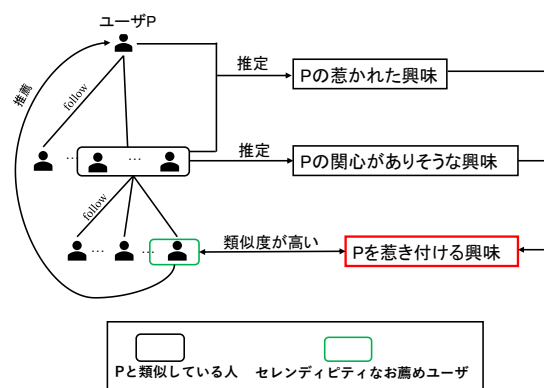


図2 提案手法の全体像

2.2 惹かれた興味の推定

retweetはユーザが興味のあるtweetに対して行われる行為なので [7]、retweetにもユーザの興味が現れているが、他人のtweetを見て、自身の興味と違う興味が惹かれ、そのツイートをretweetしたものの、時間が経てば忘れそうな興味が含まれていると考える。このような興味は、retweetに現れる興味の中でもよく出てこない、すなわち下位に属する興味であると考えられ、ユーザにとって意外性を感じると考えられる。また、ユーザベース協調フィルタリングに基づく、そのような興味の中でも、ユーザと類似している人の興味にも含まれるものは、推薦されると再び興味を持つ可能性が高い。

これに基づき、本研究では、retweet に現れる下位興味のうち、ユーザ自身の興味には含まれず、ユーザと類似している人が持つ興味に含まれるものを「惹かれた興味」として扱う。惹かれた興味は関心がありそう、かつ、意外性を感じるものであり、セレンディピティの二つの要素を同時に満足すると考えられる。

惹かれた興味を推定するために、まず、ユーザと類似している人の興味を推定する。次に、ユーザ自身の興味と下位興味を抽出した上で、惹かれた興味を推定する。

2.2.1 ユーザと類似している人の興味を推定する手法

Twitter で最も特徴的な機能である tweet と retweet は多くのユーザが常に利用しているので、これらを用いることによってユーザと類似している人を推定しやすいと考えられる。また、ユーザと類似している人と推定された人らが投稿した tweet と retweet から各人の興味を推定する。

興味のあるものは基本的に名詞に現れると考えられるため、ユーザとユーザがフォローしている人の Timeline から tweet と retweet を抽出した上で、名詞を抽出する。Timeline の抽出には Tweepy^{*1}を、名詞の抽出には MeCab^{*2}を用いる。ここで、ユーザの Timeline にはリプライもあるが、眩きやユーザ間の雑談が多いと考えられるため、本研究ではリプライを用いない。次に、抽出した名詞をユーザごとに一つの文書としてコーパスに入れた上で、各文書にある名詞の TF-IDF 値を計算し、コサイン類似度を用いてユーザ同士の類似度を計算する。

本研究ではユーザ p との類似度の最上位から 15 人をユーザと類似している人とする。また、これらのユーザの tweet と retweet から抽出した名詞を TF-IDF 値に基づいて降順に並べ替え、最上位から 20 位までの名詞を各人の興味とする。

2.2.2 ユーザ自身の興味と下位興味

ユーザ p に対して前節で述べた方法により抽出した興味に加えて、自己紹介に書いた興味は Timeline から推定する興味より重要度が高いので、自己紹介にある興味 (i 個) と Timeline から推定した興味 (20- i 個) を合わせて 20 個の興味をユーザ p 自身の興味とする。また、下位興味は retweet にあまり現れない興味であるとしているので、ユーザ p の retweet から抽出される名詞の TF-IDF 値によって最下位から 20 個までの興味を下位興味とする。

2.2.3 惹かれた興味を推定する手法

先に述べたように、ユーザ p の下位興味は忘れそうな興味であると考えられるため、自身の興味を除いて残った興味 (UI_p) は意外性を感じると考えられる。ユーザ p は自分と類似している人から興味を惹かれるので、次に、式 1 に従って p の惹かれた興味 (AI_p) を推定する。ここで、 FI は p と類似している 15 人がもつ興味の集合である。

$$AI_p = UI_p \cap FI \quad (1)$$

2.3 関心がありそうな興味の推定

上述のようにして、ユーザ p の惹かれた興味を推定するが、ユーザ p と類似している人は同じ興味をもつ可能性があるため、 FI が多くない可能性がある。それによって、 UI_p との共通部分が少なくなり、十分な数の惹かれた興味を推定できない可能性がある。したがって、ユーザ p の「関心がありそうな興味」を導入し、惹かれた興味を合わせてユーザ p を惹き付ける興味を推定する。

関心がありそうな興味とは他の人と関心が共通している興味である。ユーザ p と類似している人がもつ興味 FI は他の人よりユーザが好きになる可能性が高いが、 p にとって FI のすべては関心があるわけではないと考えられる。そのため、まず、ユーザ p 自身の興味と惹かれた興味を FI から除いて、残った興味の重要度を計算し、重要度の高い興味は関心がありそうな興味としてユーザが好きになる可能性が高いと考える。また、ユーザ自身の興味を含まないので、意外性を感じると考えられる。そのため、セレンディピティにある二つの要素を同時に満足する。

FI からユーザ p 自身の興味と惹かれた興味を除いて残った興味の重要度を式 2 [8] に従って算出し、重要度の最上位から 20 個までの興味を p の関心がありそうな興味にする。この式を用いることにより、ユーザ p と類似しているユーザの興味 i に対する評価値と、各ユーザが持つすべての興味の平均評価値を考慮した上で、ユーザ p がまだ持っていない興味 i に対して、どのぐらい関心があるか、すなわち興味の重要度を予測する。

興味の評価値には、興味ごとの TF-IDF 値 (重要度) を用いる。 $\hat{r}_{pi}$ は興味 i の重要度、すなわちユーザ p が興味 i に対して予測される評価値である。 $\bar{r}_p$ はユーザ p 自身の興味の平均評価値である。 U_p はユーザ p と類似している 15 人である。 $\text{sim}(p, v)$ はユーザ p と v のコサイン類似度である。 r_{vi} はユーザ v の興味 i に対する評価値である。ユーザ v が興味 i を持たない場合には、 r_{vi} は 0 にする。 $\bar{r}_v$ はユーザ v がもつ 20 個興味の平均評価値である。

$$\hat{r}_{pi} = \bar{r}_p + \frac{\sum_{v \in U_p} \text{sim}(p, v) \cdot (r_{vi} - \bar{r}_v)}{\sum_{v \in U_p} |\text{sim}(p, v)|} \quad (2)$$

2.4 惹き付ける興味の推定

惹かれた興味はユーザ p が従来からもつ興味なので、推定された p の関心がありそうな興味より好きになる可能性が高く、重要度も関心がありそうな興味より高くすべきであると考えられる。そこで、まず、推定したすべての惹かれた興味 (j 個) とユーザ p の関心がありそうな興味 (重要

*1 <https://www.tweepy.org/>

*2 <https://taku910.github.io/mecab/>

度上位 (20-j) 個) を合わせて p を惹き付ける興味 (20 個) にする。

次に、惹き付ける興味の中にある、惹かれた興味の元の重要度 (興味の TF-IDF 値) を一番関心がありそうな興味の重要度に足して、惹かれた興味の新しい重要度とする。つまり、惹かれた興味の重要度を関心がありそうな興味の重要度より高くするために、惹かれた興味の元の重要度を関心がありそうな興味の中で一番重要度の高い興味の重要度に加えて、惹かれた興味の新しい重要度にする。

また、惹き付ける興味の重要度は 1 を超える可能性がある。最後に、式 3 に従って最小値 0, 最大値 1 に正規化する。ここで、 Y は正規化した惹き付ける興味の重要度である。 X は各惹き付ける興味の元の重要度であり、 $xmin$ は一番重要度の低い惹き付ける興味の重要度、 $xmax$ は一番重要度の高い惹き付ける興味の重要度である。

$$Y = \frac{X - xmin}{xmax - xmin} \quad (3)$$

2.5 セレンディピティのあるユーザの推薦

まず、ユーザ p と類似している 15 人がフォローしているすべてのユーザに対して、各ユーザの tweet と retweet を含む名詞の TF-IDF 値を計算する。次に、各ユーザの名詞に p を惹き付ける興味があれば、それを抽出した上で、計算した TF-IDF 値を重要度として、ユーザごとのベクトルを作成する。最後に、コサイン類似度を用いて p を惹き付ける興味 (20 個) との類似度を計算し、類似度の高いユーザ 20 人を惹き付ける興味をもつユーザにする。このようなユーザはセレンディピティのあるユーザとしてユーザ p に推薦する。

3. 評価実験

3.1 実験目的と仮説

本実験の目的と検証する仮説は以下の通りである。本実験では、提案手法により推薦されるユーザを Twitter のお薦めユーザシステムが提供しているユーザと比較する。

● 実験目的

以下の 3 つの目的を設定する。

- (1) ユーザ自身の興味とは異なる興味はユーザが意外性を感じると想定するため、自身の興味を推定する必要がある。そこで、TF-IDF を用いて被験者の Timeline^{*3}から抽出した名詞によって被験者自身の興味を推定できるかを検証する。
- (2) Twitter において、提案手法によりセレンディピティなお薦めユーザを推定できるかを検証する。
- (3) Twitter が提供しているお薦めユーザに比べて、提案手法により推薦されるユーザに対して、セレ

ンディピティに関する性能が向上するかどうかを検証する。

● 仮説

実験目的を達成するために、以下の仮説を立てる。

- (1) TF-IDF を用いて被験者の Timeline から抽出した名詞によって被験者自身の興味を推定できる。
- (2) Twitter において、提案手法によりセレンディピティなお薦めユーザを推定できる。
- (3) 提案手法により推薦されるユーザ群に対して、セレンディピティのあるユーザの割合は Twitter が提供しているお薦めユーザより多くなる。

3.2 検証方法の設定

各仮説に対する検証方法を以下に示す。

● 仮説 1

各被験者に対して推定される自身の興味を見てもらい、「関心があるか」をアンケートで評価する。関心がある程度は 1 から 5 まで、評価値が 3 以上の場合に、関心があるとみなす。

● 仮説 2

本研究では Ge らが提案するセレンディピティの評価方法 (式 4) [9] を用いて、推薦されるユーザの中で、セレンディピティのあるユーザの割合を求める。

$$serendipity = \frac{|UNEXP \cap USEFUL|}{|N|} \quad (4)$$

式 4 に示すように、 $UNEXP$ と $USEFUL$ の両方を満足すると、セレンディピティなユーザとみなす。ここで、 $USEFUL$ は推薦されたユーザに対して、アンケートで被験者が関心を持ったと回答したユーザの集合であり、 N は推薦されたユーザの人数である。また、 $UNEXP$ は意外性を感じるユーザの集合であり、一般的には、推薦されたユーザのうち、予測の正確さを重視する手法により推薦されないユーザを意外性を感じるユーザとみなす。これによって、お薦めユーザに対して意外に感じるかどうかを評価できるが、ユーザからフィードバックを得なければ、セレンディピティの評価結果は信用できないため [10]、本研究では、推薦されるユーザに対して、式 4 に基づいて、アンケートで直接被験者に「意外性を感じるか」($UNEXP$) を回答してもらうこととする [11]。

また、式 5 と式 6 に従って $UNEXP$ と $USEFUL$ の評価値を計算する。2 つの式に示すように、アンケートでお薦めユーザに対して被験者が 3 点以上をつけると、意外性を感じる、または関心があると認める。両方の評価値が 1 であれば、セレンディピティなお薦めユーザになる。

^{*3} 本研究が利用した Timeline データはリプライを含まない

$$UNEXP = \begin{cases} 0 & (rating < 3) \\ 1 & (rating \geq 3) \end{cases} \quad (5)$$

$$USEFUL = \begin{cases} 0 & (rating < 3) \\ 1 & (rating \geq 3) \end{cases} \quad (6)$$

● 仮説 3

Twitter のお薦めユーザシステムにより推薦されたユーザに対して、仮説 2 と同じように被験者に対してアンケートを行って、セレンディピティのあるユーザの割合を計算した上で、仮説 2 の結果と比較し、提案手法の方が値が高くなるかを確かめる。

3.3 実験の設定

本節では、仮説を検証するために行う実験の設定について述べる。

3.3.1 被験者の条件

被験者が持つ情報、たとえば、投稿したツイートが少なく、提案手法によりセレンディピティなユーザを十分に推定できないため、以下の二つの条件を設定し、本研究が対象とするユーザの要件とする。

- (1) 被験者がフォローしている人数 > 80
- (2) 被験者が投稿したツイート数 > 500

3.3.2 抽出した tweet と retweet の件数と前処理

今回、ユーザの興味を推定するために、各ユーザが発信した tweet と retweet のみを抽出する必要がある。しかしながら、Tweepy はリプライも含んだ Timeline 全体を抽出するようになっている。そのため、各被験者の Timeline から 3500 件を抽出した上で、リプライを除き、tweet と retweet を合わせて 200 件を抽出する。また、被験者の惹かれた興味を推定するために、retweet のみを 200 件抽出する。最後に、被験者がフォローしているユーザ (1-hop ユーザ) と被験者と類似している人がフォローしているユーザ (2-hop ユーザ) に対して、Timeline から 600 件を抽出した上で、リプライを除き、tweet と retweet を合わせて 200 件を抽出する。

抽出された tweet と retweet には様々なノイズとなる情報がある。たとえば、URL やユーザ ID などのように興味として分析できない情報はノイズとなる情報である。このようなノイズがあると、ユーザ間の類似度と各被験者の興味をうまく推定できないと考えられる。そのため、抽出されたデータから以下の前処理によってノイズを除去する。

- URL を除去する。
- EMOJI を除去する。
- @user ID を除去する。
- 英単語は全部小文字に変換する。
- 半角に統一する。
- 数字を全部 0 に置き換える。

- SlothLib を利用し、情報のない Stop-Word を除去する。

3.3.3 興味の抽出方法

前節で述べた前処理を行ったとしても、MeCab を用いて抽出した名詞にはまだノイズがある。たとえば、「笑」や「下」など意味がない名詞が存在する。そのため、それぞれのユーザに対して抽出された名詞の中で、興味を表すのに相応しくない名詞を人手で削除した上で、残った名詞の TF-IDF 値に基づいて、最上位から各ユーザの興味を抽出する。また、実験者の retweet にある名詞に対しても同様に、興味を表すのに相応しくない名詞を人手で削除した上で、TF-IDF 値に基づき、最下位から被験者の下位興味を抽出する。

3.3.4 2-hop ユーザの人数の設定

2-hop ユーザはユーザと類似している人 (15 人) がフォローしているユーザである。Twitter API には制限があるので、2-hop ユーザの Timeline を全部抽出するのは非常に時間がかかる。そこで、本研究では、2-hop ユーザの中から、重複するユーザを除去した上で、ランダムで抽出された 2000 人を対象にする。

3.3.5 実験の内容

各被験者に対して、Twitter のお薦めユーザシステムから推薦されるユーザの最上位から 20 位までを集め、推薦リスト 1 とする。ここで、推薦リスト 1 のお薦めユーザは被験者が各自で Twitter アプリを用いて確認する。また、提案手法により推定される 20 人は推薦リスト 2 とする。推薦リスト 1 と推薦リスト 2 のそれぞれに含まれるユーザに対して、自己紹介と Timeline を 5 分間見てもらい、アンケートに記入する。

3.4 全被験者の基本的な情報

本実験の被験者は 8 人である。

まず、各被験者の Timeline からリプライを除き、上限を 200 件として抽出した tweet, retweet の数、および、下位興味の推定のために 200 件を上限として抽出した retweet の件数を表 1 に示す。retweet 件数については、3500 件を Timeline から抽出しても、上限の 200 件より少ない被験者も存在している。

表 1 抽出した Timeline と retweet の件数

被験者	Timeline	retweet
A	200	200
B	200	200
C	200	69
D	200	173
E	200	165
F	200	200
G	200	134
H	200	112

次に、各被験者がフォローしているユーザ (1-hop ユー

ザ)の中で、ロックしていないユーザの Timeline から上限を 200 件として tweet と retweet を抽出した。各被験者の 1-hop ユーザの人数、そのうちロックしているユーザの人数、および、ロックしていないユーザの人数を表 2 に示す。

最後に、各被験者と類似している人がフォローしているユーザ (2-hop ユーザ) の中から、ランダムに 2000 人を選択し、ロックしていないユーザの Timeline から 200 件を上限として tweet と retweet を抽出した。各被験者の 2-hop ユーザの人数、そのうちロックしているユーザの人数、および、ロックしていないユーザの人数を表 3 に示す。

表 2 1-hop ユーザ

被験者	1-hop	Lock	Unlock
A	105	64	41
B	296	59	237
C	215	132	83
D	314	156	158
E	374	180	194
F	323	113	210
G	383	172	211
H	384	135	249
平均値 (割合)			
	1-hop	Lock	Unlock
全被験者	299.25	126.38 (42.23%)	172.87 (57.77%)

表 3 2-hop ユーザ

被験者	2-hop	Lock	Unlock
A	2000	562	1438
B	2000	427	1573
C	2000	473	1527
D	2000	448	1552
E	2000	364	1636
F	2000	391	1609
G	2000	932	1068
H	2000	523	1477
平均値 (割合)			
	2-hop	Lock	Unlock
全被験者	2000	515 (25.75%)	1485 (74.25%)

3.5 実験結果

仮説 1 において、推定された被験者自身の興味に対して関心があるかどうかを評価した結果を表 4 に示す。また、被験者 A, B, C の自己紹介に書かれている興味の個数は 1, 5, 7 である。その他の被験者の自己紹介を書いていたので、すべての興味は tweet と retweet からのみ推定されたものである。

仮説 2 において、提案手法と比較手法^{*4}により推薦されるユーザリストの中で、各被験者にとって意外性を感じるユーザの人数を図 3 に示す。図 4 には、各被験者にとって関心があるユーザの人数を示す。

また、各被験者へ推薦される 20 人の中で、式 4 により求められるセレンディピティなユーザの占める割合を図 5 に

^{*4} Twitter のお薦めユーザシステムが利用している推薦手法

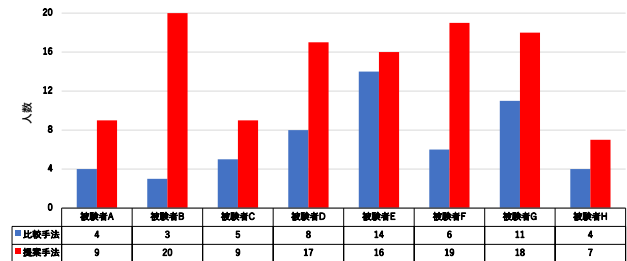


図 3 意外性を感じるお薦めユーザ

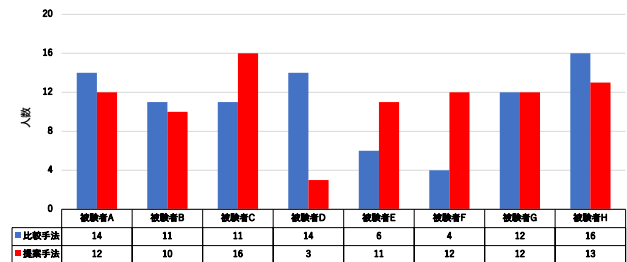


図 4 関心があるお薦めユーザ

示す。さらに、仮説 3 において、セレンディピティなユーザの占める割合に関して全被験者を対象に集計した結果を表 5 に示す。

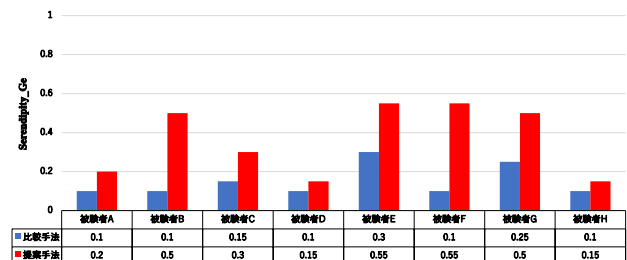


図 5 各推薦リストに対するセレンディピティなユーザの割合

表 5 セレンディピティなユーザの割合に関する集計結果

	平均値	標準偏差	中央値	有意確率 p (両側)
比較手法	0.15	0.08	0.1	0.006
提案手法	0.3625	0.18	0.4	

3.6 考察

実験結果によって、すべての仮説が成立していることが確認された。以下、実験結果について考察する。

表 1 に示されているように、Timeline から 3500 件抽出しても、8 人の被験者のうち 5 人が発信した retweet は 200 件未満であることがわかった。これにより、一部のユーザは日常的にリツイートしていない可能性があると考えられるため、retweet から推定される惹かれた興味のみに着目すると、セレンディピティなお薦めユーザを十分に推定できない可能性がある。したがって、関心がありそうな興味を導入したのは効果があったと考えられる。

表 2 に示されているように、1-hop ユーザの中でロック

表 4 推定された自身の興味に関する評価結果

被験者	推定される興味 (個)	関心があるか		関心がある興味の割合 (%)
		関心あり (個)	関心なし (個)	
A	19	15	4	78.9%
B	15	11	4	73.3%
C	13	12	1	92.3%
D	20	16	4	80.0%
E	20	18	2	90.0%
F	20	17	3	85.0%
G	20	14	6	70.0%
H	20	13	7	65.0%
平均値				
	推定される興味 (個)	関心あり (個)	関心なし (個)	関心がある興味の割合 (%)
全被験者	18.375	14.5	3.875	78.91%

しているユーザの占める割合の平均値は 42.23%である。このように、Timeline を抽出できるユーザの人数はそれほど多くないため、本実験では、各被験者と類似しているユーザのコサイン類似度は僅か 0.3 未満であった。また、表 3 に示されているように、2-hop ユーザの中から、平均では 1485 人から提案手法によるセレンディピティなお薦めユーザが推薦されている。2-hop ユーザの人数が多いため、セレンディピティに関する性能に悪い影響がある。そこで、この二つの問題を改善すれば、セレンディピティに関する性能を向上されることが考えられる。

表 4 に示されているように、推定される被験者自身の興味の中で関心がある興味の割合の平均値は 78.91%である。そのため、ユーザの Timeline から抽出した名詞に基づいて、TF-IDF を用いて、自身の興味を推定できていると考えられる。

図 3 と図 4 に示されているように、提案手法によって推定されるユーザの中で被験者にとって関心があるユーザは比較システムより少なくはないとともに、意外性を感じるユーザが大幅に多くなっていることが見られるが、全部セレンディピティなユーザになるわけではない。そこで、単に意外性を向上させるのはセレンディピティなユーザの発見に対して、役に立たないと考えられる。また、図 5 に示されているように、各被験者に対して、提案手法によって推薦されるお薦めユーザの中でセレンディピティなユーザの占める割合は比較システムより高くなっていることが見られる。

表 5 に示されているように、提案手法により推定されるユーザの中で、セレンディピティなユーザの占める割合の平均値は 0.3625 である。この平均値の意味として、たとえば、被験者に 20 人を推薦する場合に、7.25 人がセレンディピティなお薦めユーザであると認められる。また、Piao らの研究 [12] では、Twitter における 3 つの推薦方法を提案し、式 4 を用いて、被験者 3 人にそれぞれ推定された 40 個の興味を推薦した。結果として、セレンディピティな興味の割合の平均値は 3 つの方法でそれぞれ 0.092, 0.163, 0.125 であることが示されている。これによって、本研究

の提案手法は Piao らの手法よりセレンディピティに関する性能が高くなると認められ、本研究の提案手法によって、セレンディピティなお薦めユーザをより多く推定できると考えられる。最後に、提案手法は比較手法よりセレンディピティなユーザの割合に関する評価の平均値と中央値の両方が高くなることが見られる。さらに、対応のある T 検定を用いて検定した結果、有意差が認められた ($p < 0.01$)。そこで、Twitter のお薦めユーザシステムよりセレンディピティに関する性能が高くなると考えられる。

4. おわりに

推薦システムにおいては、予測の正確さ以外の特性であるセレンディピティを重視する必要があるため、本研究では、Twitter において、協調フィルタリングに基づいて、提案したユーザを惹き付ける興味を推定し、ユーザと類似している人がフォローしているユーザの中で、このような興味を多く持つユーザをセレンディピティのあるユーザとして推薦する。被験者 8 人に対して実験を実施し、ユーザの Timeline から抽出した名詞に基づいて、TF-IDF を用いて、自身の興味を推定できているだけでなく、提案手法によりセレンディピティなお薦めユーザを発見できることが示された。さらに、Twitter のお薦めユーザシステムよりセレンディピティに関する性能が高くなることが示された。

今後は、ユーザとより類似している人を発見するために、各ユーザが発信する情報の中で、言葉として異なるが類似している話題であるかどうかを判別する手法を確立した上で、ユーザがフォローしている人 (1-hop) に制限せず、より広い範囲でユーザと類似している人を推定することを試みる。

参考文献

- [1] Ricci, F., Rokach, L. and Shapira, B.: Introduction to recommender systems handbook, in *Recommender systems handbook*, pp. 1-35, Springer (2011).
- [2] Silva, A. M., Silva Costa, da F. H., Diaz, A. K. R. and Peres, S. M.: Exploring Coclustering for Serendipity Improvement in Content-Based Recommendation, in *International Conference on Intelligent Data Engineering*

- and Automated Learning*, pp. 317–327 (2018).
- [3] Zheng, Q., Chan, C.-K. and Ip, H. H.: An unexpectedness-augmented utility model for making serendipitous recommendation, in *Industrial conference on data mining*, pp. 216–230 (2015).
 - [4] Kotkov, D., Veijalainen, J. and Wang, S.: A serendipity-oriented greedy algorithm for recommendations, in *WEBIST 2017: Proceedings of the 13rd International conference on web information systems and technologies. Volume 1*, ISBN: 978-989-758-246-2 (2017).
 - [5] 徐哲林, 周娟, 高田秀志 : マイクロブログにおける「惹き付ける興味」の推測によるセレンディピティなユーザの推薦, ワークショップ 2019 (GN Workshop 2019) 論文集, 第 2019 巻, pp. 85–92 (2019).
 - [6] Kim, D., Jo, Y., Moon, I.-C. and Oh, A.: Analysis of twitter lists as a potential source for discovering latent characteristics of users, in *ACM CHI workshop on microblogging*, Vol. 6 (2010).
 - [7] Boyd, D., Golder, S. and Lotan, G.: Tweet, tweet, retweet: Conversational aspects of retweeting on twitter, in *2010 43rd Hawaii International Conference on System Sciences*, pp. 1–10 (2010).
 - [8] Ning, X., Desrosiers, C. and Karypis, G.: A comprehensive survey of neighborhood-based recommendation methods, in *Recommender systems handbook*, pp. 37–76, Springer (2015).
 - [9] Ge, M., Delgado-Battenfeld, C. and Jannach, D.: Beyond accuracy: evaluating recommender systems by coverage and serendipity, in *Proceedings of the fourth ACM conference on Recommender systems*, pp. 257–260 (2010).
 - [10] Kaminskas, M. and Bridge, D.: Diversity, serendipity, novelty, and coverage: a survey and empirical analysis of beyond-accuracy objectives in recommender systems, *ACM Transactions on Interactive Intelligent Systems (TiiS)*, Vol. 7, No. 1, p. 2 (2017).
 - [11] Shani, G. and Gunawardana, A.: Evaluating recommendation systems, in *Recommender systems handbook*, pp. 257–297, Springer (2011).
 - [12] Piao, S. and Whittle, J.: A feasibility study on extracting twitter users’ interests using nlp tools for serendipitous connections, in *2011 IEEE Third International Conference on Privacy, Security, Risk and Trust and 2011 IEEE Third International Conference on Social Computing*, pp. 910–915 (2011).