

局所差分プライバシーを用いた行列分解による ネット広告システムの提案

峯田 初音[†] 韓 耀緯^{††} 曹 洋^{††} 吉川 正俊^{††}

[†] 京都大学工学部情報学科 〒606-8501 京都市左京区吉田本町

^{††} 京都大学大学院情報学研究科 〒606-8501 京都市左京区吉田本町

E-mail: [†]{mineta,yaowei}@db.soc.i.kyoto-u.ac.jp, ^{††}{yang,yoshikawa}@i.kyoto-u.ac.jp

あらまし 本研究ではオンライン識別子用いずにパーソナライズされた広告を配信するネット広告システムを提案する。広告の選択においては、オンラインショッピングサイトなどで利用される商品推薦の概念を利用する。広告を商品、広告をクリックしたかどうかを商品に対する評価値とみなし、ユーザに最適な広告を予測する。ユーザ自身と信頼できない広告サーバ以外に信頼できる第三者を想定しないネット広告システムを提案している研究は未だかつて存在せず、この点で我々は貢献していると言える。そして、予測の正誤、真の評価値と予測の誤差について評価を行い、提案手法がプライバシーを保護しながら、高い有用性を担保していることを証明した。

キーワード 広告, パーソナライズ, プライバシ保護, 行列因子分解, 局所的差分プライバシー

1 はじめに

インターネット経済において、個人最適化されたインターネット広告は非常に重要な要素である。最近ではほとんどの広告システムに導入されており、個人最適化された広告が一般的になっている。ユーザと関連度の高い広告を選出する際には、サイトやアプリに登録されている性別や年齢などの基本的な情報に加え、トラッキング技術を用いて収集した閲覧履歴や購買履歴など大量のユーザデータが利用されている。

しかし、近年、ビッグデータの浸透に伴い、プライバシー保護の関心が高まっている。2018年にはEUでGeneral Data Protection Regulation(GDPR)[1]という法律が施行され、IPアドレスとCookieを含むオンライン識別子の利用が制限された。また、Googleは、2年以内に広告目的のCookieの利用を制限すると発表した[2]。今後ますますオンライン識別子の利用が困難になることが予想される。

そこで、本研究ではオンライン識別子を用いずに個人最適化された広告を配信するインターネット広告システムを提案する。プライバシー保護の手段として、匿名化や通常の差分プライバシーよりも強力な局所差分プライバシー[3]を用いた。最後に、システム全体が ϵ -差分プライバシーを保証していることを証明した。

本研究の貢献は、局所差分プライバシーを用いたインターネット広告システムを提案したことにある。筆者の知る限りこのようなシステムは未だかつて存在しない。

2 背景と動機

2.1 現在のインターネット広告システム

現在のインターネット広告システムではRTB(Real-Time Bidding)と呼ばれる仕組みが使われている。概要を図1に示

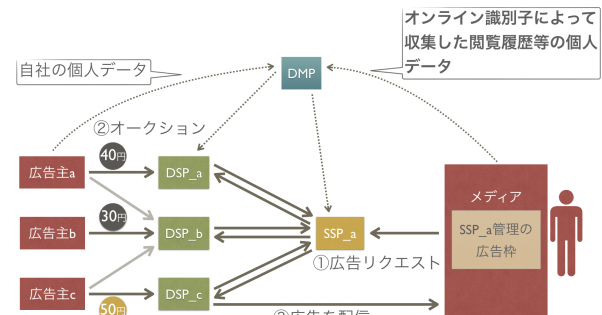


図1 現在のインターネット広告システム (RTB)

す。このシステムではDMP(Data Management Platform)と呼ばれる個人データを収集・管理する機関が重要な役割を果たしている。ユーザがサイトにアクセスすると、DMPは収集した個人データをもとにどのようなユーザがアクセスしてきたかを広告サーバに伝える。広告サーバはリアルタイムでオークションを行い、最も高価格で入札した広告が表示される。

DMPはCookieやIPアドレス等のオンライン識別子を利用してインターネット上の個人データをトラッキングしている。このようにして収集された大量の個人データを利用することで、RTBでは個人最適化された広告の配信が可能となっている。

2.2 プライバシへの関心の高まり

大量のパーソナルデータの創出に伴い、プライバシーへの関心が高まってきている。インターネット上の個人データに関する法律が世界各国で整備され始めた。2018年にはEUでGDPR(General Data Protection Regulation)[1]という法律が施行された。GDPRとは、ビッグデータ時代に個人情報を取り囲む環境変化への対応を目的として策定された、新しいデータ保護の枠組みである。GDPRにおいて保護対象となる個人

データには、Cookie と IP アドレスを含むオンライン識別子など、国内法にはないものも含まれている。日本を含む EU 域外の企業にも広く影響があり、また違反時の制裁金が高額なことから、GDPR への対応を検討する日本企業は増えている。日本でも 2020 年の個人情報保護法の改正に向けて、Cookie の取り扱いに関して議論されている [4]。

このように、オンライン識別子を含むインターネット上の個人データの利用が問題視され始めている。将来、オンライン識別子が利用できなくなる可能性があり、オンライン識別子なしで個人最適化された広告を配信する仕組みが必要である。

3 問題設定

広告配信システムには以下の 2 つのフェーズがある。

広告配信 ユーザとの関連度が高い広告を選択して配信する。
統計収集 広告の表示回数や CTR(Click Through Rate) といった統計量を計算する。広告主への料金請求や、フィードバック提供のために必要である。

この 2 つのフェーズにおいて、ユーザのプライバシーが脅かされる可能性がある。広告配信では、2.1 節で述べたようにインターネット上で収集した個人データを利用している。このような個人データにはユーザの閲覧履歴や購買履歴も含まれており、それらからユーザの性別や趣味、政治的思想が推定されてしまう可能性がある。また、2.2 節で述べたようなオンライン識別子の利用制限により、そもそも個人データを収集・利用することができなくなる可能性がある。統計収集では、ユーザにどの広告を表示したか、またユーザは広告をクリックしたかどうかを計測し、統計量を計算する。これらの統計量からも広告配信時と同様にユーザのプロファイルが推定されてしまう可能性がある。

上記を踏まえ、本研究では以下の 3 点を満足する広告システムを提案することを目標とする。

- IP アドレスや Cookie 等のオンライン識別子を使用しない

- 個人最適化されている

ユーザと関連度の高い広告を配信することで、広告費用対効果を高める。

- プライバシーが保護されている

攻撃者にユーザのプロファイルを学習されないようにする。

4 関連研究

4.1 プライバシー保護インターネット広告システム

プライバシー保護インターネット広告システムとは、広告配信や統計収集の際に、プライバシー保護技術を用いて、ユーザのプロファイルの保護を実現するものである。このようなシステムはいくつか提案されている。

PrivAd [5] では暗号化とディーラを利用した匿名化を行っている。広告配信時、ユーザは広告サーバの公開鍵を用いて暗号化したプロファイルの一部をディーラに送信し、ディーラは

表 1 関連研究と比較した際の本研究の位置付け

	個人最適化 の手法	広告配信時の プライバシー保護	統計収集時の プライバシー保護
PrivAd [5]	キーワードによる ランキング	匿名化 (ディーラを信頼)	
ObliviAd [6]	キーワードと 限られたユーザ 情報に従う	匿名化 (コプロセッサを信頼)	
Adnostic [7]	Web コンテンツ に従う	匿名化 (プロキシを信頼)	
Hardt ら の研究 [8]	限られたユーザ 情報に従う	匿名化 (プロキシを信頼)	
本研究	行列因子分解	局所差分 プライバシー	必要なし

ユーザからのリクエストにランダムな ID を付加して広告サーバに送信する。ディーラはユーザと ID のマッピング表を保存している。広告サーバはユーザプロファイルを用いて広告を選択し、ID と共にディーラに送り返す。ディーラは ID とマッピング表からユーザを特定し、広告を送信する。

ObliviAd [6] では、PrivAd におけるディーラ部分を、広告サーバのコプロセッサに置き換えた構成となっている。コプロセッサには広告サーバ自身もアクセスできないとしている。

また、Adnostic [7] や Hardt らの研究 [8] では、最終的な広告の決定はユーザデバイス上で行という形式をとっている。広告サーバは、ユーザが送信を許可したいくつかの個人データや Web サイトのコンテンツを元に複数の広告を選択しユーザに送信する。ユーザは全ての個人データを用いて 1 つの広告を選択する。ただし、この場合、1 つの広告をユーザに送信する場合に比べて通信コストが大きくなってしまう。また統計収集の際には、プロキシを用いた匿名化によりプライバシー保護を実現している。

このようにプライバシーを保護しつつ個人最適化された広告を配信するシステムは提案されてきているものの、我々は上記の手法に共通の問題が一点あると考えている。それは信頼できる第三者を想定しているということである。信頼できる第三者とは、PrivAd のディーラ、ObliviAd のコプロセッサ、Adnostic や Hardt らの研究のプロキシのことを指す。もし、なんらかの方法で広告サーバがこれらにアクセスできた場合、広告サーバ自身が保持する情報と合わせることで、ユーザのプロファイルが推測されてしまう可能性がある。

我々の研究では、プライバシー保護の手法として局所差分プライバシーを用いることで、信頼できる第三者を必要としないプライバシー保護インターネット広告システムを提案する。また本研究の位置付けを表 1 に示す。

4.2 プライバシー保護を満たす商品推薦

広告システムに限らず、商品推薦はユーザのプライバシーを脅かす可能性が大きい [9]。このようなプライバシーリスクからユーザを守るため、推薦において差分プライバシーを用いた研究が多くされてきた [10] [11] [12]。これらは推薦システムが各々のユー

ザの商品に対する評価データを収集した後、プライバシー保護を行う、というものである。一般的にこういったシステムは信頼できるサーバ上に存在しており、攻撃を受けることはないとされている。しかし全ての推薦システムが信頼できるサーバに存在しているという保証はなく、実際はサーバ上のデータを用いてユーザーの特定が行われる可能性がある、という危険性もある[13]。

こういったリスクを避ける方法の一つとして、推薦システムがユーザーのデータを収集するよりも前、つまり、ユーザーがデータを送信する際にデータのプライバシー保護を行う、というものがある。それを実現させるのが局所差分プライバシーという技術であり、その概要については5.2.2節で述べる。この局所差分プライバシーを商品推薦に適用した研究も昨今増加してきている。“ユーザーが評価した商品”と“商品に対する評価値”にノイズを加え、強力にプライバシーを保護した手法を提案しているのがShin[14]らの研究である。

プライバシー保護を満たす推薦は増加してきているものの、広告システムにおいて、信頼できる第三者の仮定なしでプライバシー保護を満たすものはこれまでにない。そこで本研究ではShinらの研究を参考として手法の提案を行うこととする。この時、“ユーザーが評価した商品”を“ユーザーが評価した広告”とし、“商品に対する評価値”を“広告に対する評価値”とする。詳細については第6章にて説明を行う。

5 準備

5.1 行列因子分解

行列因子分解 (Matrix Factorization) とは、商品推薦において最もよく使われる方法の一つである。ここで商品推薦とは、あるユーザーの未知の商品に対する評価を予測する方法である。

このとき最も良く用いられる方法の一つが協調フィルタリングというものであり、これはあるユーザーのある商品に対する評価値を予測するにあたり、他のユーザーの種々の商品への評価値を用いる方法である。協調フィルタリングでは n 人のユーザーが m 個の商品 (例: 映画, 商品) の中の任意の商品を評価している状況を考える。そしてそれぞれのユーザーの各商品への評価を、 $S \subset \{1, \dots, n\} \times \{1, \dots, m\}$ で表す。ここで評価数を $S = |S|$ 、ユーザー i の商品 j への評価を r_{ij} と表す。一般的に評価を表す行列は疎であるため、 S は nm よりもずっと小さい値であるといえる。

協調フィルタリングは現在与えられている $\{r_{ij} : (i, j) \in S\}$ を用いて、まだ評価されていない商品への評価値を予測する方法である。

一般的に通常の協調フィルタリングはユーザー数や商品数が多くなればなるほど精度が落ちてしまうという問題がある。これは評価値を表す行列が疎になってしまうためである。

この問題を解決するために用いられるのが行列因子分解という方法であり、これは行列の次元削減を行って精度を上げるというものである。行列因子分解では、評価値を表す $n \times m$ の行列 R に対して、ユーザー要素を表す $d \times n$ の行列 U と商品要素

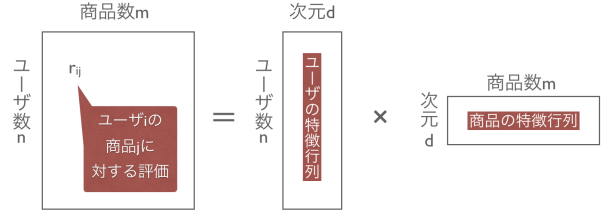


図2 行列因子分解

を表す $d \times m$ の行列 V を考える。ここで d は一般的に、 n や m よりもずっと小さい値となる。そして評価値を表す行列を R とすると、 U と V を用いた以下の近似式が成り立つ。

$$R \approx U^T V$$

図で表すと図2のような形となる。

このとき各ユーザーの各商品に対する評価値を表す行列は、行列 U^T と行列 V の積で表すことができる。具体的に、 u_i を U の i 列ベクトル、 v_j を V の j 列ベクトルとすると、ユーザー i の商品 j に対する評価値は、 u_i^T と v_j の内積から求められる。行列因子分解ではユーザー要素を $u_i \in \mathbb{R}^d$ ($i = 1, 2, \dots, m$)、商品要素を $v_j \in \mathbb{R}^d$ ($j = 1, 2, \dots, n$) と表し、この中のそれぞれの値を既知の評価値から学習する。学習においては、以下の(1)式で表される、正規化された平均二乗誤差を最小化する行列 U と行列 V を訓練データから導く。

$$\frac{1}{S} \sum_{(i,j) \in S} (r_{ij} - u_i^T v_j)^2 + \lambda_u \sum_{i=1}^n \|u_i\|^2 + \lambda_v \sum_{j=1}^m \|v_j\|^2 \quad (1)$$

ここで λ_u と λ_v は正規化を行うための値であり、正の定数とする。(1)式を最小化するための最適化問題は、最急降下法を用いて以下の式から行列 U と V の各々の値を更新することで解くことができる。

$$u_i^t = u_i^{t-1} - \gamma_t \cdot \{\nabla_{u_i} \phi(U^{t-1}, V^{t-1}) + 2\lambda_u u_i^{t-1}\} \quad (2)$$

$$v_j^t = v_j^{t-1} - \gamma_t \cdot \{\nabla_{v_j} \phi(U^{t-1}, V^{t-1}) + 2\lambda_v v_j^{t-1}\} \quad (3)$$

ここで γ_t は t 回目の反復における学習率、 $\nabla_{u_i} \phi(U, V)$ と $\nabla_{v_j} \phi(U, V)$ はそれぞれ u_i と v_j の勾配である。この勾配は(1)式で表される平均二乗誤差の導関数から求められ、以下の式で表すことができる。

$$\nabla_{u_i} \phi(U, V) = -\frac{2}{S} \sum_{j:(i,j) \in S} v_j (r_{ij} - u_i^T v_j) \quad (4)$$

$$\nabla_{v_j} \phi(U, V) = -\frac{2}{S} \sum_{i:(i,j) \in S} u_i (r_{ij} - u_i^T v_j) \quad (5)$$

これらの式を用いて(1)式を最小化する最適な行列 U と行列 V を計算し、求められた行列を用いて未知の商品に対する評価値を予測する。

5.2 差分プライバシー

本研究では 5.1 節で紹介した行列因子分解をプライバシー保護を満たす形に拡張するために、局所差分プライバシーを利用する。これは差分プライバシー [15] を拡張した技術である。差分プライバシーは、データベース中の個人データの含まれるレコードの内容を攻撃者から保護しつつも、データベース全体に対して正確に統計的解析ができるようにする仕組みである [15]。本節では、一般的な差分プライバシーの説明を 5.2.1 節で、それを拡張し、よりプライバシー保護を強固にした局所差分プライバシーについて 5.2.2 節で説明する。

5.2.1 差分プライバシー

差分プライバシー (Differential Privacy) は、あるデータベースから出力される統計結果を用いて、ユーザに関する何らかの情報を得ようとする攻撃者を想定する。そしてそれらの攻撃からユーザのプライバシー保護を行うことを目的とした技術である。この技術の一般的なアイデアは、“あるユーザが存在するデータベースと、そうでないデータベース、双方から統計的結果を得た際に、それがどちらのデータベースから得たものなのか分からないようにする” というものである。本節ではその数学的定義を述べる。

あるデータベース D を考える。 D はその要素である x_i の出現回数によって表現されることとし、これをヒストグラムと呼ぶ。ここで $i = 1, \dots, m$ である。このヒストグラムによって表現された二つのデータベース D, D' の距離を、以下に定義する l_1 ノルムを用いて定義する。

定義 1 (l_1 ノルム). m 次元ベクトル x の l_1 ノルム $|x|_1$ は次式で定義される。

$$|x|_1 = \sum_{i=1}^m |x_i|$$

この定義から、 D と D' の距離は $|D - D'|_1$ と書ける。これを用いて差分プライバシーの数理モデルを定義する。

定義 2 (差分プライバシー). 関数 $\mathcal{M}$ が以下の条件を満たすとき、 $\mathcal{M}$ は (ϵ, δ) -差分プライバシーを満たすという。任意のデータセット D, D' 及び $\forall Y \subseteq \text{Range}(\mathcal{M})$ に対して、 $|D - D'|_1 = 1$ であるとき

$$\Pr[\mathcal{M}(D) \in Y] \leq \exp(\epsilon) \Pr[\mathcal{M}(D') \in Y] + \delta$$

なお、 $\delta = 0$ の場合、 $\mathcal{M}$ は ϵ -差分プライバシーを満たすという。

ここで関数 $\mathcal{M}$ はあるデータセットを入力したときに、その統計的結果をある分布に従って変化させるものであり、用いられる分布にはラプラス分布や幾何分布などがある。また $\text{Range}(\mathcal{M})$ は、 $\mathcal{M}$ が生成する可能性のある出力の集合を指す。

ここで $|D - D'|_1 = 1$ であるとは、ある一人のユーザが評価した全ての商品とその評価値の中で一つのレコードが異なる、ということの意味する。

この一般的な差分プライバシーは信頼できるサーバを想定しているものであり、サーバはユーザから収集したノイズが加わっていない生データを収集し、タプルとしてデータベースに格納

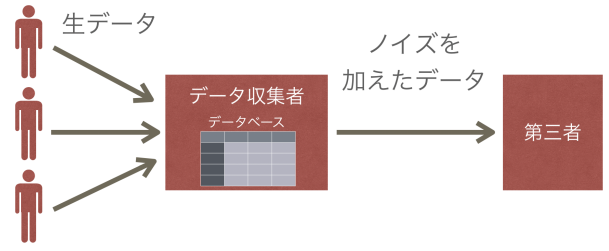


図 3 差分プライバシー

する。概要を図 3 に示す。ここでデータベースの 1 タプルは 1 ユーザに対応しているとする。そして統計結果を公開する際にデータベースの各タプルを集約し、関数を用いてノイズを加えたデータを生成し、そこから計算を行う。このデータベースが存在するサーバが本当に信頼できるものであれば、プライバシー保護は確実になされていると言える。しかしこのサーバが信頼できないものである可能性もある。そのような場合にユーザのプライバシーを保護するためには、各タプルに格納されるデータは収集される前にノイズを加える必要がある。その場合に利用されるのが局所差分プライバシーであり、それについて次節で紹介する。

5.2.2 局所差分プライバシー

局所差分プライバシー (Local Differential Privacy) のアイデアは、“各ユーザが各々のデータに対して差分プライバシーを満たすノイズを加える” というものである。ここで参考のために、通常の差分プライバシーは、“各ユーザのデータを格納したサーバが全てのデータに対してノイズを加える” ものである。局所差分プライバシーを実現するためのアイデアは、“ある 1 ユーザに関して、あるデータを持っているかどうかに関わらず、統計的結果が変化しないようにする” というものである。局所差分プライバシーの定義を以下に示す。

定義 3 (局所差分プライバシー). x, x' を任意のタプルとする。関数 $\mathcal{M}$ が以下の条件を満たすとき、 $\mathcal{M}$ は ϵ -局所差分プライバシーを満たすという。

$$\Pr[\mathcal{M}(x) \in Y] \leq \exp(\epsilon) \Pr[\mathcal{M}(x') \in Y]$$

ここで $\text{Range}(\mathcal{M})$ は、 $\mathcal{M}$ が生成する可能性のある任意の出力を指す。一般的な差分プライバシーと局所差分プライバシーは、入力として考えているものが異なる。前者はあるタプル中の 1 つのレコードのみが異なる 2 つのデータセットを、後者は任意の 2 つのタプルを考えている。そのため、収集されたデータから情報を得る際に、加わるノイズの量が異なってくる。一般的な差分プライバシーでは計算結果にノイズを加えるが、局所差分プライバシーでは、それぞれのユーザのデータ毎にノイズを加え、それらを用いて計算を行う。概要を図 4 に示す。よって、局所差分プライバシーのノイズは一般的な差分プライバシーよりも大きくなる傾向がある。しかしその分、より強固にプライバシー保護を行うことができる。

ここで、局所差分プライバシーを実現する方法について紹介す

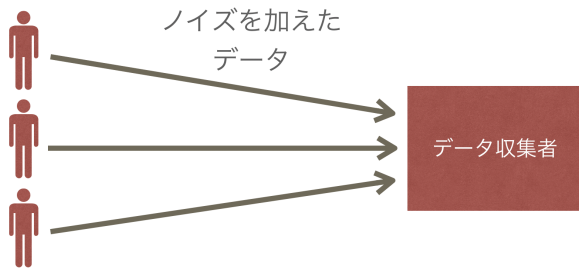


図 4 局所差分プライバシー

表 2 Randomized Response の決定確率の例

		答える確率	
		yes	no
真の値	yes	0.3	0.7
	no	0.4	0.6

表 3 本稿で扱う記号とその意味

記号	意味
n	ユーザの数
m	広告の数
d	次元数
k	繰り返しの回数
U	ユーザ行列. 行を $u_i \in \mathbb{R}^d$ とする
V	広告行列. 行を $v_j \in \mathbb{R}^d$ とする
r_{ij}	ユーザ i の広告 j に対する評価値
y_{ij}	ユーザ i の広告 j を評価したかどうかを表す値
∇_V^*	ノイズが付加された, 行列 V の勾配

る. 最もよく使われるのが Randomized Response [16] という方法であり, これはある値を入力した際, どの値を出力するかを指定した確率に基づいて決定するという方法である. 具体例として, ある質問に対して “yes” または “no” で答えた場合のノイズの加え方を表 2 に示す. この確率はそれぞれのユーザが任意に決定できるため, 各々のユーザは自らのプライバシー意識に応じてデータにノイズを加えることができる.

6 提案手法

6.1 手法の概要

本研究では, プライバシー保護を満たすインターネット広告システムを提案する. 具体的には, ユーザが Web サイトを訪れた際, 推薦によって配信する広告を選択し, また, 広告の表示回数とクリック数を計測する. 提案手法を説明する際に扱う記号の一覧を表 3 に示しておく.

6.2 想定する攻撃者

本研究ではインターネット広告システムが信頼できないサーバ上に存在すると考える. したがって攻撃者はこのシステムにアクセスできるものとし, 広告配信や統計収集の際にシステムが受け取るデータをもとに, あるユーザのプロファイルを知ろうとする者とする. この攻撃からユーザのプライバシーを守るため, システムに送信するデータにノイズを加えることとする.

6.3 広告配信

局所差分プライバシーを適用した行列因子分解を利用してユーザに最適な広告を選択して配信する. この時, プライバシー保護が必要な場面が 2 つ存在する.

- 学習時. つまり, 行列因子分解を行う時である. 分解後の 2 つの行列 (広告行列・ユーザ行列) を更新する際の勾配は, 式 (4), 式 (5) で示したように互いの行列を用いている. そのことから, ユーザ行列をサーバに推測されてしまう恐れがある. ユーザ行列が知られると, サーバが保持している広告行列からユーザの広告に対する評価値も推測できてしまうのでプライバシー保護が必要である.

- 配信時. ユーザデバイス上で, パブリックである広告行列と自身のユーザ行列 (ベクトル) から, 各広告に対する評価値を推測する. 評価値が最も高い広告が, ユーザに最も適している広告であるが, それをそのままサーバに通知するとユーザの趣向が推測される恐れがある. よってここでもプライバシー保護が必要である.

それぞれの手法について述べる.

6.3.1 学習

商品を広告, 商品に対する評価値を広告に対する評価値と置き換えて, 行列因子分解を適用する. まず通常の行列因子分解の手順を説明する.

(1) 広告システムがユーザに対してユーザ行列 U , 広告行列 V を送信する.

(2) ユーザが (1) と自身のデータを用いて行列 U , V を更新するための勾配を計算し, 広告システムに送信する.

(3) 推薦システムが (2) で受け取った勾配を用いて行列 U , V を更新する.

(4) (1) ~ (3) を複数回繰り返す.

提案手法ではユーザのデータに関してプライバシー保護を行うために, 2 点変更を加える. 1 つ目は, ユーザ行列 U の更新はユーザデバイス上で行うものとする. 2 つ目は, 上記の (2) のステップにおいてデータにノイズを加えるものとする. したがって提案手法は以下のような手順となる. 手法の概要を図 5 に示す.

(1) 推薦システムがユーザに対して広告行列 V を送信する.

(2) ユーザが (1) と自身のデータを用いて行列 U , V を更新するための勾配を計算し, 行列 U の自身の行を更新する. 行列 V の勾配に, ϵ -局所差分プライバシーを満たすノイズを加えたものを広告システムに送信する.

(3) 推薦システムが (2) で受け取った勾配を用いて行列 V を更新する.

(4) (1) ~ (3) を複数回繰り返す.

ここで, 上記を k 回繰り返すとすると, その都度勾配は異なるものとなるはずである. したがって $1 \leq s < t \leq k$ としたとき, s 回目送信した勾配と t 回目送信したものを攻撃者が比較することによって, その差分から何らかの情報を得る可能性がある. このプライバシーリスクを避けるために, 全ての勾配を比べた際に, 全ての組み合わせについて ϵ -差分プライバシーが

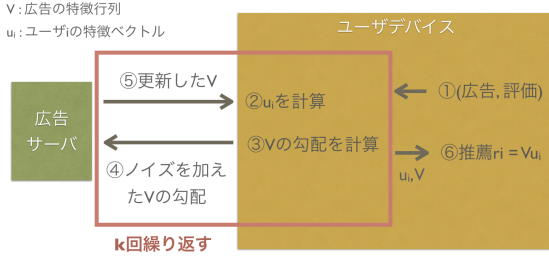


図5 学習の概要図

満たされる必要があると考える。

“ユーザが評価した広告”と“広告に対する評価値”の2つの情報についてプライバシー保護を行うための方法について述べる。ここで広告に対する評価値とは、ユーザの閲覧履歴や基本情報などユーザデバイスが持つ様々な情報から複合的に計算された値とする。

以下で2つのバージョンを提案する。バージョン1はサーバに送信する勾配にノイズを加えるというシンプルなものである。バージョン2では勾配の全ての次元をサーバに送信するのではなく、ランダムに選択された1次元のみノイズを加えて送信する。

a) バージョン1

まず、“ユーザが評価した広告”つまりユーザがある広告を評価したかどうか、という情報にノイズを加える方法について述べる。ユーザ i がある広告 j を評価したかどうかを示す値を y_{ij} とし、評価した場合を1そうでない場合を0とする。ここでユーザ i が広告 j に対する評価値 r_{ij} を考えたとき、 $\mathcal{M}$ に含まれる要素は全て0よりも大きいはずである。したがって次のような式変形が成り立つ。

$$\sum_{(i,j) \in \mathcal{M}} (r_{ij} - u_i^T v_j)^2 = \sum_{i=1}^n \sum_{j=1}^m y_{ij} (r_{ij} - u_i^T v_j)^2$$

したがって式(5)は以下のように変形できる。

$$\nabla_{v_j} \phi(U, V) = -\frac{2}{n} \sum_{i=1}^n y_{ij} u_i (r_{ij} - u_i^T v_j) \quad (6)$$

このとき、どの広告を評価したかという情報をプライバシー攻撃から守るためには、ユーザが広告を評価したかどうかを表すベクトル $Y_i = (y_{ij})_{1 \leq j \leq m}$ にノイズを加える必要がある。このときにノイズを加えるために、5.2.2節で紹介した Randomized Response を利用する。これを用いた場合、 y_{ij} にノイズを加えた値 y_{ij}^* は以下のように求められる。

$$y_{ij}^* = \begin{cases} 0, & \text{確率 } p/2 \\ 1, & \text{確率 } p/2 \\ y_{ij}, & \text{確率 } 1-p \end{cases}$$

次に、“広告に対する評価値”にノイズを加える方法について述べる。このとき、 d 次元ベクトル $g_{ij} = (g_{ij1}, \dots, g_{ijd}) = -2u_i(r_{ij} - u_i^T v_j)$ に対してノイズを加える。このノイズ η_{ijl} は

ラプラス分布 $Lap(\sigma)$ に基づくものとし、これによってノイズが加わった値 g_{ijl}^* は以下のように表される。

$$g_{ijl}^* = g_{ijl} + \eta_{ijl}$$

以上の方法でそれぞれのユーザが自身のデータにノイズを加え、 $\{(y_{ij}^*, g_{ij1}^*, \dots, g_{ijd}^*) : j = 1, \dots, m\}$ をサーバに送信したとする。このとき行列 V の更新のために用いる勾配はそれぞれのユーザの勾配の平均値となるため、 $\nabla v_j = n^{-1} \sum_{i=1}^n y_{ij} g_{ij}$ と表される。そしてこの値にノイズが加わったものは以下のよう表される。

$$\nabla_{v_j}^* = \frac{1}{n} \sum_{i=1}^n \left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^*$$

このようにノイズが加わったデータを用いて計算された勾配を用いて行列の値を更新する、という操作を k 回繰り返し、評価値を予測するための行列を最適化する。

b) バージョン2

プライバシー保護を強化し行列 V の誤差を小さくする[14]ために、ランダム化を用いる。勾配 $g_{ij} = (g_{ij1}, \dots, g_{ijd}) (j = 1, \dots, m)$ を全て広告システムに送信するのではなく、ランダムに選択した1次元のみを送信するというものである。手順は以下の通りである。

- (1) j と l をそれぞれ、 $\{1, \dots, m\}$ と $\{1, \dots, d\}$ からランダムに選択する。
- (2) $(x_i)_{jl} = -2y_{ij}u_{il}(r_{ij} - u_i^T v_j)$ を計算する。
- (3) $(x_i)_{jl}$ を $[-1, 1]$ に射影する。
- (4) ベルヌーイ分布 $Bern(\frac{(x_i)_{jl}(e^{\epsilon/k}-1)+e^{\epsilon/k}+1}{2(e^{\epsilon/k}+1)})$ に従う確率変数 T に基づいて $(x_i^*)_{jl}$ を決定する。

$$(x_i^*)_{jl} = \begin{cases} md \frac{e^{\epsilon/k}+1}{e^{\epsilon/k}-1}, & T = 1 \\ -md \frac{e^{\epsilon/k}+1}{e^{\epsilon/k}-1}, & T = 0 \end{cases}$$

- (5) $g_i^* = 0 + (x_i^*)_{jl}$ とする。

この時、勾配は以下のように求められる。

$$\nabla_V^* = \frac{1}{n} \sum_{i=1}^n g_i^*$$

各ユーザはランダムに選択された広告に対してのみノイズが付加された勾配を送信する。また、その値は $B (= md \frac{e^{\epsilon/k}+1}{e^{\epsilon/k}-1})$ または $-B$ のいずれかであるため、攻撃者はユーザがどのアイテムを評価したか、またそれらに対する評価値は知ることができない。

6.3.2 配信

配信については以下のような手続きが必要である。

- (1) 行列因子分解を用いてユーザデバイス上で評価値の最も高い広告を選択する。
- (2) 選択した広告の配信を広告サーバにリクエストする。
- (3) 広告サーバはユーザに広告を配信する。

この時、2.で推薦結果をそのまま広告サーバに送信すると、推薦された広告のカテゴリの偏りからユーザのプロファイルが

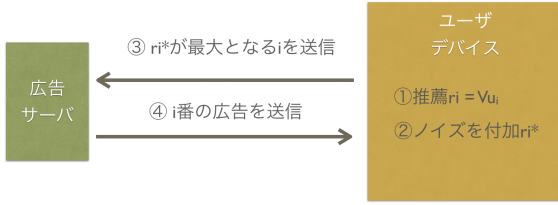


図6 配信の概要図

推測されてしまう可能性がある。プライバシー保護を行うために、Noisy Max [17] メカニズムを用いて推薦結果にノイズを加える。各評価値に対してランダムに生成されたラプラスノイズ $Lap(1/\epsilon)$ を加算し、その上で最大値のインデックスを広告サーバに送信する。概要を図 6.3.2 に示す。各評価値に付加するノイズを $\tau_j(1, \dots, m)$ とすると、求めるインデックス a は式 (7) で表される。

$$a = \arg \max_j r_{ij} + \tau_j \quad (7)$$

6.4 統計収集

広告の表示回数とクリックしたユーザの人数を計測する。

表示回数については、6.3.2 節で述べたように、ユーザが広告サーバに配信リクエストをする時点でプライバシーが保護されているため、実際に配信した回数を数え上げれば良い。

クリック数について、ユーザデバイスからクリック通知を送信してしまうと、ユーザがクリックした広告のカテゴリの偏りから、ユーザプロフィールが広告サーバに知られてしまいプライバシー保護が満たされなくなる。一定期間にユーザ i が広告 j をクリックしたかどうかを示す値を c_{ij} とし、クリックした場合を 1 そうでない場合を 0 とする。ユーザがどの広告をクリックしたかという情報のプライバシーを保護するために、ベクトル $C_i = (c_{ij})_{1 \leq j \leq m}$ にノイズを加えて広告サーバに送信する。広告サーバは各広告が何回クリックされたか、つまり各広告に対するクリック数の総和が分かれば良く、どのユーザがどの広告をクリックしたかを知る必要はない。よって、ノイズを加える方法として、ここでも Randomized Response を利用する。概要を図 6.4 に示す。 c_{ij} にノイズを加えた値 c_{ij}^* は以下のように求められる。

$$c_{ij}^* = \begin{cases} 0, & \text{確率 } p'/2 \\ 1, & \text{確率 } p'/2 \\ c_{ij}, & \text{確率 } 1 - p' \end{cases}$$

広告 j をクリックしたユーザの人数の推定値を c_j^* とすると、 $\sum_{i=1}^n c_{ij}^*$ の期待値は以下のような。

$$\sum_{i=1}^n c_{ij}^* = (1 - p')c_j^* + np'/2$$

これを c_j^* について解くことで、広告 j をクリックしたユー

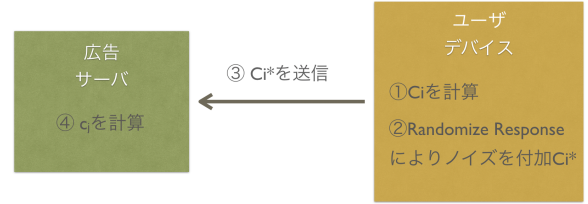


図7 統計収集の概要図

ザの人数が推定できる。

$$\begin{aligned} (1 - p')c_j^* &= \sum_{i=1}^n c_{ij}^* - np'/2 \\ c_j^* &= \frac{\sum_{i=1}^n c_{ij}^* - np'/2}{1 - p'} \\ c_j^* &= \frac{\sum_{i=1}^n (c_{ij}^* - p'/2)}{1 - p'} \\ c_j^* &= \sum_{i=1}^n \left(\frac{c_{ij}^* - p'/2}{1 - p'} \right) \end{aligned}$$

本研究では、簡単のため、クリックについての統計を求めたが、 $\{0, 1\}$ で表現できる情報であれば、多種多様な広告形態に合わせた統計量を測定することができる。具体的には、“広告をスキップしたかどうか”や“広告内ゲームをプレイしたかどうか”などがある。

7 検証

広告配信フェーズの学習時に ϵ_α -差分プライバシーを満たし、配信時に ϵ_β -差分プライバシーを満たすことを証明する。また、統計収集フェーズで ϵ_γ -差分プライバシーを満たすことを証明する。そのことにより、 $\epsilon = \epsilon_\alpha + \epsilon_\beta + \epsilon_\gamma$ とすると、提案システム全体では ϵ -差分プライバシーを満たす。

7.1 広告配信

7.1.1 学習時

a) バージョン 1

y_{ij} については、確率 $p = 2/(1 + e^{\epsilon_1/km})$, $w \in \{0, 1\}$ とすると、

$$\Pr[y_{ij}^* = w|y_{ij}] = \frac{p}{2} + (1 - |y_{ij} - y_{ij}^*|)(1 - p) \quad (8)$$

と表せる。

また、 $(g_{ijl})_{1 \leq l \leq d}$ については、ラプラス分布のスケールパラメータを $\sigma = \frac{\Delta dkm}{\epsilon_2}$, $\Delta = \max_{i,j} r_{ij} - \min_{i,j} r_{ij}$, $\eta_{ijl} \sim Lap(\sigma)$, $w \in \mathbb{R}^d$ とすると、 $g_{ijl}^* = g_{ijl} + \eta_{ijl}$ となるとき、 $\eta_{ijl} = g_{ijl}^* - g_{ijl}$ より、

$$\Pr[g_{ijl}^* = w|g_{ijl}] = \frac{1}{2\sigma} \exp\left(\frac{-|g_{ijl}^* - g_{ijl}|}{\sigma}\right)$$

と表せる。よって、 $\tilde{w}_{ji} \in \mathbb{R}^d$, $w_y \in \{0, 1\}$, $w_g \in \mathbb{R}^d$ とすると、以下の式が成り立つ。

$$\begin{aligned}
& \frac{\Pr \left[\left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = \tilde{w}_{ji} | y_{ij}, g_{ij} \right]}{\Pr \left[\left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = \tilde{w}_{ji} | y'_{ij}, g'_{ij} \right]} \\
&= \frac{\Pr [y_{ij}^* = w_y | y_{ij}] \Pr [g_{ij}^* = w_g | g_{ij}]}{\Pr [y_{ij}^* = w_y | y'_{ij}] \Pr [g_{ij}^* = w_g | g'_{ij}]} \\
&\leq \frac{\max_{y'_{ij}} p/2 + (1 - |y_{ij} - w_y|)(1-p)}{\min_{y'_{ij}} p/2 + (1 - |y'_{ij} - w_y|)(1-p)} \cdot \prod_{l=1}^d \exp \left(\frac{|g_{ijl} - g'_{ijl}|}{\sigma} \right) \\
&= \frac{2-p}{p} \cdot \prod_{l=1}^d \exp \left(\frac{|g_{ijl} - g'_{ijl}|}{\sigma} \right) \\
&= \frac{2 - 2/(1 + e^{\epsilon_1/km})}{2/(1 + e^{\epsilon_1/km})} \cdot \prod_{l=1}^d \exp \left(\frac{\epsilon_2 |g_{ijl} - g'_{ijl}|}{\Delta dkm} \right) \\
&= e^{(\epsilon_1 + \epsilon_2)/km}
\end{aligned}$$

よって, $w \in \mathbb{R}^{m \times d}$ とすると, 繰り返しの各回で広告システムが得る ∇_V^* について以下の式が成り立つ.

$$\begin{aligned}
& \Pr [\nabla_V^* = w | D] \\
&= \prod_{j=1}^m \Pr [\nabla_{v_j}^* = w_j | D] \\
&= \prod_{j=1}^m \Pr \left[\sum_{i=1}^n \left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = n w_j | D \right] \\
&= \prod_{j=1}^m \sum_{\substack{\tilde{w}_{j1}, \dots, \tilde{w}_{jn} \\ \tilde{w}_{j1} + \dots + \tilde{w}_{jn} = n w_j}} \prod_{i=1}^n \Pr \left[\left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = \tilde{w}_{ji} | D \right] \\
&= \prod_{j=1}^m \sum_{\substack{\tilde{w}_{j1}, \dots, \tilde{w}_{jn} \\ \tilde{w}_{j1} + \dots + \tilde{w}_{jn} = n w_j}} \prod_{i=1}^n \Pr \left[\left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = \tilde{w}_{ji} | y_{ij}, g_{ij} \right]
\end{aligned}$$

データベース D と D' はユーザ q のデータのみが異なっているとすると,

$$\begin{aligned}
& \frac{\Pr [\nabla_V^* = w | D]}{\Pr [\nabla_V^* = w | D']} \\
&\leq \prod_{j=1}^m \max_{\substack{\tilde{w}_{j1}, \dots, \tilde{w}_{jn} \\ \tilde{w}_{j1} + \dots + \tilde{w}_{jn} = n w_j}} \prod_{i=1}^n \frac{\Pr \left[\left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = \tilde{w}_{ji} | y_{ij}, g_{ij} \right]}{\Pr \left[\left(\frac{y_{ij}^* - p/2}{1-p} \right) g_{ij}^* = \tilde{w}_{ji} | y'_{ij}, g'_{ij} \right]} \\
&= \prod_{j=1}^m \max_{w_q} \frac{\Pr \left[\left(\frac{y_{qj}^* - p/2}{1-p} \right) g_{qj}^* = \tilde{w}_{jq} | y_{qj}, g_{qj} \right]}{\Pr \left[\left(\frac{y_{qj}^* - p/2}{1-p} \right) g_{qj}^* = \tilde{w}_{jq} | y'_{qj}, g'_{qj} \right]} \\
&\leq \prod_{j=1}^m e^{(\epsilon_1 + \epsilon_2)/km} \\
&= e^{(\epsilon_1 + \epsilon_2)/k}
\end{aligned}$$

となる. これを k 回繰り返すので, $\epsilon_\alpha = \epsilon_1 + \epsilon_2$ とすると, 最終的に得られる行列 V について以下の式が成り立つ.

$$\begin{aligned}
\frac{\Pr [V = \bar{V} | D]}{\Pr [V = \bar{V} | D']} &\leq \max_{a_1, \dots, a_k} \prod_{t=1}^k \frac{\Pr [\nabla_V^{*,t} = a_t | D]}{\Pr [\nabla_V^{*,t} = a_t | D']} \\
&\leq e^{\epsilon_1 + \epsilon_2} \\
&= e^{\epsilon_\alpha}
\end{aligned}$$

よって最終的に得られる V は ϵ_α -差分プライバシーを満たす.

b) バージョン 2

$$\begin{aligned}
& (x_i)_{j,l} \in [-1, 1] \text{ で, } T \text{ はベルヌーイ分布 } \text{Bern}\left(\frac{(x_i)_{j,l}(e^{\epsilon_\alpha/k} - 1) + e^{\epsilon_\alpha/k} + 1}{2(e^{\epsilon_\alpha/k} + 1)}\right) \\
& \text{に従うので, } w \in \mathbb{R}^{m \times d} \text{ とすると,} \\
& \frac{\Pr [x_i^* = w | x_i]}{\Pr [x_i^* = w | x'_i]} = \frac{1/m \cdot 1/d \cdot \Pr [T = 1 | x_i]}{1/m \cdot 1/d \cdot \Pr [T = 1 | x'_i]} \\
&\leq \frac{\max_{x_i} \Pr [T = 1 | x_i]}{\min_{x'_i} \Pr [T = 1 | x'_i]} \\
&= \frac{\max_{x_i \in [-1, 1]} \left\{ (x_i)_{j,l} (e^{\epsilon_\alpha/k} - 1) + e^{\epsilon_\alpha/k} + 1 \right\}}{\min_{x'_i \in [-1, 1]} \left\{ (x'_i)_{j,l} (e^{\epsilon_\alpha/k} - 1) + e^{\epsilon_\alpha/k} + 1 \right\}} \\
&= e^{\epsilon_\alpha/k}
\end{aligned}$$

となる. よって, $w \in \mathbb{R}^{m \times d}$ とすると, 繰り返しの各回で広告システムが得る ∇_V^* について以下の式が成り立つ.

$$\begin{aligned}
\Pr [\nabla_V^* = w | D] &= \Pr \left[\sum_{i=1}^n x_i^* = n w | D \right] \\
&= \sum_{\substack{\tilde{w}_1, \dots, \tilde{w}_n \\ \tilde{w}_1 + \dots + \tilde{w}_n = n w}} \prod_{i=1}^n \Pr [x_i^* = \tilde{w}_i | D] \\
&= \sum_{\substack{\tilde{w}_1, \dots, \tilde{w}_n \\ \tilde{w}_1 + \dots + \tilde{w}_n = n w}} \prod_{i=1}^n \Pr [x_i^* = \tilde{w}_i | x_i]
\end{aligned}$$

データベース D と D' はユーザ q のデータのみが異なっているとすると,

$$\begin{aligned}
\frac{\Pr [\nabla_V^* = w | D]}{\Pr [\nabla_V^* = w | D']} &\leq \max_{\substack{\tilde{w}_1, \dots, \tilde{w}_n \\ \tilde{w}_1 + \dots + \tilde{w}_n = n w}} \prod_{i=1}^n \frac{\Pr [x_i^* = \tilde{w}_i | x_i]}{\Pr [x_i^* = \tilde{w}_i | x'_i]} \\
&= \max_{\tilde{w}_p} \frac{\Pr [x_p^* = \tilde{w}_p | x_p]}{\Pr [x_p^* = \tilde{w}_p | x'_p]} \\
&\leq e^{\epsilon_\alpha/k}
\end{aligned}$$

これを k 回繰り返すので, 最終的に得られる行列 V について以下の式が成り立つ.

$$\begin{aligned}
\frac{\Pr [V = \bar{V} | D]}{\Pr [V = \bar{V} | D']} &\leq \max_{a_1, \dots, a_k} \prod_{t=1}^k \frac{\Pr [\nabla_V^{*,t} = a_t | D]}{\Pr [\nabla_V^{*,t} = a_t | D']} \\
&\leq e^{\epsilon_\alpha}
\end{aligned}$$

よって最終的に得られる V は ϵ_α -差分プライバシーを満たす.

7.1.2 配信時

配信時については, Noisy Max のノイズが従う分布をラプラス分布 $Lap(1/\epsilon_\beta)$ とすると, 選択される広告番号は ϵ_β -差分プライバシーを保証する [17].

7.2 統計収集

$p' = 2/(1 + e^{\epsilon\gamma})$, $w \in \{0, 1\}$ とすると, 式 (8) より,

$$\begin{aligned} \frac{\Pr[y_{ij}^* = w|y_{ij}]}{\Pr[y_{ij}^* = w|y'_{ij}]} &\leq \frac{\max_{y_{ij}} p/2 + (1 - |y_{ij} - w|)(1 - p)}{\min_{y'_{ij}} p/2 + (1 - |y'_{ij} - w|)(1 - p)} \\ &= \frac{2 - p}{p} \\ &= \frac{2 - 2/(1 + e^{\epsilon\gamma})}{2/(1 + e^{\epsilon\gamma})} \\ &= e^{\epsilon\gamma} \end{aligned}$$

よって, 得られる c_j^* は ϵ_γ -差分プライバシーを保証する.

8 おわりに

本研究では, プライバシー保護の手法として局所差分プライバシーを用いたインターネット広告システムを提案している. 提案システムでは, 学習時・配信時・統計収集時の3箇所に局所差分プライバシーを用いた.

インターネット広告システムにプライバシー保護手法の中でも強力な局所差分プライバシーを適用した研究は著者の知る限りでは未だ存在せず, プライバシー保護インターネット広告システムという分野において貢献している.

今後の課題については以下にまとめる.

(1) 提案システムにおいて, プライバシー保護を満たすことはできたものの, 有用性については, 改善の余地があると考えている.

(2) 提案システムはプライバシー保護と個人最適化を重要視しているため, 現存の広告システムに關与しているあらゆるサービスには, 必ずしも適応していない. 現存の広告システムのフレームワークを保持したまま, 提案システムを適用する方法が必要である.

9 謝 辞

本研究は JSPS 科研費基盤研究 (S) No. 17H06099, (A) No. 18H04093, 若手研究 No. 19K20269 の助成を受けたものです. 本研究を進めるにあたり, 非常に多くの方に御世話になりました. ここに深く感謝の意を表します.

文 献

- [1] Council of the European Union and European Parliament: REGULATION (EU) 2016/ 679 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL - of 27 April 2016 - on the Protection of Natural Persons with Regard to the Processing of Personal Data and on the Free Movement of Such Data, and Repealing Directive 95/ 46/ EC (General Data Protection Regulation).
- [2] Reuters Staff: グーグル、2年以内に広告目的のクッキー利用制限へ, <https://jp.reuters.com/article/alphabet-google-privacy-idJPKBN1ZD2UJ> (2020).
- [3] Evfimievski, A., Gehrke, J. and Srikant, R.: Limiting Privacy Breaches in Privacy Preserving Data Mining, *Proceedings of the Twenty-Second ACM SIGMOD-SIGACT-SIGART Symposium on Principles of Database Systems*, PODS '03, Association for Computing Machinery, p. 211–222 (2003).
- [4] 個人情報保護委員会: 「個人情報保護法 いわゆる3年ごと見直し制度改正大綱(骨子)」の公表について, <https://www.ppc.go.jp/news/press/2019/20191129/>.
- [5] Guha, S., Cheng, B. and Francis, P.: PrivAd: Practical Privacy in Online Advertising, *Proceedings of the 8th USENIX Conference on Networked Systems Design and Implementation*, NSDI'11, Boston, MA, USENIX Association, pp. 169–182 (2011).
- [6] Backes, M., Kate, A., Maffei, M. and Pecina, K.: ObliviAd: Provably Secure and Practical Online Behavioral Advertising, *2012 IEEE Symposium on Security and Privacy*, pp. 257–271 (2012).
- [7] Tulabandhula, T., Vaya, S. and Dhar, A.: Privacy-Preserving Targeted Advertising, *arXiv:1710.03275 [cs]* (2017).
- [8] Hardt, M. and Nath, S.: Privacy-Aware Personalization for Mobile Advertising, *Proceedings of the 2012 ACM Conference on Computer and Communications Security*, CCS '12, New York, NY, USA, ACM, pp. 662–673 (2012).
- [9] Calandrinio, J. A., Kilzer, A., Narayanan, A., Felten, E. W. and Shmatikov, V.: "You Might Also Like:" Privacy Risks of Collaborative Filtering, *2011 IEEE Symposium on Security and Privacy*, pp. 231–246 (2011).
- [10] Balu, R. and Furon, T.: Differentially Private Matrix Factorization Using Sketching Techniques, *Proceedings of the 4th ACM Workshop on Information Hiding and Multimedia Security*, pp. 57–62 (2016).
- [11] Berlioz, A., Friedman, A., Kaafar, M. A., Boreli, R. and Berkovsky, S.: Applying Differential Privacy to Matrix Factorization, *Proceedings of the 9th ACM Conference on Recommender Systems*, Association for Computing Machinery, Inc, pp. 107–114 (2015).
- [12] Liu, Z., Wang, Y.-X. and Smola, A. J.: Fast Differentially Private Matrix Factorization, *arXiv:1505.01419 [cs]* (2015).
- [13] Narayanan, A. and Shmatikov, V.: Robust De-Anonymization of Large Sparse Datasets, *2008 IEEE Symposium on Security and Privacy (Sp 2008)*, pp. 111–125 (2008).
- [14] Shin, H., Kim, S., Shin, J. and Xiao, X.: Privacy Enhanced Matrix Factorization for Recommendation with Local Differential Privacy, *IEEE Transactions on Knowledge and Data Engineering*, Vol. 30, No. 9, pp. 1770–1782 (2018).
- [15] Dwork, C.: Differential Privacy: A Survey of Results, *Theory and Applications of Models of Computation* (Agrawal, M., Du, D., Duan, Z. and Li, A.(eds.)), Lecture Notes in Computer Science, Berlin, Heidelberg, Springer, pp. 1–19 (2008).
- [16] Warner, S. L.: Randomized Response: A Survey Technique for Eliminating Evasive Answer Bias, *Journal of the American Statistical Association*, Vol. 60, No. 309, pp. 63–66 (1965).
- [17] Dwork, C. and Roth, A.: The Algorithmic Foundations of Differential Privacy, *Foundations and Trends® in Theoretical Computer Science*, Vol. 9, No. 3–4, pp. 211–407 (2013).