

SNS 記事におけるプライバシー侵害に関わる特徴量の推定

An Estimate of Privacy Features in SNS article

輪島 幸治 (Koji WAJIMA)¹・古川 利博 (Toshihiro FURUKAWA)²・佐藤 哲司 (Tetsuji SATOH)³¹筑波大学 図書館情報メディア研究科・²東京理科大学 工学部情報工学科 教授³筑波大学 図書館情報メディア系 教授

[Abstract] The social concern with Online media's communication and Privacy have been growing for the last several years. A large number of unspecified users are investigating Online media's communication in the Internet. Therefore, inappropriate behavior in online media leads to Enjyo and Framing. Enjyo and Framing lead to privacy issue in Online Media. Additionally, GDPR(General Data Protection Regulation) has provoked a great deal of controversy in regards to that privacy issue. For this purpose, We need to prevent the occurrence of privacy issues in online media. The feature quantity in Online media needs to be examined in detail now. This paper is intended as an investigation of feature quantity for the discovery of privacy issue. In this Paper, We investigate online media's invasion of privacy on SNS article. We comprehensively extracted text information's feature quantity using existing research. Feature quantity is used, for 10 groups, 33 types and 2,073 dimensions in text information. Proposed Method consists of feature transformation using NMF and the NMF base on evaluation using SVR. The validity of the feature quantity is verified by classifier and Bayesian Network. It was found from the result that has an impact on Privacy. The classifier's evaluation yielded the robust result. And, The Bayesian Network's evaluation has been gained causal relationship feature quantity. We report that research result.

[キーワード] プライバシー侵害, NMF, SVR, Bayesian Network, LDA

1. はじめに

スマートフォンなどの常時携帯型の端末が広く普及したことで、オンライン上でのコミュニケーション (SNS: Social Network Service) が急速に広まっている。バイラル・マーケティング[1]などユーザ相互の口コミ情報などを SNS 上で拡散する情報伝搬手法を用いた商品やサービスの宣伝活動も認識されてきている。一方で、SNS など不特定多数による情報発信を基盤とするメディアでは、ふとしたことから、他者のプライバシーを侵害したり[2]¹、誹謗中傷に加担してしまうことが問題となっている。特に、バイラル・マーケティングの過程などで、このような問題が発生すると、一個人の問題では済まなくなり、社会的な損失も無視できない規模に拡大することがある[3]。また、プライバシーに関しては、2018年5月25日に適用されたEUの一般データ保護規則 (General Data Protection Regulation: GDPR)²がある。GDPRは、オンライン・サービスも含め、個人データの処理、および個人データを欧州経済領域 (European Economic Area : EEA) から第三国に移転するための法的要件の規定である。このように、近年、プライバシーに関わるデータの処理には、より慎重に扱い、かつ速やかな対応が求められている。本研究では、オンライン・サービスの一つである SNS で投稿者が意図したか否かに関わらず、共有されたコンテンツが、他者に不愉快な感情を与えるコンテンツであるか、あるいは SNS で共有されたコンテンツが、他者がプライバシー侵害と判断するかを、コンテンツから判断できるかを論じる。

提案手法では、テキスト情報より得られるテキストの特性を表す単語や単語集合である特徴量を用いて評価を行う。既存研究で明らかになっている様々な特徴量をコンテンツから抽出し、抽出した特徴量からプライバシー侵害に影響がある特徴量を判別するための手法を提案する。数1,000次元に及ぶコンテンツの特徴量を、非負値行列因子分解 (NMF: Nonnegative Matrix Factorization) によって次元圧縮を行う。そして、SVRを用いてプライバシー侵害に強く影響している基底を推定する。最後に、提案手法を実際の投稿記事集合に適用して、分類器評価およびBayesian Networkを用いてプライバシー侵害に影響がある特徴量を論じる。

¹ <http://sp.trendmicro.co.jp/jp/about-us/press-releases/articles/20130829025049.html>² <https://www.jetso.go.jp/world/europe/eu/gdpr/>

提案手法は、既存研究の特徴量を用いて NMF, SVR を組み合わせて評価を行っているという特徴を有している。また、特徴量は既存研究のテキスト情報の特徴量を網羅的に抽出しており、33 個の特徴量、2,073 次元の特徴量を評価した。加えて、提案手法の基底評価では、55 名に実施したプライバシー侵害のアンケートを用いている。提案手法を実装し、プライバシー侵害の原因となる特徴量の推定を行った。得られた結果の評価を行い、提案手法の有効性と限界を論じる。

以下、本論文では、2 章で関連研究を概観し、本研究の位置づけを明らかにする。3 章で特徴量の評価手法を提案し、4 章で実験環境と評価方法を述べる。そして 5 章で実験結果と考察を示し、最後に 6 章でまとめと今後の課題を示す。

2. 関連研究

2. 1 プライバシーに関する研究

プライバシーの概念には、多義性と文脈依存性があることが知られており[4]、特定の単語が出現するか否かだけではプライバシーに言及しているかを論じることはできない。プライバシー保護の原則に、OECD プライバシー・ガイドラインに基づいて 1982 年に行政管理庁のプライバシー保護研究会が公表したプライバシー保護研究会報告書の 5 つの基本原則がある[5]。基本原則は、(1)収集制限の原則、(2)利用制限の原則、(3)個人参加の原則、(4)適正管理の原則、(5)責任明確化の原則である。基本原則に則ってプライバシー保護に関する様々な研究[6]がなされてきた。プライバシー保護に関する既存研究では、プライバシーに関連する投稿の検索[7]など情報収集に関する研究、情報開示抵抗感・二次利用侵害懸念の調査分析[8]など情報処理に関する研究。また、プライバシーの漏洩検知と公開範囲の設定[9]など情報拡散に関する研究として広く知られている。特に情報拡散の研究では、個人のプライバシーに関する情報が流用・歪曲される場合が多いことから、漏洩検知と公開範囲の設定[9]だけでなく、プライバシー侵害シーンの抽出[10]など多岐にわたる。

本研究で評価対象とする SNS の投稿記事におけるプライバシー侵害は、情報の収集、処理、拡散が容易に行えることから、炎上やフレーミングと言う事象がしばしば発生している。ここで、炎上とは、コミュニケーション相手と異なる第三者が、逸脱や不適切と判断されるメッセージを問題視し、SNS 上で指摘することに端を発し、批判が集中する事象である[3]。一方、フレーミングとは、コンピュータ上のコミュニケーションで当事者のいずれかが徹底的で攻撃的な相互行為を取る事象である[11]。炎上やフレーミングの多くは SNS のコミュニケーションで過激な発言や行為を繰り返す事象であり、いじめ、嫉妬、批判などによって、冷静さを失った投稿者が行うことが多い。そして意図せず、他者のプライバシー侵害に相当する記事を投稿し、社会的な損失も無視できない規模の問題へ発展する場合がある。

2. 2 テキスト情報の特徴量に関する研究

本研究では、投稿者が意図したか否かに関わらず、SNS 上で共有されたコンテンツが、他者のプライバシー侵害に該当するかを、投稿記事の内容であるコンテンツから判断する方法を論じる。提案手法では、情報検索[12][13]で広く使用されているテキスト情報を評価の特徴量とする。提案手法は既存研究のテキスト情報の特徴量を網羅的に抽出し、有効性を論じているところに特徴がある。類似の研究には、SNS 投稿記事テキスト情報とパーソナルデータの組み合わせによるプライバシー侵害を検知[14][15]の研究がある。本研究では、類似研究と比較し、2.3 節、2.4 節で後述する特徴量も含めると特徴量の多さで既存研究を凌駕している。本研究で使用するテキスト情報の特徴量を表 1 に示す。表 1 の 8 種類の特徴量でコンテンツである文面の表層の特徴を捉えることとした。本研究の文字種には、ひらがな、カタカナ、漢字、アルファベット、数字、半角記号、空白記号、全角記号を採用した。

表 1 表層情報の特徴量の詳細

特徴量名	次元数	値の定義/例	文献/脚注
文の数	1	整数値(1 文中の[.][?][!]の合計頻度)	-
読点	1	整数値(1 文中の[,]の頻度)	-
句点	1	整数値(1 文中の[.])	-
文の長さ	1	整数値(1 文中の総バイト数)	-
文字種(頻度)	8	整数値(ひらがな, カタカナ, 漢字)など	[16] ³
文字種(比率)	8	実数値(同上)	[16] ³
読点間距離	1	実数値(算出)	[17]
漢字含有率	1	実数値(算出)	[17]

2. 3 話題抽出に関する研究

テキスト情報の特徴量に関する研究に、話題(トピック)を抽出する研究がある。テキスト情報のトピックの特徴量とは、テキスト情報が言及する話題である。トピックを抽出する手法では、テキストは複数の単語(語彙)が集まったもの(Bag-of-words)として捉える。そして、類似した語彙の集合で話題(トピック)が形成されるとしている。話題抽出で用いられる基本的なアルゴリズムには、LSI (Latent Semantic Indexing) や LDA (Latent Dirichlet Allocation) [18]がある。LSI は、文書集合を行列 N として、低ランク行列の積 $U^T H$ にフロベニウスノルムが最小になるように行列分解する手法である。LDA は、文書集合の行列 N の文書が低ランク行列 θ と Φ をパラメータとして持つカテゴリ分布 Λ から生成されると仮定する手法である。LSI を図 1 に LDA を図 2 に示す。

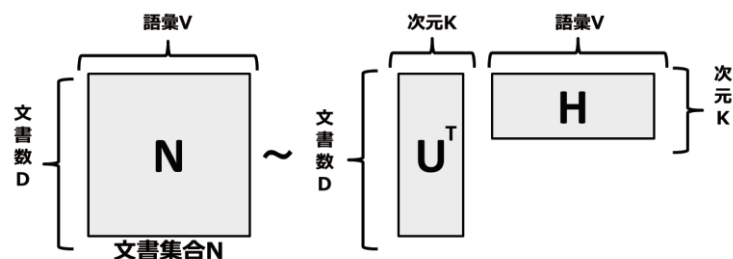


図 1 LSI (Latent Semantic Indexing)

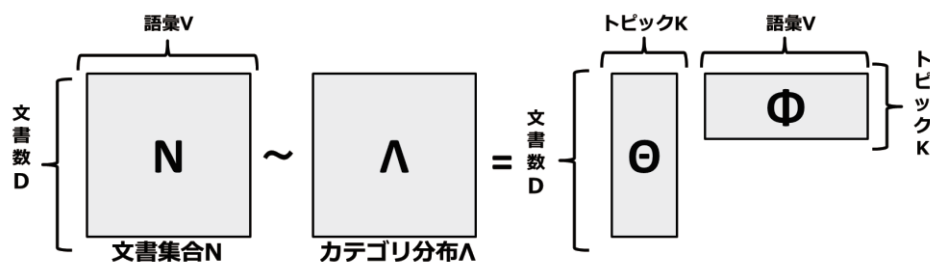


図 2 LDA (Latent Dirichlet Allocation)

ここで、文書集合である行列 N は 1 行が 1 文書、各列は語彙に相当し、要素は語彙 V の出現頻度である。

LSI の場合、次元 K が話題であり、文書に対する次元 K の重みである U の要素が特徴量に相当する。

LDA の場合、トピック K が話題であり、文書に対するトピック K の割り当て確率であるトピック分布 θ の要素が話題に相当する。アルゴリズムでは、LSI の次元数や LDA のトピック数は任意の数の設定が行える。4 章で後述するが、LSI, LDA の次元数には各 300 を設定した。

³ Unicode 10.0 Character Code Charts : <http://www.unicode.org/charts/>

4 語彙の特徴を表す辞書

既存研究の辞書では、意味情報の特徴量を抽出する。意味情報の抽出では、個別の特徴量抽出の研究が行われているが、本研究では、ニクラス・ルーマンの社会システム理論のゼマンティックという概念に基づき、既存研究の辞書を用いる[19]。ゼマンティックとは、高度に一般化され、状況に依存せずに使用できる意味である。既存研究の辞書を用いることで、一意的な意味情報の特徴量抽出が行える。著者らが知る限り、従来研究で作成された日本語のテキストの既存の辞書は、21 個ある。加えて、プライバシーに関する単語を抽出した既存の辞書が2 個ある。したがって、本研究では表2 に示す 23 個の辞書を用いる。

表2 既存研究の特徴量の詳細

特徴量名	次元数	語数	値の定義/例	文献/脚注
語種	7	6, 519	和語, 漢語, 外来語	脚注 ⁴
基本語(1)	2	697	基本語二千に選定	脚注 ⁴
基本語(2)	6	6, 519	外国語学習基本語彙	脚注 ⁴
基本語(3)	2	424	基本語二千・六千に選定	脚注 ⁴
意味分類コード(1)	233	697	体の類・抽象的關係	脚注 ⁴
意味分類コード(2)	487	6, 519	(同上)	脚注 ⁴
意味分類コード(3)	307	424	(同上)	脚注 ⁴
機能表現	122	29, 262	O, B判断	[20], 脚注 ⁵
文末モダリティ	32	32	可能性, 表出	[21]
質問文末表現	38	38	質問, 回答	[22]
名詞比率	1	–	実数(算出)	[23]
MVR	1	–	実数(算出)	[23]
IPA 品詞	14	–	連体詞, 接頭詞, 名詞	脚注 ⁶
固有名詞	4	–	実数(算出)	脚注 ⁶
拡張固有表現	132	18, 075	数値表現	脚注 ⁷
PS ワード	1	20	プライバシー, 肖像権	[23]
LPSW	1	49	プライバシー, 侵害	[7]
評価極性情報 (用言編)	4	5, 280	ネガ(経験)	[24], 脚注 ⁵
(名詞編)	51	13, 314	～になる(状態) 客観	[25], 脚注 ⁵
感情極性 (頻度)	2	110, 250	ポジティブ・ネガティブ	[26], 脚注 ⁸
(比率)	2	–	実数(算出)	[26][27], 脚注 ⁷
(平均値)	1	–	実数(算出)	[26][27], 脚注 ⁷
評価値表現	1	5, 234	–	[28], 脚注 ⁹

表2のうち、単語感情極性対応表は、日本語と英語の辞書があり、本研究では日本語の辞書を用いた。また、特徴量中のための単語には、辞書の見出し語と読みを用いた。表2のNAIST-JENEは、Wikipediaの見出し語に対し、関根の拡張固有表現階層の定義に基づき、固有表現クラスを付与した辞書である。なお、表2の日本語教育基本語彙、意味分類体語彙表、分類項目一覧表は調査研究目的に作成された一覧表である。次ページに記載の基準で、加工・見出し語を選定した。

⁴ 『日本語教育のための基本語彙調査』データ：<http://mmsrv.ninjal.ac.jp/bvjsl84/>

⁵ 機能表現タグ付与コーパス、日本語評価極性辞書：<http://www.cl.ecei.tohoku.ac.jp/index.php>

⁶ ipadic version 2.7.0 ユーザーズマニュアル：<http://chasen.naist.jp/snapshot/ipadic/ipadic/doc/ipadic-ja.pdf>

⁷ NAIST Japanese ENE Dictionary on Wikipedia：<https://github.com/masayu-a/NAIST-JENE>

⁸ 単語感情極性対応表：http://www.lr.pi.titech.ac.jp/~takamura/pndic_ja.html

⁹ 評価値表現辞書：http://www.syncha.org/evaluative_expressions.html

(1) 見出し語の加工

空白記号, 「-」 「0」 「など」 を除去

(2) 対象外の見出し語

「」 「-」 「[]」 「→」 「/」 「()」 「・」 「その他」 を含む語

3. 提案手法

3. 1 概要

本研究では, SNS 投稿記事の特徴量を用いて, プライバシー侵害と因果関係のある特徴量を推定する. 提案手法ではプライバシー侵害の判断は明示されない言語外知識であり, 言語として記述された表層情報で直接評価することは困難である. そこで, 複数の特徴量や手法を組み合わせ, 抽出した特徴量からプライバシー侵害に影響がある SNS 投稿記事を判別するための手法を提案する. 本研究では, 既存研究で明らかになっている様々な特徴量をコンテンツから抽出し, 特徴量変換を行い, プライバシー侵害に影響の大きい特徴量を明らかにする. 提案手法の概要を図3に示す. 図3の(1)から(3)は, 提案手法の手順を示すステップである.

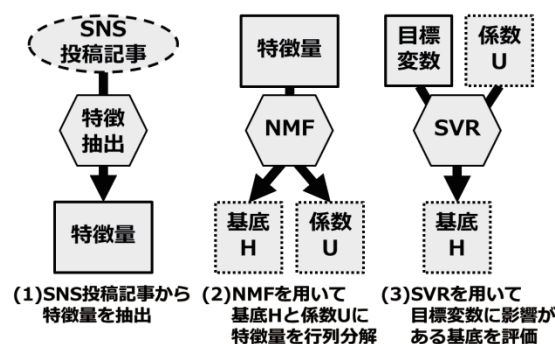


図3 提案手法の概要

図3(1)の特徴量の抽出では2.2節, 2.3節, 2.4節に示した特徴量を抽出している. 抽出する特徴量の数は33個, 各特徴量を水平方向に連結した際の次元数は2,073次元である. 採用した特徴量は既存研究でテキスト情報の評価の研究で用いられている特徴量を網羅的に抽出した. (2)および(3)が, 本論文で提案する, プライバシー侵害に関わる基底の推定手順である. プライバシー侵害に関わる基底の評価で, プライバシー侵害に影響の大きい特徴量が明らかになる. プライバシー侵害に影響の大きい特徴量を評価するために本研究の目標変数には, プライバシー侵害のアンケート結果を用いた. 提案手法では, 抽出した特徴量をNMFを用いて行列分解を行い, 特徴量の変換を行う. そして, SVRを用いて, プライバシー侵害に影響の大きいNMFの基底の推定を行う.

提案手法は, プライバシー侵害の SNS 投稿記事の評価している従来研究 [7][14][15][29]と比較し, SNS 投稿記事からテキスト情報の特徴量を網羅的に抽出している. また, プライバシー侵害の従来研究で, NMFによる行列分解で特徴量変換し, プライバシー侵害に該当するテキスト情報の特徴量を評価する研究は著者の知る限り, 行われていない. 加えて, 本研究では行列分解の結果得られる基底の評価を係数値による重み付けだけでなくSVRと組み合わせて評価している. 本章では, まず, 特徴量の変換手法を示す. そして, プライバシー侵害に影響の大きい基底の評価手法を述べる.

3. 2 NMF を用いた特徴量の変換

プライバシー侵害に影響の多い特徴量の評価を行うため, 本研究では, 特徴量の変換を行う. 特徴量変換を行うことで, 影響の多い変換された特徴を明らかにし, 変換された特徴を基にした特徴量の評価や因果関係の推定が行える. 本研究の特徴量の変換には, 非負値行列因子分解 (NMF: Nonnegative Matrix Factorization) を用いる[30]. NMFは観測行列 Y を基底行列 H と係数行列 U の積に分解するアルゴリズムである.

以後, 文書の特徴量の行列を観測データ行列と見なし, データをベクトルで表記する. したがって, 観測行列の要素である特徴量の頻度や値はベクトル要素である. 本研究でNMFの適用を行う観測行列は, 2.2節, 2.3節, 2.4節に示した特徴量を SNS 投稿記事より抽出し, 水平方向に連結した行列である. NMFの簡略化した分解表現を式(3)に示す.

$$Y \approx HU \quad (1)$$

式(1)の観測行列 $\mathbf{Y}$ は、文書数($i = 1, \dots, N$)、特徴量の次元数($j = 1, \dots, K$)から構成される N 行 K 列の長方形行列である。4章で後述するが、 N 行は文書数であり96、次元数 K は特徴量の数であり2,073である。したがって、本研究の観測行列 $\mathbf{Y}$ は、96行2,073列の長方形行列である。ここで、観測行列 $\mathbf{Y}$ 、基底行列 $\mathbf{H}$ の要素はベクトルである。NMFでは、観測行列 $\mathbf{Y}$ の次元数 K よりも、基底数 M を小さく設定することで、特徴量変換が行える。また、NMFでは、観測行列 $\mathbf{Y}$ が、相関行列や分散共分散行列であることを前提としていない。したがって、主成分分析(Principal Component Analysis)や因子分析(Factor Analysis)などと異なり、観測行列 $\mathbf{Y}$ の各次元の特性に依存せず適用できる。基底行列 $\mathbf{H}$ の基底数を($m = 1, \dots, M$)とした際のNMFによる観測行列の分解を式(2)に示す。

$$(y_{ij})_{NK} \simeq \sum_{m=1}^M h_{j,m} u_{m,i} \quad (2)$$

式(2)の $(y_{ij})_{NK}$ は観測行列 $\mathbf{Y}$ を表す。また、 $h_{j,m}$ は基底行列 $\mathbf{H}$ の成分、 $u_{m,i}$ は係数行列 $\mathbf{U}$ の成分を表す。ここで、 $\sum_{m=1}^M h_{j,m} u_{m,i}$ は、 $h_{j,m}^T u_{m,i}$ に相当する。したがって、 $h_{j,m}^T u_{m,i}$ は、観測行列 $(y_{ij})_{NK}$ と等しくなるべき値である。しかし、行列分解では一般に誤差が発生する。また、行列分解される行列 $\mathbf{H}, \mathbf{U}$ は、一意に決まらない。このため、NMFの行列分解は、観測行列 $\mathbf{Y}$ と $\mathbf{H}, \mathbf{U}$ の誤差を最小化する行列 $\mathbf{H}, \mathbf{U}$ を求める最適化問題である。NMFでは、乖離度規準を定義し、規準に基づく更新式の反復で最適解を求める。NMFの乖離度規準は、観測行列 $\mathbf{Y}$ の生成プロセスによって異なる。本研究では、観測行列 $\mathbf{Y}$ の生成プロセスに、非負の整数の確率分布であるPoisson分布を仮定した。したがって、乖離度規準には、一般化Kullback-Leiblerダイバージェンス[31]を採用した。また、最適化で用いる更新式の反復には、Multiplicative Update[32]を用いている。乖離度規準を式(3)に、乖離度規準に基づく最適化する更新式を式(4)に示す。

$$D(\mathbf{Y}, \mathbf{H}\mathbf{U}) = y_{ij} \log \frac{y_{ij}}{h_{j,m}^T u_{m,i}} - y_{ij} + h_{j,m}^T u_{m,i} \quad (3)$$

$$h_{j,m} \leftarrow h_{j,m} \frac{\sum_i y_{ij} u_{m,i} / x_{k,n}}{\sum_i u_{m,i}}, \quad u_{m,i} \leftarrow u_{m,i} \frac{\sum_i y_{ij} h_{j,m} / x_{k,n}}{\sum_i h_{j,m}} \quad (4)$$

NMFでは、ランダムな非負値で初期化した行列 $\mathbf{H}, \mathbf{U}$ に式(4)の更新式を収束するまで繰り返し、適用する。更新の収束で、行列分解結果の基底行列 $\mathbf{H}$ 、係数行列 $\mathbf{U}$ が得られる。

3. 3 SVRを用いた基底の評価

NMFでは、 $M < \min(K, N)$ のとき、観測行列 $\mathbf{Y}$ は、低ランク行列の積で近似することに相当することから、基底行列 $\mathbf{H}$ と係数行列 $\mathbf{U}$ の線形結合で表現できる。線形結合の例を式(5)に示す。

$$\mathbf{Y} \simeq \sum_{m=1}^M h_{j,m} \mathbf{u}_{m,i} = h_{j,1} \begin{pmatrix} u_{1,i} \\ u_{2,i} \\ \vdots \\ u_{M,i} \end{pmatrix} + \dots + h_{j,M} \begin{pmatrix} u_{1,i} \\ u_{2,i} \\ \vdots \\ u_{M,i} \end{pmatrix} \quad (5)$$

式(7)の基底行列 $\mathbf{H}$ の基底 m は、寄与率の高い特徴量がグルーピングされた情報である。また、係数行列 $\mathbf{U}$ の成分 $u_{m,i}$ の値は文書 i の基底 m への重みを表す。したがって、成分 $u_{m,i}$ の値を用いることで、基底 m の評価が行える。本研究では、NMFの基底の評価に、非線形回帰手法の一つであるサポートベクター回帰モデル(SVR: Support Vector Regression) [33]を用いる。SVRは、入力 $\mathbf{x}_i$ を特徴空間へ非線形写像し、特徴空間で線形回帰を行うモデルである。SVRの回帰関数を式(6)に示す。

$$f(\mathbf{x}) = \sum_{i=1}^l (\alpha_i + \alpha_i^*) k(\mathbf{x}_i, \mathbf{x}) + b \quad (6)$$

α_i, α_i^*, b は学習によって決まるパラメータである。 $k(\mathbf{x}_i, \mathbf{x})$ は入力 $\mathbf{x}_i$ を p 次元の特徴空間へ写像するカーネル関数である。本研究では、RBFカーネル[34]を用いる。RBFカーネルを式(7)に示す。

$$k(\mathbf{x}_i, \mathbf{x}) = \exp(\|\mathbf{x}_i - \mathbf{x}\|^2 / 2\sigma^2) \quad (7)$$

4 章で後述するが、本研究の SVR の目的変数には、プライバシー侵害のアンケート結果を用いた。また、説明変数 $\mathbf{x}_i$ には、NMF の基底 $\mathbf{H}$ に対する重みである係数行列 $\mathbf{U}$ の値を用いた。本研究の基底の評価では、各説明変数 $\mathbf{x}_i$ を順に除外した場合で、プライバシー侵害のアンケート結果に対する各予測誤差を算出した。予測誤差が大きい入力の組み合わせは、影響大きい説明変数が除外された組み合わせである。したがって、説明変数 $\mathbf{x}_i$ から目標変数であるプライバシー侵害を判別できる。予測誤差には、MAE (Mean Absolute Error) および RMSE (Root Mean Squared Error) を用いた。MAE を式(8)に RMSE を式(9)に示す。 n は予測対象のデータ数、 y_i は実績値、 $\hat{y}_i$ は予測値である。

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (8)$$

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (9)$$

上述の、3.2 節、3.3 節のとおり、提案手法は、テキスト情報から抽出した特徴量を既存アルゴリズムの NMF で行列分解(3.2 節)した分解結果を用いて、特性を表す基底を SVR で評価(3.3 節)する基底の評価手法である。NMF の基底の評価を行うことで、プライバシー侵害に影響が大きい特徴量を推定できる。

4. 実装と評価

4. 1 実験環境と評価方法

本研究の実装は、プログラム言語 Python¹⁰を用いた。単語分割、品詞判定は、MeCab¹¹を用いた。アルゴリズムの実装は、Gensim¹², scikit-learn^{13,14}を用いた。テキスト情報の前処理は、文字コードを UTF-8 に変換、空白記号・改行コードの除去を実施した。バイト列の欠損行は対象外とした。テキスト情報の等価な文字は Normalization Form KC (NFKC)¹⁵で正規化を実施した。また、形態素解析では、活用形を標準形に変換し、1 文字単語などを削除する前処理を実施した。本研究で話題の特徴量抽出に用いる LSI の次元数や LDA のトピック数では、任意で値の設定を行う。LSI の K の最適値は 300 から 500 の範囲が提案されている[35]。本研究では $K=300$ を設定した。LDA に関しても同様に $K=300$ を設定した。本研究の話題抽出の特徴量を表 3 に示す。

表 3 話題の特徴量の詳細

特徴量名	次元数	値の定義	文献
潜在意味解析(LSI)	300	文書の重み付き係数値	[18]
トピックモデル(LDA)	300	文書のトピック分布	[18]

提案手法で用いる NMF の観測行列 $\mathbf{Y}$ は、($j = 1, \dots, K$)から構成される N 行 K 列の長方形行列である。ここで、特徴量の次元数 K は、表層情報の次元数 22, アルゴリズムの次元数 600, 辞書の次元数 1, 451 の合計値 2, 073 である。また、特徴変換を行う際の基底数 M は 10 である。観測行列 $\mathbf{Y}$ 作成は、Pandas¹⁶および Numpy¹⁷を用いた。特徴量は 0 から 1 の間に正規化し、各変数の計測尺度の違いを考慮し、L1 正則化による制約補正を実施した。

提案手法で基底評価に用いる目標変数には、SNS 投稿記事を基に作成したプライバシー侵害のアンケートを用いた。目標変数の文書数である N は 96 である。アンケートでは、収集した Twitter の静止画付き投稿が、プライバシー侵害に該当しているかのアンケートを行い、アノテーションを実施した。アンケート実施人数は 55 名、

¹⁰ Welcome to Python.org : <https://www.python.org>

¹¹ MeCab : <http://taku910.github.io/mecab/>

¹² gensim : <https://radimrehurek.com/gensim/>

¹³ scikit-learn : <http://scikit-learn.org/stable/>

¹⁴ scikit-learn : <https://github.com/scikit-learn>

¹⁵ Unicode Technical Reports : <https://unicode.org/reports/>

¹⁶ Pandas : <http://pandas.pydata.org/>

¹⁷ NumPy : <http://www.numpy.org>

収集対象としたハッシュタグは「#写真, #集合写真, #隠し撮り, #盗撮, #無音カメラ, #バカ発見器」である。

提案手法の有用性の評価では、文書分類の評価および因果関係がある特徴量の推定を行う。提案手法の評価方法を図4に示す。(1)では、基底に基づく特徴量を用いて分類器による評価を行う。(2)では、ベイジアンネットワークを用いて目標変数に因果関係のある特徴量を推定する。

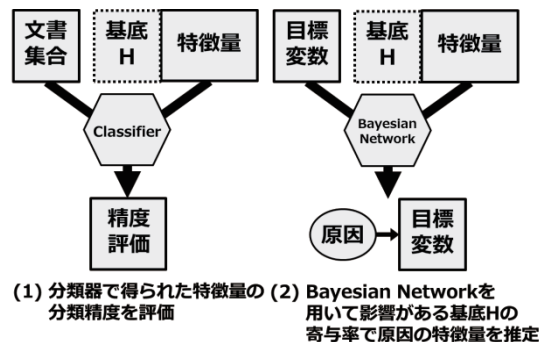


図4 基底の評価方法

4. 2 文書分類による提案手法の評価

本研究では、提案手法で得られた基底の評価を分類器で行う。NMFの基底は、寄与率の高い特徴量がグルーピングされた情報である。そこで、提案手法で得られたプライバシーの影響が大きい基底の寄与率上位の特徴量を用いてテキスト集合を分類し、提案手法の有用性を評価する。提案手法で基底選択に用いたアンケートと同様のハッシュタグのデータを抽出し、文書分類で評価を行う。評価に用いたデータは2008年から2012年までのTwitterのSNS投稿記事を基にしており、収集期間の投稿記事のデータ総数は約12億件である。評価に用いたデータを表4に示す。本研究では、収集期間の投稿記事より、ファイルの拡張子「.png」「.jpg」「.jpeg」「.gif」および「twimg」が含まれる投稿記事を抽出し、画像付き投稿記事とした。

表4 分類評価に用いたデータ

評価対象のSNS投稿記事	Twitterのツイートデータ
収集期間	2008年10月 - 2012年5月 (但し、一部期間は欠落あり)
収集期間の投稿記事の総数	1,198,426,092件(約12億件)
抽出した画像付き投稿記事数	5,249,748件(約520万件)
評価対象タグの投稿記事数	423件
評価対象タグ以外の投稿記事数	5,249,325件

分類評価では、アンケート評価と同様の評価対象タグの投稿記事を「プライバシー侵害に該当した投稿記事」とした。評価対象のタグは「#写真, #集合写真, #隠し撮り, #盗撮, #無音カメラ, #バカ発見器」である。評価対象タグ以外の画像付き投稿記事は「プライバシー侵害に該当していない投稿記事」としてタグ付けを行った。評価実験では、評価対象タグ以外の画像付き投稿記事から、4万件をランダムに抽出し、評価に用いた。

また、本研究では、分類器に「Ada Boost」「Random Forest」「MLP」「KNeighbors」を用いた。分類器による評価では、「適合率」「再現率」「F-measure」の評価指標を用いて評価する。適合率は不正解データを正解データと判定しないようにする評価指標である。再現率は、分類器がすべての正解データを判定する評価指標である。またF-measureは適合率と再現率の調和平均値である。評価実験では「適合率」「再現率」「F-measure」は層化5分割交差検証で算出した。

4. 3 因果関係のある特徴量の推定

提案手法で、プライバシー侵害に影響が大きい基底が得られる。したがって、基底に対する寄与率の高い特徴量が明らかになる。一方で、因果関係の推定の場合、基底への寄与率のみでは、プライバシー侵害の原因となる特徴量とは言えない。そこで本研究では、分類評価に加え、寄与率上位の特徴量を用いて、特徴量間およびプライバシー侵害の判断との因果関係を明らかにする。因果関係の評価には、ベイジアンネットワーク(Bayesian

Network) を用いる[36]. ベイジアンネットワークは、ベイズの定理(Bayes rule)に基づいたモデルである.

ベイジアンネットワークの例を図5に示す. 図5の例では X_1 が親ノードであり原因, X_2 や X_3 が子ノードであり結果である. X_2 や X_3 のように, 子ノードは, 他の子ノードである X_j の親ノードにもなる.

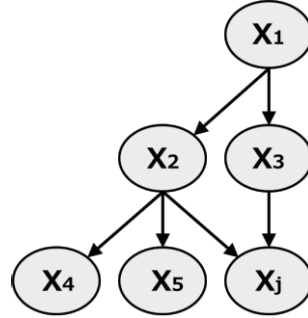


図5 Bayesian Network

ここで, 事象 $P(X_1)$ を原因, 事象 $P(X_j)$ が結果である場合, 原因によって結果が起きる確率は $P(X_j|X_1)$ で表される. ベイズの定理では, この $P(X_j|X_1)$ に加えて, $P(X_1)$, $P(X_j)$ が明らかな場合, 結果に対する原因の確率である $P(X_1|X_j)$ を推定できる. 原因の事象間の関係性を多数のノードに拡大したモデルがベイジアンネットワークである. ベイジアンネットワークは, 各変数は確率変数, ノード間の矢印は因果関係, 条件付き確率で定量化されている有向非循環グラフで表される. ベイジアンネットワークでは, 結果である子ノードを X_j とした際, 原因 $P_\alpha(X_j)$ は, 親ノードの集合 $(x_1^j, \dots, x_i^j)$ である. したがって, 依存関係は $P(X_j|P_\alpha(X_j))$ と表記できる. このため, すべての確率変数の同時確率分布は式(10)と表記できる.

$$P(X_1, \dots, X_n) = \prod_{j=1}^n P(X_j|P_\alpha(X_j)) \quad (10)$$

すべての確率変数の同時確率分布が明らかな場合, 確率変数の依存関係は各子ノードとその親ノードの間にリンクを張って, ベイジアンネットワークのグラフ構造で表される. 本研究の確率変数 $(X_1, \dots, X_n)$ には, 特性の表す基底に対する特徴量の寄与率の結果および, プライバシー侵害のアンケート結果を用いてベイジアンネットワークを構築した. したがって, 構築したグラフ構造から変数間の依存関係が明らかになり, プライバシー侵害のアンケート結果のノードから, 原因ノードである特徴量のノードを推定できる. このため, ベイジアンネットワークで, プライバシー侵害の原因の可能性がある特徴量が推定できる.

本研究のベイジアンネットワークのモデル構築では, 山登り法(Hill Climbing), グラフ構造の評価に用いる評価基準には, AIC(Akaike's Information Criterion)を用いて構築を行った. ノードは評価を行う NMF の基底の寄与率上位 100 個の特徴量を用いた. また, ベイジアンネットワークの変数値は 2 値の名義尺度に変換し, モデル構築を行った. ここで, プライバシー侵害のアンケート回答結果 U は, $(i = 1, \dots, N)$ から構成される. N は SNS 投稿記事数である. 結果ノードであるプライバシー侵害の評価ノード $X_j (X_{j,i}, \dots, X_{j,N})$ を式(11)に示す.

$$X_{j,i} = \begin{cases} \frac{\sum_{n=1}^N U_n}{N} \leq U_i, & X_{j,i} \in \text{Set } 1 \\ \frac{\sum_{n=1}^N U_n}{N} > U_i, & X_{j,i} \in \text{Set } 0 \end{cases} \quad (11)$$

式(11)の U_i はアンケートの結果, SNS 投稿記事である文書 i をプライバシー侵害であると判定した数である. プライバシー侵害のアンケート結果は, 人手によるプライバシー侵害の評価の基準が正規分布にしたがって生成されると仮定し, 平均値を基準とした. SNS 投稿記事がプライバシー侵害であると判定した数が平均値を超えた場合, 文書 i におけるプライバシー侵害の評価ノードの要素である $X_{j,i}$ は「プライバシーに該当」として 1 をタグ付けし, 平均値未満の場合は「プライバシーに該当しない」として 0 をタグ付けした. また, 本研究の評価では, N は 96 である. アンケート結果のタグ付けでは, アンケートに用いた SNS 投稿記事の 96 件中 46 件がプライバシー侵害であると評価された.

特徴量のノードのタグ付けは, SNS 投稿記事は内容が各記事で異なり, 特徴量の値は大きく異なる場合がある. したがって, 評価の基準値には平均値を用いず, 中央値を基準にタグ付けを行った. タグ付けは SNS 投稿記事である文書 i の特徴量 k の評価ノードである $X_{k,i}$ の値で行う. 評価ノード $X_{k,i}$ の値が特徴量 k の値の中央値を超

えた場合は「寄与率が高い」1 をタグ付けし中央値未満の場合は「寄与率が小さい」として0 をタグ付けした。

5. 評価実験

5. 1 基底の評価結果

提案手法の結果である基底 H を SVR で評価した結果を表5 および図6 に示す。結果、表5 および図6 より、最も MAE および RMSE の値が大きくなるのは、基底3 の説明変数を除いて SVR を行った場合である。したがって、最もプライバシー侵害に影響が大きい基底は基底3 である。同様に、基底2 および基底9 も影響が大きい基底であることが明らかになった。各基底を特徴量の種別ごとにヒートマップを表したものを図7 に示す。既存研究の特徴量は、メルヴィル・デューイの十進分類法[37]に基づき、表層情報(Surface Layer Information)、話題(TOPIC)、語種(Word Type)、基本語(Basic Vocabulary)、意味属性(Semantic Attributes)、言語表現(Verbal Expression)、文末表現(Sentence End Expression)、品詞(Pos Type)、固有表現(Unique Expression)、評価表現(Evaluation Expression)に分類した。図7 の色の濃さは、抽出した基底に対する種別の係数値が大きいことを示している。

表5 SVR でMAE およびRMSE を算出した結果

	除去 なし	基底の除去									
		0	1	2	3	4	5	6	7	8	9
MAE	0.8627	0.8602	0.8388	0.8775	0.9091	0.8506	0.8380	0.8395	0.8493	0.8204	0.8672
RMSE	1.1470	1.1343	1.0624	1.1613	1.1946	1.1196	1.1075	1.0719	1.1165	1.0683	1.1564

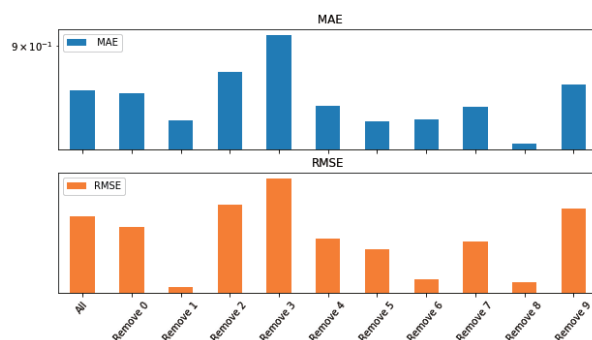


図6 基底Hの評価

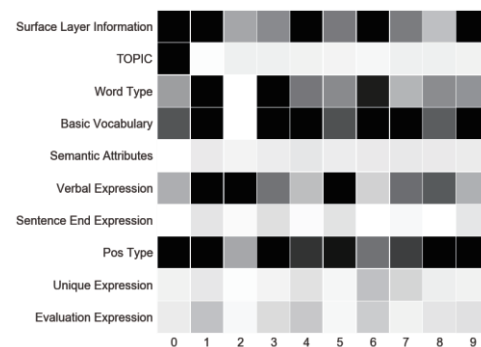


図7 基底Hのヒートマップ

図7 から、基底3 は語種、基本語、品詞の寄与率が高い。ここで、最も影響が大きい基底は基底3 である。提案手法の結果、得られた基底の特徴量の上位20 個を表6 に示す。

表6 基底3 の寄与率上位20 の特徴量

特徴量名	
意味分類体語彙表 意味分類コード(1. 12)	品詞特徴量(副詞)
分類項目一覧表 意味分類コード(1. 503)	評価値表現辞書
分類項目一覧表 意味分類コード(2. 503)	単語感情極性対応表(ネガティブ単語頻度)
分類項目一覧表 意味分類コード(3. 503)	機能表現タグ(B-否定)
分類項目一覧表 基本語(1)	文字種特徴量(頻度 - カタカナ)
TOPIC(LSI)	文字種特徴量(比率 - カタカナ)
日本語教育基本語彙 意味分類コード(1461)	機能表現タグ(B-順接限定)
意味分類体語彙表 意味分類コード(1. 461)	LPSW
語種判定(3)	機能表現タグ(I-否定)
日本語教育基本語彙 基本語(0)	意味分類体語彙表 基本語(*)
分類項目一覧表 意味分類コード(1. 503)	品詞特徴量(副詞)

寄与率による評価では、語種判定(3)、文字種特徴量(頻度 - カタカナ)、文字種特徴量(比率 - カタカナ)など外来語の特徴量が寄与率が高い結果であった。また、評価値表現辞書や単語感情極性対応表(ネガティブ単語頻度)など評価表現の特徴量も寄与率が高い結果となっている。加えて、従来法である文献[7]の特徴量であるLPSWなどが寄与率上位に評価されている。したがって、基底3はプライバシーに影響が高い基底であると解釈できる。

5.2 分類実験による特徴量の評価結果

5.1節で提案手法を適用し、得られたプライバシーの影響が大きい基底3の寄与率上位100個の特徴量を用いて分類器による基底の評価を実施した。分類器による分類結果を表7に示す。

表7 分類器による評価

	適合率	再現率	F-measure
Ada Boost	0.88	0.85	0.87
Random Forest	0.90	0.64	0.75
MLP	0.87	0.72	0.78
KNeighbors	0.89	0.68	0.77

結果、すべての分類器で「適合率」の評価基準で高い結果が得られた。最も「適合率」が高い分類器は「Random Forest」の「0.90」である。一方で、「再現率」では「Ada Boost」が「0.85」、ついで「MLP」では「0.72」が得られた。

結果、適合率と再現率の調和平均値である「F-measure」では、「Ada Boost」が最も高い「0.87」を示した。従来法である文献[7]のプライバシーに関する投稿記事の抽出評価では、「F-measure」で最も良い評価結果が「0.73」である。したがって、従来法と比較しても良い評価結果が得られた。このため、提案手法によって得られた基底の特徴量を用いることで、プライバシー侵害に影響があるSNS投稿記事の分類が行えることがデータで示された。

5.3 因果関係がある特徴量の評価結果

アンケート結果を目的変数として、ベイジアンネットワークを構築し、プライバシー侵害の変数ノードの原因ノードとなった特徴量を表8に示す。ベイジアンネットワークによる因果関係のある特徴量の推定では、5.1節で評価した基底3に加えて、重み付け係数値上位の基底である基底2、基底9の評価も行った。

表8 プライバシー侵害記事の評価結果の特徴量

No.	基底2	基底3	基底9
1	TOPIC(LSI)	機能表現タグ (B-不許可)	TOPIC(LSI)
2	文字種 (比率 - カタカナ)	意味分類コード(2) (1.564)	単語感情極性 (ポジティブ比率)
3	語種判定(3)	-	品詞特徴量(動詞)
4	固有名詞(一般)	-	-

提案手法で評価した基底3の特徴量の寄与率で因果関係を推定した際は、機能表現タグ (B-不許可)と、意味分類コード(2)の(1.564)がプライバシー侵害の原因として評価された。機能表現タグ (B-不許可)の辞書は「て、ちゃ」である。一方で、意味分類コード(2)の(1.564)は分類では「体の類・自然物および自然現象」である。しかし、詳細な分類では「魚(魚類・円口類)」に該当し、辞書の単語は「魚、たい、あじ」などである。したがって、ベイジアンネットワークの推定では基底3とアンケート結果を用いて推定された辞書の特徴量がプライバシー侵害と因果関係があるとは十分に言えない。

一方で、係数値上位の基底 2 の特徴量の寄与率で因果関係を推定した場合、カタカナの比率や語種判定(3)など、外来語の特徴量がプライバシー侵害の原因として評価された。推定された特徴量である語種は「ホテル・旅館・やどや」など、同じ意味であっても、イメージが異なる特徴量である[38]。推定結果である「語種判定(3)」は外来語である。既存研究の評価では、外来語は「女性・ファッション」の要因が相対的に高い[39]。しかし、肯定的なイメージを付加する反面、「異質性、斬新な響き、ステイタス顕示」など否定的な印象で評価される要因もある[40]。また、「プライバシー」という単語も外来語である。したがって、基底 2 で評価された外来語はプライバシー侵害に関連し、否定的な印象で評価される外来語があると推察される。また、固有名詞も因果関係がある結果となっていた。このため、特定の固有名詞がプライバシー侵害と関連していると推察される。

基底 9 では動詞の比率がプライバシー侵害の原因として評価された。文書に動詞が多い場合、「行動、変化を述べる」などの動きの描写がある動詞的な記述の文章であるとされている[41]。また、動詞は既存研究のコーパス間比較では、論文、新聞と比較し、日記が比率が高い[42]。したがって、行動、変化を述べる、日記のような個人の投稿がプライバシー侵害と関連していると推察される。基底 2、基底 9 では TOPIC (LSI) が原因として評価された。したがって、プライバシー侵害に該当する話題があると推察される。

ここで、提案手法の結果、得られた基底は基底 3 である。しかし、ベジアンネットワークの推定では、基底 3 では因果関係のある特徴量の推定は十分ではない。一方で、係数値上位の基底 2 および基底 9 では、因果関係があると推定される特徴量が得られた。結果、ベジアンネットワークで NMF の基底から因果関係のある特徴量を推定する場合は、単一の基底でなく、複数基底の評価が必要であると推察される。

6. まとめ

本研究では、SNS の投稿記事を対象に、プライバシー侵害に関わる特徴量を推定した。提案手法では、NMF、SVR を組み合わせ、数 1,000 次元に及ぶコンテンツの特徴量を次元圧縮し、プライバシー侵害に影響の大きい特徴量がグルーピングされた情報である基底を評価した。結果、得られた基底に寄与率の高い特徴量では、「意味分類コード」、「基本語」「語種」「カタカナ」「評価表現」「従来法で得られたプライバシー侵害の辞書の単語」などがプライバシー侵害に影響がある特徴量が得られた。

得られた結果を用いて、プライバシー侵害に影響がある SNS 投稿記事の分類実験を行い、「Ada Boost」で適合率「0.88」、再現率「0.85」、F-measure「0.87」の評価基準で高い結果が得られた。したがって、得られた特徴量は、プライバシー侵害に影響がある SNS 投稿記事の分類で有用性があることが示された。また、ベジアンネットワークによる因果関係のある特徴量の推定では、外来語や行動、変化を述べる、日記のような個人の投稿がプライバシー侵害と因果関係がある特徴量であると推定された。

今後の課題は、提案手法で特徴量変換を行う際の最適な基底数の決定や写真や動画などのメディアの特徴量を組み合わせた評価も行いたい。また、得られた結果をデジタル・デバイドによるプライバシー侵害投稿の発生防止に活用したい。

謝辞

本研究では、文献[7]および文献[23]で使用されたデータを用いた。ここに記してデータ提供して頂いた嶋田 茂先生に謝意を示します。

[参考文献]

- [1] 小蔵 正輝, テンポラルネットワーク上の伝播過程, 計測と制御, 公益社団法人 計測自動制御学会, 2016, 55, 11, 942-947
- [2] 高橋 正和, SNS の安全な歩き方 : セキュリティとプライバシーの課題と対策 (特集 情報セキュリティ技術の国際標準化動向の最新報告), 映像情報 medical , 産業開発機構, 2012-08, 44, 9, 737-742
- [3] 山口 真一, 実証分析による炎上の実態と炎上加担者属性の検証, 情報通信学会誌, 公益財団法人 情報通信学会, 2015, 33, 2, 53-65,
- [4] 大谷 卓史, プライバシーの多義性と文脈依存性をいかに取り扱うべきか : Nissenbaum の文脈的完全性と Solove のプラグマティズム的アプローチの検討, 吉備国際大学研究紀要 (人文・社会科学系), 2016-03-18, 26, 41-62
- [5] 堀部 政男 : プライバシー・個人情報保護の新課題, 商事法務, 2010
- [6] 佐久間 淳, プライバシー保護データマイニング(私のブックマーク), 人工知能学会誌, 人工知能学会, 2011-09-01, 26, 5, 533-536

- [7] 佐生 明陽, 輪島 幸治, 雨車 和憲, 田中 勇帆, 嶋田 茂, 小川 誠巳, 古川 利博, SNS 記事の語彙的結束性の分析によるプライバシー関連語の抽出精度の向上, DEIM2015 第 7 回データ工学と情報マネジメントに関するフォーラム 講演論文集, 日本データベース学会, 2015, pp.E2-1
- [8] 高崎 晴夫, パーソナライズド・サービスに対する消費者選好に関する研究, 情報通信学会誌, 公益財団法人情報通信学会, 2016, 34, 3, 25-39
- [9] 町田 史門, 梶山 朋子, 嶋田 茂, 越前 功, SNS におけるセンシティブデータの漏洩検知に基づく公開範囲の設定方式, 情報処理学会論文誌, 2014-09-15, 55, 9, 2092-2103
- [10] 高田 さとみ, 周 子胤, 高田 美樹, 大本 茂史, 岸本 拓也, 奈良 育英, 嶋田 茂, キャプションテキスト感情の分析によるプライバシー侵害シーン抽出, 研究報告自然言語処理 (NL), 一般社団法人情報処理学会, 2014-01-30, 2014, 6, 1-5,
- [11] 平井 智尚, なぜウェブで炎上が発生するのか: 日本のウェブ文化を手がかりとして, 情報通信学会誌, 公益財団法人 情報通信学会, 2012-03-25, 29, 4, 61-71
- [12] 山名 早人, 村田 剛志, 検索エンジン 2005 -Web の道しるべ: 1. 検索エンジンの概要, 情報処理, 一般社団法人情報処理学会, 2005-09-15, 46, 9, 981-987
- [13] 外資系コンサルのリサーチ技法, アクセンチュア 製造・流通本部 一般消費財業界グループ, 東洋経済新報社
- [14] 高田 さとみ, 小山 貴之, 町田 史門, 宋 洋, 嶋田 茂, SNS 画像投稿時に発生するプライバシー侵害の要因分析, 電子情報通信学会技術研究報告, 一般社団法人電子情報通信学会, 2012-08-20, 112, 189, 69-74
- [15] 大本 茂史, 岸本 拓也, 高田 美樹, 高田 さとみ, 奈良 育英, 周 子胤, 嶋田 茂, ウェアラブルカメラによる SNS 記事投稿時に発生するプライバシー侵害の特徴分析, DEIM2014 第 6 回データ工学と情報マネジメントに関するフォーラム 講演論文集, 日本データベース学会, 2014, pp.B2-1
- [16] 家辺 勝文, インターネットの技術基盤の国際化と日本語文書処理 (特集 デジタル時代の日本語), 情報の科学と技術, 0913-3801, 社団法人情報科学技術協会, 2014, 64, 11, 450-455
- [17] 福田 誠, 吉田 武尚, 吉田 誠, 樫岡 健史, 中学校技術・家庭科教科書 電気領域の表記・表現について, 日本教科教育学会誌, 日本教科教育学会, 2000, 23, 3, 37-42
- [18] 岩田 具治: トピックモデル (機械学習プロフェッショナルシリーズ), 講談社, 2015
- [19] 高橋 徹, 社会システム分化とゼマンティック: ルーマンにおける社会変動論の一視角, 社会学評論, 日本社会学会, 1999-03-30, 49, 4, 620-634
- [20] 松吉 俊, 江口 萌, 佐尾 ちとせ, 村上 浩司, 乾 健太郎, 松本 裕治, テキスト情報分析のための判断情報アノテーション, 電子情報通信学会論文誌 D, 情報・システム, 人電子情報通信学会, 2010-06-01, 93, 6, 705-713
- [21] 福田 一雄, 日本語モダリティ覚え書き(その 1), 言語の普遍性と個別性, 新潟大学大学院現代社会文化研究科「言語の普遍性と個別性」プロジェクト, 2014-03, 5, 1-13
- [22] 西原 陽子, 松村 真宏, 谷内田 正彦, Q&A コミュニティでの質疑応答パターンの理解, 人工知能学会全国大会論文集, 人工知能学会, 2008, 22, 1-4
- [23] 高田 さとみ, 小山 貴之, 町田 史門, 宋 洋, 嶋田 茂, SNS 画像投稿時に発生するプライバシー侵害の要因分析, 電子情報通信学会技術研究報告, 一般社団法人電子情報通信学会, 2012-08-20, 112, 190, 69-74
- [24] 小林 のぞみ, 乾 健太郎, 松本 裕治, 立石 健二, 福島 俊一, 意見抽出のための評価表現の収集, 自然言語処理, 言語処理学会, 2005, 12, 3, 203-222
- [25] 東山 昌彦, 乾 健太郎, 松本 裕治, 述語の選択選好性に着目した名詞評価極性の獲得, 言語処理学会第 14 回年次大会論文集, 2008, 584-587
- [26] Hiroya Takamura, Takashi Inui, and Manabu Okumura. 2005. Extracting semantic orientations of words using spin model. In Proceedings of the 43rd Annual Meeting on Association for Computational Linguistics (ACL '05). Association for Computational Linguistics, Stroudsburg, PA, USA, 133-140.
- [27] 乾 孝司, 奥村 学, テキストを対象とした評価情報の分析に関する研究動向, 自然言語処理, 言語処理学会, 2006-07-10, 13, 3, 201-241
- [28] 小林 のぞみ, 乾 健太郎, 松本 裕治, 意見情報の抽出/構造化のタスク仕様に関する考察, 情報処理学会研究報告自然言語処理 (NL), 情報処理学会, 2006-01-13, 2006, 1, 111-118

- [29] 町田 史門, 小山 貴之, 宋 洋, 高田 さとみ, 嶋田 茂, 越前 功, SNS 写真投稿に起因するプライバシー侵害の類型化とその保護策, 電子情報通信学会技術研究報告, 一般社団法人電子情報通信学会, 2012-09-27, 112, 226, 5-10
- [30] 亀岡 弘和, 非負値行列因子分解, 計測と制御, 公益社団法人 計測自動制御学会, 2012-09-10, 51, 9, 835-844
- [31] CSISZAR, I. , I-divergence geometry of probability distributions and minimization problems, The Annals of Probability, 1975, 3, 1, 146-158
- [32] Fevotte, Cedric, and Jerome Idier. "Algorithms for nonnegative matrix factorization with the β -divergence." Neural computation 23.9 (2011): 2421-2456.
- [33] Alex J. Smola and Bernhard Scholkopf. 2004. A tutorial on support vector regression. Statistics and Computing 14, 3 (August 2004), 199-222.
- [34] 小林 正幸, 小西 康夫, 藤田 貞雄, 石垣 博行, サポートベクタ回帰モデルを用いた超音波モータの位置決め制御, 精密工学会誌論文集, 公益社団法人精密工学会, 2006-05-05, 72, 5, 596-601
- [35] Roger B. Bradford. 2008. An empirical study of required dimensionality for large-scale latent semantic indexing applications. In Proceedings of the 17th ACM conference on Information and knowledge management (CIKM '08). ACM, New York, NY, USA, 153-162.
- [36] 本村 陽一, ベイジアンネットによる確率的推論技術, 計測と制御, 公益社団法人 計測自動制御学会, 2003, 42, 8, 649-654
- [37] 光富 健一, デューイ十進分類法(DDC) (特集 分類について考える), 情報の科学と技術, 一般社団法人 情報科学技術協会, 1989, 39, 11, 478-483
- [38] 菊地 悟, 大学生の外来語意識(1) : イメージ・表記・語種意識の調査から, 岩手大学教育学部附属教育実践研究指導センター研究紀要, 岩手大学, 1994, 4, 61-73
- [39] 山崎 誠; 小沼悦. 現代雑誌における語種構成. 言語処理学会第 10 回年次大会発表論文集, 2004, 670-673.
- [40] 堀切 友紀子, 外来語に関する研究動向 : 使用意識と言語接触の視点から, お茶の水女子大学人文科学研究, お茶の水女子大学, 2013-03, 9, 113-124
- [41] 中尾 桂子, 品詞構成率に基づくテキスト分析の可能性 : メール自己紹介文、小説、作文、名大コーパスの比較から, 大妻女子大学紀要 文系, 大妻女子大学, 2010-03, 42, 128-101
- [42] 石田 栄美, 安形 輝, 野末 道子, 文体からみた学術的文献の特徴分析, 三田図書館・情報学会研究大会発表論文集, 三田図書館・情報学会, 2004, 33-36

(2018年7月15日受理)

(2018年8月17日修正)