

読み手の個性を考慮したニュース記事の感情応答モデルの構築

森山皓未^{1,†} 新 恭兵^{1,†} 小山 聡^{2,†} 栗原正仁^{2,†}

† 北海道大学大学院情報科学研究科 〒060-0814 札幌市北区北14条西9丁目

E-mail: ¹{a.moriyama_ast, atarashi_k}@complex.ist.hokudai.ac.jp,

²{oyama, kurihara}@ist.hokudai.ac.jp

あらまし 近年ではネット上のサービスが多様化していることから、様々な形で情報発信をすることが増えている。その中でも文書は一般的な媒体であり、この解析のために文書の感情推定に関するニーズも高まっている。しかしながら、文書の感情推定における既存の手法の多くは、各ユーザーの文書に対する反応の違いを考慮していないという欠点がある。この問題に対して、本研究では文書とユーザーが持つ2つの潜在的な特徴を組み合わせることでユーザー毎に異なる感情の分類器を作成し、文書に対する感情応答を推測するモデルを構築する。提案するモデルはトピックモデルを基盤としており、文書の潜在特徴としてはトピックの分布、ユーザーの潜在特徴としては実際にユーザーが付けたラベルを利用して学習を行う。またクラウドソーシングによってニュース記事に感情ラベルを付与したデータセットを用いて、潜在特徴の解釈性の確認や各ユーザーの傾向を考慮した応答が得られているかの実験を行う。

キーワード 感情推定, 文書処理, トピックモデル, クラウドソーシング

1 はじめに

近年ではネット上のサービスが多様化していることから、様々な形で情報発信をすることが増えている。その中でも文書は一般的な媒体であり、ニュース記事や広告、SNSへの投稿など、様々な情報が文書で発信されている。

しかし、文書で発信された情報が、いわゆる炎上といった予想外の読者の反応を引き起こし、問題となることがある。このようなリスクを避けたり、書き手の意図した効果を読者にもたすために、文書の感情推定に関するニーズも高まっている。

これまでの感情手法は一つの単語・文書に対してただ一通りの出力を持つことが多い。その一方で同じ出来事であっても内容によって、それぞれの受け手ごとに感情は異なる場合が多い。そのため文書を用いた感情推定においてもそれぞれ読み手ごとに異なる感情の傾向を考慮できるモデルを用いる必要がある。これにより、読み手に応じて良い感情をもたらす文書を提示したり、逆に不快な感情をもたらす恐れのある文書の提示を避けたりすることが可能となると期待できる。また、ある文書を特定の人々の集団に提示した場合、何割ぐらいの人がどのような感情を持つか、感情の分布を推定することも可能になる可能性がある。

そこで本研究では文書と読み手が持つ潜在的な特徴を組み合わせることによって、読み手毎に異なる感情の分類器を作成するモデルを考案した。また本論文では文書の一つとして、ニュース記事に関するユーザーの感情応答を推測するモデルを構築した。

本論文で使用する記号を表1に示す。

2 関連研究

文書における感情推定の既存研究の一つとして Ptaszynski ら [1] が用いた、感情語の辞書を利用した感情推定がある。[1] においては2つの手法が述べられており、一つは ML-ask データベース [2] のように単純に辞書内の感情語と単語や表現を比較する手法がある。

またもう一つは上位語や下位語、同族語というような言語が持つ階層構造の特性を利用することによって、単語がどの感情に含まれるか分類を行う手法が述べられている。無償利用可能な大規模シソーラスとして、英語については Wordnet [3]、日本語については日本語 Wordnet [4] などがある。

3 トピックモデル

トピックモデルとは1つの文書が複数のトピックによって構成されていると仮定する考え方である。一例として「羽生善治 竜王が史上初の永世7冠を獲得して国民栄誉賞を受賞した」という内容の文書を考えたとき、この文書は「将棋」や「政治」など、複数のトピックを持つと考えられる。

トピックモデルでは、まず文書集合 D^* のある文書 d は K 次元のトピックの分布 θ_d を保持しており、その分布に従って n 番目の単語 w_{dn} にトピック z_{dn} が割り当てられる。また一方で文書全体の単語の分布から各トピックにおける単語分布 ϕ_k が求められる。トピックモデルにおいて文書中の単語は、単語ごとに割り当てられたトピックの単語分布 $\phi_{z_{dn}}$ に基づいて生成されている。このときの文書 $\mathbf{w}_d$ の生成確率は式 (1) で示される。

$$p(\mathbf{w}_d|\theta_d, \Phi) = \prod_{n=1}^{N_d} \sum_{k=1}^K p(k|\theta_d) p(w_{dn}|\phi_k) \quad (1)$$

そのグラフィカルモデルは図 1，生成過程は以下のように表される．

- (1) For $k = 1, \dots, K$
 - (a) 単語分布の生成 $\phi_k \sim \text{Dirichret}(\beta)$
- (2) For $d = 1, \dots, D$
 - (a) トピック分布の生成 $\theta_k \sim \text{Dirichret}(\alpha)$
 - (b) For $n = 1, \dots, N_d$
 - i. トピックの生成 $z_{dn} \sim \text{Categorical}(\theta_d)$
 - ii. 単語の生成 $w_{dn} \sim \text{Categorical}(\phi_{z_{dn}})$

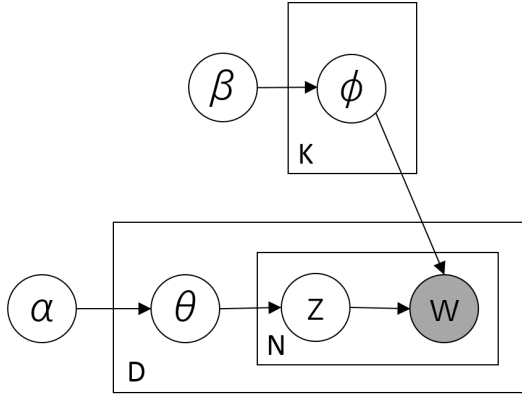


図 1 トピックモデルのグラフィカルモデル

トピックモデルの推定を行う手法のうち，最もよく知られるものの一つとして LDA (Latent Dirichlet Allocation) [7] というモデルがある．LDA とはトピック分布の事前分布としてディリクレ分布を仮定し，ベイズ推定を用いて推定を行う確率的生成モデルである．以降，本論文においてはトピックモデルの推定手法として LDA を用いることとする．

記号	説明
D	全文書の数
N_d	文書中の全単語数
K	トピック数
θ_d	文書 d におけるトピック分布
z_{dn}	文書 d における n 番目の単語のトピック
w_{dn}	文書 d における n 番目の単語
α, β	ディリクレ事前分布のハイパーパラメータ
ϕ_k	トピック k における単語分布
H	全ユーザーの数
E	全感情の数
ψ_{keh}	トピック k における読み手 h の感情 e の分布
l_{deh}	文書 d における読み手 h の感情 e のラベル
γ	ベータ事前分布のハイパーパラメータ

表 1 本論文における記号

4 提案モデル

提案モデルは前章で触れたトピックモデル (LDA) を基盤として読み手の潜在的な特徴を求める部分について拡張したモデルである．本論文では文書の潜在的な特徴の一例としてニュース記事のトピックの分布と，読み手の各記事に対する感情に関する潜在的な特徴の 2 つを組み合わせることによって，ニュース記事に対する感情について読み手ごとに異なる分類器を得るモデルを作成する．

4.1 グラフィカルモデルと生成過程

提案モデルのグラフィカルモデルは図 2 で表され，その生成過程は以下ようになる．

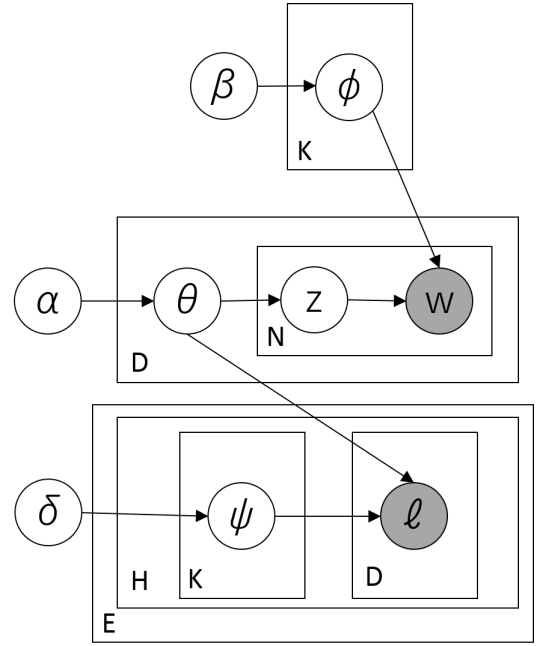


図 2 提案モデルのグラフィカルモデル

- (1) For $k = 1, \dots, K$
 - (a) 単語分布の生成 $\phi_k \sim \text{Dirichret}(\beta)$
 - (b) For $e = 1, \dots, E$
 - i. For $h = 1, \dots, H$
 - A. ユーザーの各感情分布の生成
 $\psi_{keh} \sim \text{Beta}(\gamma)$
- (2) For $d = 1, \dots, D$
 - (a) トピック分布の生成 $\theta_k \sim \text{Dirichret}(\alpha)$
 - (b) For $n = 1, \dots, N_d$
 - i. トピックの生成 $z_{dn} \sim \text{Categorical}(\theta_d)$
 - ii. 単語の生成 $w_{dn} \sim \text{Categorical}(\phi_{z_{dn}})$
 - (c) For $e = 1, \dots, E$
 - i. For $h = 1, \dots, H$
 - A. ユーザーの各感情ラベルの生成
 $l_{deh} \sim \text{Bernoulli}(\theta_d \cdot \Psi_{eh})$

4.2 実装

提案モデルは確率的生成モデルであるためサンプリング法や変分ベイズ推定を利用して学習を行う。本研究では[8]を元に pymc3 による Automatic Differentiation Variational Inference (ADVI) と呼ばれる手法を用いて実装した。ADVI とは、潜在変数の事後分布の近似分布をガウス分布と仮定したのちに変数変換と確率的最適化を行うことによって、変分ベイズ推定を自動で行う手法である。これを用いることで近似分布の学習においてモデル固有のパラメータ更新式の導出を行わなくても推定できるという利点がある。

読み手の集合 H^* のうちある読み手 h が文書 d に付与したラベルを l_{dh} とし、トピック k における読み手 h の感情分布を ψ_{kh} とする。このとき文書集合 D^* と読み手の感情ラベル集合 L^* に関する事後分布の関係は式 (2) で表される。

$$p(\Theta, \Psi | D^*, L^*) \propto p(D^* | \Theta) \cdot p(L^* | \Theta, \Psi) \quad (2)$$

両辺について対数を取ることで右辺は和の形を取ることが可能であり、右辺第一項については前節で述べたトピックモデルの式 (1) と同様である。

また右辺第二項について、感情 e にラベルを付与した読み手の集合を $\overline{H_e^*}$ としたとき、式 (3) で表すことができ、これを用いることによって文書集合と読み手の感情ラベル集合の事後分布が得られる。

$$p(l_d | \theta_d, \Psi) = \prod_{e=1}^E \prod_{h=1}^{\overline{H_e^*}} \sum_{k=1}^K p(k | \theta_d) p(l_{deh} | \psi_{keh}) \quad (3)$$

4.3 類似手法との比較

提案モデルでは、複数の読み手の複数の文書に対するラベル付けの結果から、文書と読み手に関する特徴を獲得することで、未知の文書に関して読み手の反応を予測することができる。一方で、複数のアノテーターが複数のデータ・タスクに与えたラベルの集合から、アノテーターの能力やデータが持つとされる真のラベルを推定する方法に関する研究は盛んに行われている。しばしばこれはクラウドソーシングのワーカーのラベル付け結果の統合に用いられるが、本研究で扱う読み手の反応の予測にも用いることもできる。本節ではこれらの既存研究で提案されたモデルについて、提案モデルとの比較も含めて述べる。

まず一つ目の先行研究として Dawid と Skene らの方法[10] (今後本モデルを DS とする) がある。この手法ではワーカーがラベルを付与する能力とそのラベルの正確さについて依存関係があるという仮定を置いている。より具体的には、各ワーカーは各々混合行列を持ち、その混合行列に従ってラベルを付与していると仮定している。この仮定を元にワーカーのパラメータと潜在変数である真のラベルについて Expectation-Maximization (EM) アルゴリズムを用いて推定を行うことで、各ワーカーについてその混合行列を獲得し、ワーカーが実際にラベルを付与していないタスクについても推定を行う手法であ

る。提案モデルとの違いとしては、ワーカーのパラメータと付与したラベルのみを用いて推定を行うことからタスクの中身が考慮されないことと、タスクに一つもラベルが付与されていないような状況だと真のラベルが推定できないため評価ができないという点が挙げられる。

次に梶野らの方法[11] (Personal Classifier (PC モデル)) について述べる。梶野らは、各ワーカーは各々分類器 (Personal Classifier) を持ち、タスクが与えられると、与えられたタスクを表現する特徴ベクトルを自身の持つ分類器に投入し、その出力に基づいてラベルを付けていると仮定している。

ここで梶野らの手法の特徴的な点として、各ワーカーのパラメータ w_h について、ベースとなるパラメータ w_b と各ワーカーの持つ能力や特性 v_h を用いて $w_h = w_b + v_h$ で表される点とパラメータの最適化問題が凸最適化になる点の2つがある。前者の利点としてはあるワーカーがパラメータの学習を行う際に、他のワーカーがラベル付けした結果を利用することが可能になる点が挙げられる。後者の利点としては大域的な最適値が得られるため、EM アルゴリズムを用いたときのように出力が初期値に大きく左右されない利点がある。本手法でパラメータの推定を行う際にはベースのパラメータ w_b とワーカーのパラメータ集合 W についてそれぞれ片方を固定して最適化するアルゴリズムを行う。提案モデルとの違いとして、潜在的特徴を特徴ベクトルから学習するわけではないため、ラベルが付与されていない文書を学習に用いることができない。

またこれらの手法に加えて、PC モデルにおいてベースとなる分類器の存在を仮定せずに、各ワーカーについて独立にロジスティック回帰を学習をさせるという方法をベースラインとして考えることができる。この方法を Independent PC (IPC) モデルとして、表 2 で提案モデルとの差異についてまとめた。

手法	ワーカー毎の評価	タスクの考慮	ラベル無し文書の学習への利用
DS モデル	○	×	×
PC モデル	○	○	×
IPC モデル	×	○	×
提案モデル	○	○	○

表 2 提案モデルと類似手法の比較

5 実験・結果・考察

5.1 データセットについて

使用する文書については[12]の livedoor ニュースコーパスより一部を抜粋して使用した。また文書に対する感情ラベルについては、クラウドソーシングサービスの一つである Lancers を活用することで、文書に対して怒り・悲しみ・喜び・不快・驚き・恐怖の6つの感情についてマルチラベルでラベル付けを行ってもらった。なお各読み手は記事に対して適した感情ラベルが存在しない場合においても、最も近い感情1つにラベル付けを行うこととした。

実際に今回収集したデータに関して、先述の文書コーパスよ

り 220 個の記事を抜粋し、それぞれの記事について 30 回ずつラベル付けを行った。なお各読み手が付けたラベルの数は一定ではないが、記事 10 個のラベル付けを 1 工程として、最低 1 工程読み手にラベル付けを行ってもらった。その結果収集されたラベルは、読み手の総数が 95、読み手の付けたラベルの数の平均値・中央値・最頻値は、それぞれ順に 69.47、30、10 となった。

感情ラベルをつけた場合を正例、つけなかった場合を負例としたとき、今回得られた各感情ラベルの正例と負例の数は表 3 のようになった。

	負例	正例
怒り	1738	260
悲しみ	1742	256
喜び	1861	137
嫌悪	1248	750
驚き	1067	931
恐怖	1836	162

表 3 文書に付与された感情ラベル数

また提案モデルにおいてはラベルのない文書を学習に利用できるという特徴があるため、上述のラベル付けを行った文書の他にラベル無し約 550 文書 (以後、UL 文書とする) を準備した。

5.2 実験 1

本実験では提案モデルを介して得られた文書と読み手の 2 つの潜在特徴を用いて、推定した読み手のトピックの分布が実際に解釈可能なものであるかを確認を行った。つまり、「ある読み手があるトピックのときに特定の感情を抱きやすいか」ということについて確認した。ここで解釈性についての例として、PC モデルでは文書中に出現した単語に対する反応が得られるため、単語レベルでの解釈性があると考えられる。

今回は文書と読み手の 2 つの潜在特徴がどの割合でトピックに反応しているかについて、ある読み手が「驚き」のラベルを付与した 10 文書の平均を用いて比較を行った。図 3 がその結果である。

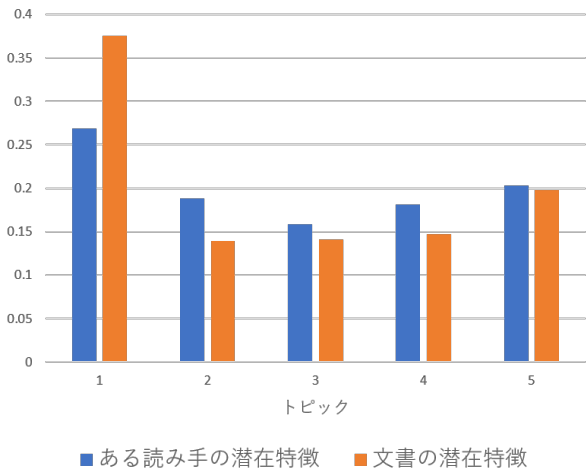


図 3 ある読み手と文書の「驚き」の潜在特徴分布

5.3 実験 1 の結果と考察

実験 1 の結果について、グラフから読み手と文書が一番反応しているトピックは共に 1 であることが見て取れる。したがって、ある読み手について各潜在特徴と感情が対応していると考えられる。このときトピック 1 の上位単語は「韓国・中国・フジテレビ・日本」となっており、「外国」のトピックが抽出されていると考えられる。これらのことから、提案モデルでは少なくとも部分的に、トピックレベルでの感情の解釈性を持った潜在特徴が得られていると考えられる。

5.4 実験 2

本節ではテストデータとして誰も評価を行っていない未知の文書を利用して、それに対する感情の推定結果の評価と比較を行った。5.1 節で述べた通り提案モデルではラベル無しデータを学習に使用することが可能である。そのため、テストデータの「文書」を学習に利用する場合としない場合の 2 つで評価を行った。また同様に UL 文書についても使用する場合としない場合について別途評価を行った。

今回比較する手法は、提案モデル・PC モデル・IPC モデルの 3 つである。なお DS モデルについては 4.3 節で述べたように誰も評価を行っていない文書についてのラベル推定ができないため、本実験においては比較を行うことができない。

今回のデータセットでは多くの感情においてラベルの正例と負例がインバランスであったことから、各感情ごとの ROC-AUC 値を評価値として扱うこととした。表 4 がその結果である。

	怒り	悲しみ	喜び	嫌悪	驚き	恐怖	平均
PC	0.569	0.539	0.635	0.641	0.563	0.599	0.591
IPC	0.556	0.578	0.565	0.604	0.584	0.554	0.573
提案 (文書有)	0.655	0.669	0.564	0.581	0.669	0.665	0.634
提案UL (文書有)	0.644	0.688	0.602	0.577	0.634	0.676	0.637
提案 (文書無)	0.650	0.685	0.570	0.597	0.675	0.661	0.640
提案UL (文書無)	0.613	0.687	0.559	0.567	0.643	0.683	0.625

表 4 各手法での ROC-AUC 値の比較

5.5 実験 2 の結果と考察

実験 2 の結果から、提案モデルの ROC-AUC 値は複数の感情において既存手法の結果を上回っており、また平均についても上回っていることが見て取れる。このことから提案モデルは誰にもラベル付けされていない未知のデータに関する感情推定について既存手法よりも優れていると考えられる。

また他にラベル無しデータを学習に使用した場合とテストデータの文書を学習に使用した場合について、それぞれ評価値に大きく差がないことが見て取れる。これについて、ラベル無しデータに対して学習データが十分に用意されていることが考えられる。このような場合においては学習データから得られない特徴量をラベル無しデータから補完する必要がなくなること

が考えられ、よりラベル無しデータを増やした環境において再度実験を行い、比較する必要があると考えられる。

5.6 全体を通した考察

データセットについて、クラウドソーシングを利用して文書にラベル付けを行ってもらったが、読み手1人あたりのラベル付与数の最頻値は10であり、非常に少ないことが見て取れる。このような場合だと、提案モデルにおいて、読み手の潜在的特徴は記事に対するラベルと記事のトピック分布を用いて学習されることから、ラベル付与数の少ない読み手に関して十分に読み手の分類器の学習ができないことが考えられる。そのため、梶野ら[13]のPCモデルをクラスタ化したClustered PCモデルへの拡張のように、読み手の潜在特徴についてDuanら[14]の手法などを用いて似通ったラベルの付与を行った読み手をクラスタ化することによって、一読み手クラスタあたりのラベル付与数を増加することなどによる工夫が考えられる。

5.7 ま と め

本研究では文書と読み手が持つそれぞれの潜在的な特徴を組み合わせることによって、文書に対する読み手の傾向を考慮した感情応答が得られるモデルを作成した。提案モデルは潜在特徴の解釈性が高く、未知のデータに対しての推定が優れているが、ラベル無しデータの大規模化を行った際の考察や、読み手のクラスタ化を用いた場合などの潜在特徴の工夫によってさらにモデルが改良できると考えられる。

謝 辞

本研究の一部はJSPS科研費JP15H02782およびJP18H03337の支援を受けた。

文 献

- [1] Michal Ptaszynski, Hiroaki Dokoshi, Satoshi Oyama, Rafal Rzepka, Masahito Kurihara, Kenji Araki, and Yoshio Momouchi. Affect analysis in context of characters in narratives. *Expert Systems with Applications*, Vol. 40, No. 1, pp. 168–176, 2013.
- [2] Michal Ptaszynski, Pawel Dybala, Rafal Rzepka, Kenji Araki, and Fumito Masui. MI-ask: Open source affect analysis software for textual input in Japanese. <https://openresearchsoftware.metajnl.com/articles/10.5334/jors.149/>.
- [3] Wordnet | a lexical database for English. <https://wordnet.princeton.edu/>.
- [4] 日本語 Wordnet - Nanyang Technological University. <http://compling.hss.ntu.edu.sg/wnja/>.
- [5] Peter Welinder, Steve Branson, Pietro Perona, and Serge J. Belongie. The multidimensional wisdom of crowds. In *Proceedings of the 23rd Annual Conference on Neural Information Processing Systems (NIPS 2010)*, pp. 2424–2432, 2010.
- [6] 岩田具治. 機械学習プロフェッショナルシリーズ トピックモデル. 講談社, 2015.
- [7] David M Blei, Andrew Y Ng, and Michael I Jordan. Latent dirichlet allocation. *Journal of Machine Learning Research*, Vol. 3, No. Jan, pp. 993–1022, 2003.
- [8] The PyMC Development Team. Automatic autoencoding variational Bayes for latent Dirichlet allocation with PyMC3 - PYMC3 3.6 documentation. <https://docs.pymc.io/notebooks/lda-advi-aevb.html>.
- [9] Alp Kucukelbir, Rajesh Ranganath, Andrew Gelman, and David Blei. Automatic variational inference in Stan. In *Proceedings of the 28th Annual Conference on Neural Information Processing Systems (NIPS 2015)*, pp. 568–576, 2015.
- [10] A. P. Dawid and A. M. Skene. Maximum likelihood estimation of observer error-rates using the em algorithm. *Journal of the Royal Statistical Society. Series C (Applied Statistics)*, Vol. 28, No. 1, pp. 20–28, 1979.
- [11] Hiroshi Kajino, Yuta Tsuboi, and Hisashi Kashima. A convex formulation for learning from crowds. In *Proceedings of the 26th AAAI Conference on Artificial Intelligence (AAAI 2012)*, pp. 73–79, 2012.
- [12] Ltd. RONDHUIT Co. ダウンロード - 株式会社ロンウイット. <https://www.rondhuit.com/download.html#1dcc>.
- [13] Hiroshi Kajino, Yuta Tsuboi, and Hisashi Kashima. Clustering crowds. In *Proceedings of the 27th AAAI Conference on Artificial Intelligence (AAAI 2013)*, pp. 1120–1127, 2013.
- [14] Lei Duan, Satoshi Oyama, Haruhiko Sato, and Masahito Kurihara. Separate or joint? estimation of multiple labels from crowdsourced annotations. *Expert Systems with Applications*, Vol. 41, No. 13, pp. 5723–5732, 2014.