

潜在的ディリクレ配分法を用いたネガティブ要因分析

輪島 幸治[†] 小河 誠巳[‡] 古川 利博[‡] 嶋田 茂[†]

[†]産業技術大学院大学 〒140-0011 東京都品川区東大井 1-10-40

[‡]東京理科大学 〒162-8601 東京都新宿区神楽坂 1-3

E-mail: [†] wajimax0513@gmail.com

あらまし 企業には、問い合わせとして製品の使い方など様々な質問が寄せられる。その質問の中には、製品に対するネガティブな意見やクレームも少なくない。質問の内容がネガティブなほど、顧客離れを招く可能性があり、早期の対応が求められる。現状、質問のネガティブさの判断と対応の優先付けは対応者の経験則によって行われている。それはネガティブさの度合いが一意的でなく、質問ごとにあいまいであり、単語のような表層的な情報だけでは判別が困難であることに起因している。そこで本研究では、質問内に潜在的に含まれているネガティブさの因子特定を目的とする。因子を特定することで、“質問のネガティブさ”が定量的に計測でき、適切な対応の優先付けが行える。今回は、感情極性による極性値ごとにレベル分けした文書から、潜在的ディリクレ配分法を用いた潜在トピックを推定し、トピックの分布について評価実験を実施した。

キーワード トピックモデル, 感情極性分類, 顧客関係管理, 潜在的ディリクレ配分法, テキストマイニング

Specific Negative Factors Using Latent Dirichlet Allocation

Koji WAJIMA[†] Tomomi OGAWA[‡] Toshihiro FURUKAWA[‡] Shigeru SHIMADA[†]

[†] School of Industrial Technology, Advanced Institute of Industrial Technology

10-40, Higashi-Ooi, Shinagawa-ku, Tokyo, 140-0011 Japan

[‡] Faculty of Engineering, Tokyo University of Science 1-3 Kagurazaka, Shinjuku-ku, Tokyo, 162-8601 Japan

E-mail: [†] wajimax0513@gmail.com

Abstract Customers ask various questions to the company. Negative questions are included in the question. If the questions are more negative, the company respond it more quickly. A decision of that the questions are negative or not negative is judged by as a general rule of thumb. The problem is that decision is difficult to determine from the superficial text. And Negative factors are fuzzy factors each the questions. This paper aiming at specifying Negative fuzzy factors. This paper specify Negative fuzzy factors using Latent Dirichlet Allocation and Sentiment Analysis. The computational simulation shows the efficiency of the proposed method.

Keyword : Topic Model, Semantic Analytics, Customer Relationship Management, Latent Dirichlet Allocation, Text Mining

1. はじめに

1.1. 背景

顧客関係管理は最も重要な企業活動の一つである。顧客との窓口となる、ヘルプデスクが受け取る顧客からの質問（以下、質問文書）は、製品に関する使用方法に関する質問の他、製品に対する改善要望やクレームなど、ネガティブな内容も少なくない。

だが現状、質問文書のネガティブさの判断と対応の優先付けは対応者の経験則によって行われている。それは“ネガティブ”の定義が一意的でなく、質問ごとにあいまいであり、単語のような表層的な情報では判別が困難であることに起因しているためである。

そこで本研究では、この潜在的に存在するネガティブな因子を以後、深刻度と定義する。深刻度が明らかになれば、ヘルプデスクをはじめとした顧客関係管理の分野においても対応者に依存せずとも、適切な質問文書の対応を行うことが実現可能となる。

1.2. 本稿の構成

本稿ではまず関連研究について述べる。その後、本研究の目的及び提案手法について述べる。そして計算機シミュレーションによる予備実験及び本実験を実施し、それぞれの実験結果から本提案手法の有用性を評価する。

2. 関連研究

本研究に関連する研究には、クレームを対象とした研究として、原田ら(2008)[1]の報告がある。

原田らは日本語の文書が”客観的な事柄”と”書き手の主観的表現”から構成されることに着目し、企業に寄せられるクレーム文書を分類化する手法を提案している。提案された手法では、語の意味や語間の関係を中心とした分類アプローチを実施し、クレームクラスと呼ばれる定義されたクラスに分類に成功している。

提案された手法は、企業への質問文書を対象とし、クレームというネガティブな要因の分類という観点で本研究と類似している。しかし原田らは、分類に着眼点を置いているため、分類された文書の優先付けが考慮されていない。また文書の表現など表層的な情報に着眼点を置いている。原田らの他、感情情報に基づいた質問文と回答文のマッチングに関する研究として呉ら(2013)[2]の報告もあった。

またこれまでに筆者らも質問文書からネガティブな要因明らかにするために様々な検討を行った。

まず質問文書に出現する単語を特徴ベクトルとし、教師付き機械学習手法の1つである Support Vector Machine (以下, SVM) を用いて、質問文書を模したテストデータに対し分類を試みた[3]。評価の結果、7割程度の精度にてネガティブな文書の分類できた。

次に異なる教師付き機械学習手法の検討として Naive Bayes classifier (以下, NB) を用いてネガティブな質問文書の分類について試みた[4]。

NB はメールのスパムフィルタへの適用を始め、非常に高い分類精度を誇り、実用性が高い。評価の結果、ネガティブな特徴を持つデータと既存の手法を教師データとしてネガティブな質問文書の分類に一定の精度を上げた。

検討結果から考察するに、教師付き機械学習の手法では、一定の分類が行えるものの、教師データに依存し未知語の対策に課題が残る。またネガティブな文書という抽象的なカテゴリへの分類しか行えない。加えてネガティブな文書というカテゴリに分類された文書は、早急に対応すべき文書と必ずしも言えない。

またこれまでの手法では質問文書の単語や熟語を始めとした表層的な情報に着眼点を置いている。質問文書の要因分析では、質問者と回答者のマッチングに横山ら(2013)の報告 [5]があるが、こちらも”感情語”表層的な情報から判別している。

高村ら[6]らによれば複数語表現のネガティブ・ポジティブの度合いは、その構成語の極性の単純な和ではないと主張している。例えば”リスク”と”低い”と単語はそれぞれがネガティブな単語だが、”リスクが低い”という文書になれば、ポジティブな文書に反転する。

このことから質問文書におけるネガティブさの度合いは表層的な情報のみでは判別が困難であると思われる。そこで本研究では質問文書に含まれる潜在的な要因に着目し、質問文書の解析を試みることにした。

3. 本研究の目的

本研究では、質問内に潜在的に含まれているネガティブさの因子である深刻度の傾向を明らかにすることを目的とする。

質問文書の深刻度を傾向が明らかになることで、質問のネガティブさが一意的となり、顧客関係管理分野への応用や自動応対などの人工知能分野への応用も期待できる。本研究では、各質問文書のネガティブさの傾向を明らかにすること及び潜在的な情報抽出を提案手法によって実現する。提案手法については次章で述べる。

4. 提案手法

提案手法について述べる。提案手法は二つの方法によって構成される。

一つ目は質問文書におけるネガティブさの傾向解析である。ネガティブさの傾向解析には、定量化の実施が必要である。定量化においては評価極性分類を採用する。評価極性分類は製品のレビュー評価を始めとした評判分析に用いられており[7][8], ネガティブさなど、人間の感情に基づいた定量化の手法としては有用である。

二つ目は潜在的な情報の解析である。潜在的な情報は単語を始めとした表層的な情報のように一意的でない。そのため、様々なアプローチが考えられるが、近年、自然言語処理の分野では、文書が複数の潜在的な情報から構成されていることが明らかになり、その潜在的な情報をトピックと呼んでいる[1]。トピックを特徴的な要素に要約する研究は盛んに行われているが、本稿ではそのうちの代表的なモデルの一つである潜在的ディリクレ配分法を採用する。

潜在的ディリクレ配分法は、文書はある特徴的な単語の分布である”トピックの混合分布”が存在していることを仮定し、”トピックの混合分布”から文書が生成されるという考えに基づいている。

この特徴的な要素の傾向とネガティブさの度合いの傾向が明らかになれば、ネガティブさの潜在的な情報の分布が明らかになり、深刻度の特定につながると推察される。本章では提案手法で採用した手法について述べる。

4.1. 評価極性分類

評価極性分類について説明する。評価極性分類とは、

単語の評価極性の傾向を示す値に着目し、文書内に含まれる単語の極性値の平均値から文書全体がネガティブかポジティブかを判定する手法である。

評価極性値とは、単語が持つポジティブ・ネガティブの度合いを数値化した値である。

評価極性値を用いた代表的なものには、文書全体の評価極性の平均極性値に基づいて評価する手法にTurney(2002)の報告がある[7][8]。

Turneyの研究は、レビュー文書などの製品に対する評判が記載された文書を対象とし、文書に含まれる感情極性値から肯定的なレビューと否定的なレビューに分類可能であるとしている。

尚、評価極性分類では、ポジティブ・ネガティブの度合いとなる単語の極性値を求め、判定の尺度とする評価極性値が必要とされている。

4.1.1. 評価極性分類における尺度

極性値とは単語のポジティブおよびネガティブ度合いを数値化したものである。前述のTurneyらの研究では、コーパスから得られる共起情報を利用して語句の評価極性値を計測している。

極性値は評価表現となる単語 t の評価極性値 $SO\text{-score}(t)$ の算出方法に関して、以下の式に示す。

$$SO\text{-score}(t) = PMI(t, \text{excellent}) - PMI(t, \text{poor})$$

上記の PMI (Pointwise Mutual Information)は、2つの単語間の共起を測る尺度である。提案された手法ではポジティブの基準に"excellent"、ネガティブの基準に"poor"を設定し、それぞれとの共起尺度からスコアを算出している。次に PMI の算出方法に関して以下に示す。

$$PMI(a, b) = \log_2 \frac{P(a, b)}{P(a) * P(b)}$$

上記の $P(a, b)$ はコーパス内において単語 a と単語 b が同一文で共起する確率である。また $P(X)$ は、単語 X をコーパス内で出現する確率を表している。

PMI はコーパス内におけるそれぞれの単語の生成確率から同時生成確率のスコアを算出している。

Turney らの手法では、信頼性の高い共起情報を得るために巨大なコーパスが必要とされている。そこで我々は既存の評価表現辞書のうち日本語に対応し一定の成果が報告されている高村ら[9]の手法を採用した。

高村らが提案した手法はスピンモデルにより単語の極性値を計測し、評価極性辞書とした。高村らの評価表現辞書を以下に示す(図1参照)。

優れる	: すぐれる	: 動詞	: 1
良い	: よい	: 形容詞	: 0.999995
喜ぶ	: よろこぶ	: 動詞	: 0.999979
褒める	: ほめる	: 動詞	: 0.999979
めでたい	: めでたい	: 形容詞	: 0.999645
(～中略～)			
ない	: ない	: 助動詞	: -0.999997
酷い	: ひどい	: 形容詞	: -0.999997
病氣	: びょうき	: 名詞	: -0.999998
死ぬ	: しぬ	: 動詞	: -0.999999
悪い	: わるい	: 形容詞	: -1

図 1 高村らの評価表現辞書[9]

図 1 では、「良い」の評価極性値が[0.999995]で positive な極性を持っていることを示している。また「悪い」、「酷い」などは負値であるため negative な極性を持っていることを示している。

高村らの評価表現辞書では 1 から -1 を範囲とし、小数点第 6 位までの値を持っている。単語数は、英語(見出し語)・日本語(見出し語・読み)を合わせると、198265 個の単語が登録されている。

4.2. トピックモデル

トピックとは、話題や分野など文書に含まれる特徴である。統計的言語モデルのトピックモデルには、文書ごとに 1 つのトピックを選ぶユニトピックモデルと文書と単語ごとにトピックが生成されるマルチトピックモデルがある。

今回は、データマイニングとしての有用性があり[10]、且つ潜在的な情報から、トピックを抽出することが可能な、潜在的ディリクレ配分法によるトピック抽出手法を採用した。

4.2.1. 潜在的ディリクレ配分法 (LDA)

潜在的ディリクレ配分法: Latent Dirichlet Allocation (以下, LDA) について述べる。LDA は Blei ら[11]によって提案されているマルチトピック抽出モデルの一つである。LDA は、文書はある特徴的な単語の分布である"トピックの混合分布"から文書が生成されるという考えに基づいている。今回は LDA の文書生成の流れに着目した。トピックモデルでは、トピックの混合分布にもとづき文書のトピックが生成され各トピックの単語分布にもとづき、単語が生成される。

この生成過程に着目し、ネガティブなトピックより多く生成された文書は、ネガティブの度合いが高い文書となることを期待した。尚、LDA のトピック分布を求める手法には Blei らの変分ベイズと呼ばれるベイズ事後分布から近似して求める手法があるが、今回は Griffiths ら[12]が提案し、実用性が高いとされているギブスサンプリングの手法を用いた推定手法を用いる。

4.3. 評価実験の流れ

まず予備実験として評価極性分類の手法により質問文書を評価極性値により定量化し、ネガティブさの傾向を解析する。

その後得られた結果を元に、提案手法である LDA と評価極性値を組み合わせたトピックの解析手法を本実験において実施し、トピックの分布に関して、結果を評価する。

5. 予備実験

本実験に先立ち、質問文書の傾向を解析するために、本稿の評価対象に対して、既存の評価極性分類の手法を用いて質問文書を評価する。

対象とする質問文書は製品に対するクレーム等、ネガティブな傾向を明らかにするためにも、企業に寄せられる質問により近いデータが最適である。そこで我々はサポートコミュニティの質問文書に着目した。

サポートコミュニティとは、企業が公式のサポートセンターの他に、製品を利用しているユーザ同士の相互サポートを促進するために提供しているインターネットコミュニティである。

特定の製品群に依存するものの、単純な質問からトラブルシューティングまで、サポートコミュニティのデータは企業に寄せられる文書に非常に近い。

そこで今回は、公開されているサポートコミュニティのうち、Apple Inc.によって提供されている Apple サポートコミュニティ[13]のデータを使用する。

サポートコミュニティ内の iPhone カテゴリのうち、"iPhone の使い方"及び、"iPhone ハードウェア"のデータ（計 10365 件）を今回の解析対象とした。

5.1. クレンジング処理

評価対象とする文書に対し、質問文書として不要な箇所を取り除くため、評価対象に対し、データのクレンジング処理を行った。

クレンジング処理とは、テキストマイニングの対象に対して行う不要な文字列の除去等を指す。本稿におけるクレンジングの手順だが、まず対象となるサポートコミュニティ内の質問 1 件を、タイトル及びその説明文をマージさせ 1 つの質問文書として取り扱った。

そして文書を単語（形態素）の集まりとして表現するため形態素解析を行った。尚、形態素解析には工藤ら(2004)の開発した形態素解析エンジン MeCab[14]を用いた。その後、評価極性値算出のため、形態素解析した結果と高村らの評価表現辞書を突合させ、評価表現辞書に登録されている単語以外を排除した。

5.2. 評価極性分類による評価実験

クレンジング処理後の評価対象に対し、評価極性分類を用いて定量化し、質問文書群のネガティブさ傾向の評価した。分類の基準は、0 をニュートラル、1 をポジティブ、0 をネガティブとして評価した。

その結果、ポジティブとして判定された文書は 10365 件のうち、わずか 30 件であり、残りの 10000 件以上はネガティブと判定された。図 2 及び表 1 から文書が持つ評価極性値の分布状況に関して示す。

図 2 の縦軸は質問文書の持つ文字数、横軸は文書の平均評価極性値である。図から見て取れるとおり、ほとんどの質問文書が 0 以下の数値を示している。また図 2 から見て取れる通り、評価対象とした質問文書では分布は平均感情極性値(-1.00)から(-0.00)の区間に集中した。

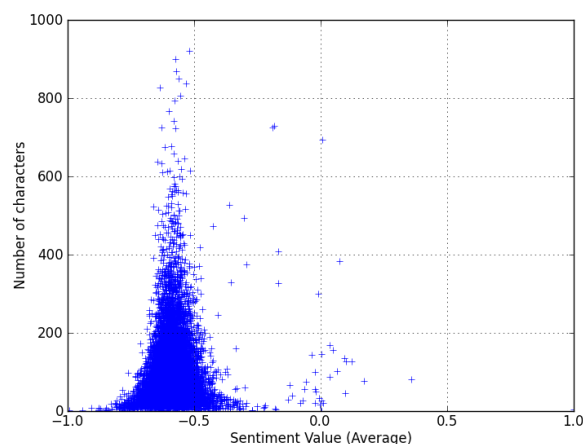


図 2 質問文書の文字数と平均評価極性値
表 1 評価極性分類による実験結果

感情極性値		文字数	
平均	-0.5876	平均	101
最大	1	最大	920※1
最小	-0.9949	最小	1
中央値	-0.5912	中央値	77
最頻値	-0.61※2	最頻値	30

※1 文字数の最大値は 1 件 7461 文字が観測されたがはずれ値とした

※2 最頻値は小数点 2 位までを基準としてカウントした

5.3. 評価極性分類による結果考察

結果から、少なくとも質問文書は評価極性においてはネガティブな値を持つことが明らかになった。また文字数と評価極性値の関係にガウス分布の形が取られていることから相関関係は見受けられない。

これは質問者が質問を行う際、自身では解決困難な内容に基づいて文書を作成するため、ネガティブな表現の品詞を多用するためであることが考察される。

この結果から、質問文書はレビュー文書とは異なる特性があることが明らかになった。

また単語という表層的な情報から文書全体の評価評価極性値を求める既存の評価極性分類の評価軸では、質問文書の分類は困難であると考察される。

6. 本実験

予備実験の結果を基に、今回の提案手法を用いて、潜在的な情報を評価極性を基に解析する。

潜在的な情報の解析にあたっては、文書群の全体から評価する他、文書群のネガティブな度合いに応じて、潜在的な情報の分布の変化の解析のため、文書群を複数にレベル分けし、評価することとした。またレベル分けの方法も3パターン用意した。

6.1. 質問文書のレベル分け

まず質問文書のレベル分けについて説明する。書群のネガティブな度合いに応じて対象とする質問文書を5段階のレベル（Level1 - Level5）にレベル分けする。

レベル分けの定義に関しては、評価表現辞書をもとに次の2つの観点でレベル分けする。尚、評価表現辞書に関しては次項で述べる。解析対象は傾向解析で対象としたクレンジング処理済みの質問文書群とし、評価表現辞書も同様に高村らの評価表現辞書を用いた。

6.1.1. 評価表現辞書の単語（パターン1）

まず評価表現辞書内の単語が同数となるよう、5段階に分割した。高村らの評価表現辞書に定義されている単語は日本語(名詞・よみ)及び英語の総品詞(198265)個がある。今回は評価極性値が0未満の単語(145871)個を対象に感情極性値の値が大きい順に(29174)個ごとに評価表現辞書の単語を5段階にレベル分けし、質問文書がどのレベルの単語を最も用いているかを質問文書のレベル分けの分類軸とした。これをレベル分けのパターン1とする。

6.1.2. 質問文書の平均極性値（パターン2）

傾向解析の際に算出した質問文書の平均極性値を元に、文書数が同数となるよう5段階に質問文書をレベル分けした。これをレベル分けのパターン2とする。

6.1.3. レベル分けなし（パターン3）

文書群全体解析とするため、レベル分けを行わないパターンとして定義する。これをパターン3とする。

6.1.4. レベル分けの結果と考察

前述に定義したパターン別にレベル分けされた結果を表2に示す。

パターン2は同数、パターン3は分類なしのため、結果は明らかであったが、感情極性辞書を元にしたパターン1では、図及び表から見て取れるように偏った分類結果となった。これは質問文書の平均感情極性がガウス分布をとっていたことから、文書内に生成される単語にも偏りが見られたためだと推察される。

表2 パターン別 レベル分け結果

	パターン1	パターン2	パターン3
Level1	22	2073	
Level2	64	2073	
Level3	735	2073	
Level4	7194	2073	
Level5	2350	2073	10365

6.2. トピックモデルによる評価実験

各パターンにレベル分けした文書群から、LDAによってトピックを抽出し、解析した結果について述べる。

解析方法はトピックに所属する単語から得られる評価極性値とトピック内の品詞構成比率とした。

6.3. 評価極性値によるトピック分布

トピックの平均評価極性値を算出し、得られた結果を以下に記す。尚、本章図中の Topic Sentiment Value はトピックの平均感情極性値を示し、Topic Index は評価極性値の降順にトピックをソートした際に得られた Index 番号である。また図中ではパターンの表記を "Pattern1" のように表記している。

6.3.1. パターン別トピック比較

パターン1およびパターン2文書群に対して実施したトピックの分布を次ページ図3及び図4に示す。

パターン1の分布には以下のような傾向が見られた。

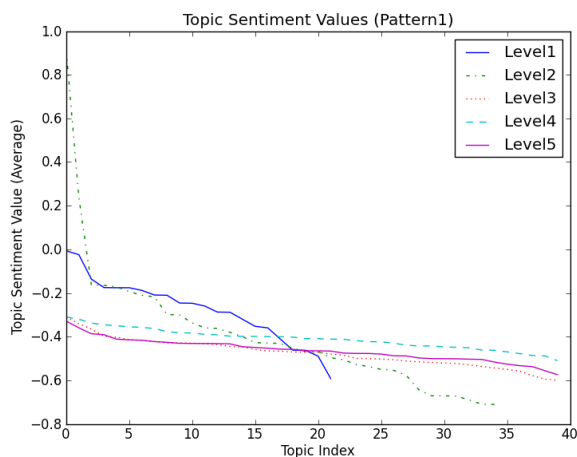


図3 パターン1トピック分布

<パターン1のトピック分布>

- ・レベル別で統一した傾向が結果が見受けられない
- ・Level2の極性値は最大約0.8、最小約0.5と広範囲

パターン 2 の分布には以下のような傾向が見られた。

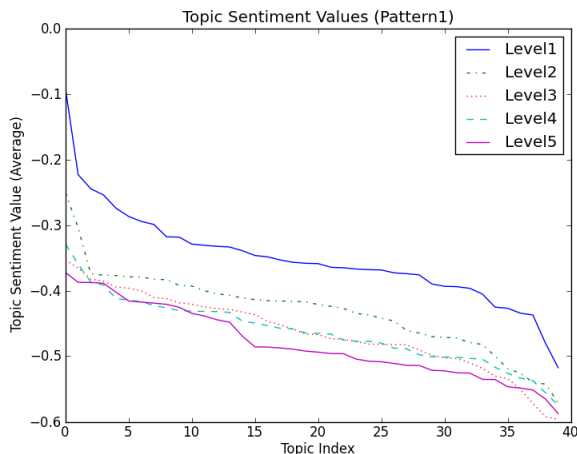


図 4 パターン 2 トピック分布

<パターン 2 のトピック分布>

- ・レベル別で統一した傾向が結果が見受けられる
- ・Level1 は Level2 から Level5 の極性値と比較して差が大きい
- ・Level5 は Index13 以降, Level2 から Level4 と比較して差が大きい
- ・グラフ終端箇所である Index35 以降すべてのレベルにおいて, 最小値である -0.6 付近に収束している

6.3.2. パターン別比較の考察

文書群を評価極性値に基づき分類したことから, 例外を考慮しても, レベル別の結果に統一した傾向が現れることが一致することが容易に推測される. だがパターン 1 においては, それが確認できず, 高い評価極性値を持つレベルであるはずの Level2 トピックが Level5 のトピックの感情極性を下回った結果となった.

この結果から, 評価表現辞書の単語に着目したレベル分けは適切ではなかったことがわかる.

パターン 2 に関しては, レベル別の分布に統一した傾向が見て取れてわかり, Level1~Level5 まで概ね予想どおりの結果となった.

しかしながら Level5 の Index14 付近の極性値の低下や, すべてのレベルで終端部分に統一して見られた極性値の低下など, ある箇所で評価極性値の急落がみられた. これは評価極性値の低い単語の集まりが特徴的に表れたものと考察される.

6.3.3. レベル別トピック比較

次に, レベル分けのパターンごとの同一レベルの比較を行う. レベルごとのトピックを以下, 図 5 及び図 6 に示す.

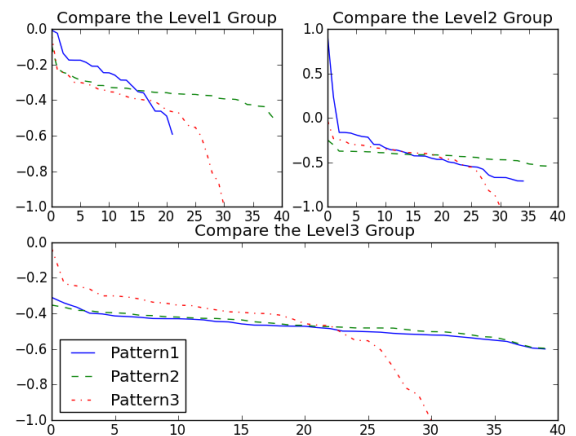


図 5 Level 1- Level 3 のトピック分布

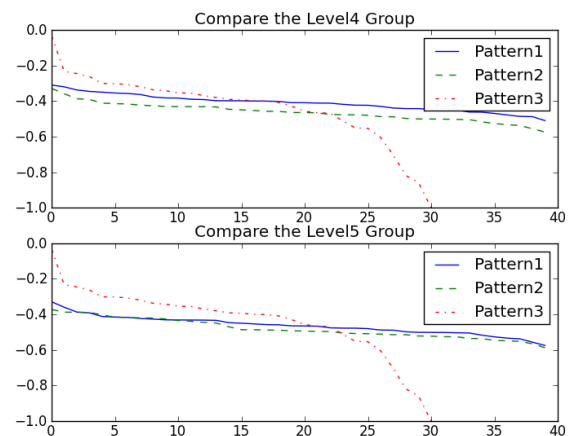


図 6 Level4 - Level 5 のトピック分布

<Level1 - Level3 のトピック分布>

- ・Level1 ではパターン 1 が急落している
- ・パターン 3 のトピック分布が Index25 以降, パターン 1 及びパターン 2 と比較して差が大きい

<Level4 - Level5 のトピック分布>

- ・パターン 1 とパターン 2 で差が小さい
- ・Level4 ではややパターン 1 の平均極性値が高い

6.3.4. 感情極性値による評価結果の考察

レベル別の比較では, Level4 と Level5 のトピックの分布が近いという結果となった. これはパターン 1 のレベル分けで多くの文書が Level4 および Level5 に分類されたことから, トピック分布が妥当な形をとったためだと考察される.

またパターン 3 のレベル分けにおいても, Index25 以降, 急落を確認できたことから, トピックの評価極性値は, 一定の Index 番号以降のネガティブな単語が局所的に集まる可能性があることが考察される.

6.3.5. 得られたトピックの考察

本実験によって得られたトピックの一例を示す。

Positive Topic		Negative Topic	
“撮影”	“アプリ”	“サービス”	“ソフトウェア”
camera	twitter	キャンペーン	OS
レンズ	drop,box	不愉快	バグ
ビデオ	yahoo	最悪	不自然
ズーム	newsstand	シヨップ	待つ
直射	bookmark	直営	操作

図 7 トピック抽出結果の一例

※1 抽出トピックから第一著者が特徴的と思われる箇所を抜粋した。

※2 図中の 1 行目および 2 行目はラベルであり，第一著者が付与した。

本実験によって得られたトピックから，ネガティブな要因と見受けられる特徴が抽出できた。このことから，トピックはネガティブな要因の抽出が期待でき，トピックの評価によって深刻度の傾向が明らかになると期待される。

7. まとめ

今回の成果について述べる。質問文書の定量化によって質問文書の平均評価極性値はネガティブな評価極性をとることが分かった。また潜在的な情報であるトピックを評価極性に基づいてトピック分布を評価した結果，ある特定の箇所から評価極性値の急落を確認できた。これはネガティブな極性を持つ単語の局所的な集合である可能性が考察される。加えて，実験から得られたトピックからネガティブな要因と見受けられる特徴を得ることができた。

今後は，得られたトピックの表層的な情報と潜在的な情報との関連性について評価していきたい。

参 考 文 献

- [1] 原田 実, 川又 真綱, "クレーム内容の自動分類", 言語処理学会年次大会発表論文集, 14th, 293-294, 2008 年 03 月 17 日.
- [2] 呉 鍾勲, 鳥澤 健太郎, 橋本 力, 川田 拓也, デサーガ ステイン, 風間 淳一, 王 軼謳, "意味的極性と単語クラスを用いた Why 型質問応答の改善", 情報処理学会論文誌 54(7), 1951-1966, 2013-07-15.
- [3] 輪島幸治, 小河誠巳, 古川利博, "テキスト評価分析を用いたヘルプデスク効率化手法の提案", 2013 年度 経営情報学会 春季全国研究発表大会, 2013-06.
- [4] 輪島幸治, 小河誠巳, 古川利博, "ヘルプデスクにおける製品事故に関するクレーム情報の分類～感情極性と製造物責任法に関する訴訟情報から～", 電子情報通信学会技術研究報告 信学技報 113(327), 9-14, 2013-11-28.
- [5] 横山友也, 宝珍輝尚, 野宮浩揮, 佐藤哲司, "文章の特徴量を用いた質問回答文の印象の因子得点の推定", 日本感性工学会論文誌 12(1), 15-24, 2013.
- [6] 高村 大也, 乾 孝司, 奥村 学, "隠れ変数モデルによる複数語表現の感情極性分類", 情報処理学会論文誌 47(11), 3021-3031, 2006-11-15.
- [7] 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", Journal of natural language processing 13(3), 201-241, 2006-07-10.
- [8] Turney, P. D. (2002). "Thumbs up? thumbs down? Semantic Orientation Applied to Unsupervised Classification of Reviews." In Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics (ACL-2002), pp. 417-424.
- [9] 高村大也, 乾孝司, 奥村学, "スピンモデルによる単語の感情極性抽出", 情報処理学会論文誌ジャーナル, Vol.47 No.02 pp. 627--637, 2006.
- [10] 奥村 学, "マイクロブログマイニングの現在", 電子情報通信学会技術研究報告. NLC, 言語理解とコミュニケーション 111(427), 19-24, 2012-01-26.
- [11] Blei et al, "Latent Dirichlet Allocation", Journal of Machine Learning Research, Vol. 3, pp. 993-1022, 2003.
- [12] T. L. Griffiths and M. Steyvers, "Finding scientific topics," Proc. of the National Academy of Sciences of the United States of America, vol.101, pp.5228-5235, 2004.
- [13] アップルジャパン合同会社, "Apple サポートコミュニティ", <https://discussions.apple.com/>, 最終閲覧日: 2014/1/5.
- [14] Taku Kudo, Kaoru Yamamoto, Yuji Matsumoto: Applying Conditional Random Fields to Japanese Morphological Analysis, Proceedings of the 2004 Conference on Empirical Methods in Natural Language Processing (EMNLP-2004), pp.230-237 (2004.)