

SNS を利用した 協調的サーベイ支援のための関連文献組織化の提案

Proposal of Related Paper Organization for Supporting Cooperative Literature Survey Using SNS

蜂谷 聖未^{1*} 高間 康史¹

Masami Hachiya¹ and Yasufumi Takama¹

¹ 首都大学東京大学院システムデザイン研究科

¹Graduate School of System Design, Tokyo Metropolitan University

Abstract: 文献サーベイは重要な研究活動の一つであるが、グループ内での各研究者の研究テーマは互いに関連があると考えられるため、協調することで効率よく行えることが期待できる。本稿では、各研究者がサーベイした文献をグループメンバー限定の SNS 上で共有し、個人に応じた関連文献の組織化を行い提示する事で文献サーベイを支援する手法を提案する。観点に基づく要約の作成、コメントによる論文の関連づけなどにより効率的な組織化を実現する手法について述べる。

1 はじめに

文献サーベイは重要な研究活動の一つである。その進め方としては、個人で行うサーベイから、複数人で同じ文献をサーベイする輪講など様々であるが、本稿では関連ある研究内容について、個々人がサーベイしつつも情報交換するような、グループとしての文献サーベイに着目する。大学の研究室や、企業の部署などのグループ内での各研究者の研究テーマは互いに関連があると考えられるため、協調することで効率よく文献サーベイを行えることが期待できる。

本稿では、各研究者がサーベイした文献をグループメンバー限定の SNS 上で共有し、利用時には研究者個人に応じた関連文献の組織化を行い提示する事で文献サーベイを支援する手法を提案する。提案手法では、観点に基づく要約の作成、コメントによる論文の関連づけなどにより、効率的な組織化を実現する。提案手法を用いることで、関心のある観点において類似する文献をまとめて確認することができる。また、グループ内の他者が読んだ関連文献もサーベイの対象とすることができる。これらの特徴により、効率的なサーベイが可能になると考える。

本稿では、提案手法の概要説明、および開発中のプロトタイプシステムについて説明する。2 節では文献サーベイ及びその関連研究についてまとめる。3 節で提案手法について説明し、4 節では開発中のプロトタイプシステムについて説明する。

2 関連研究

文献サーベイとは、既存の資料や研究を調査することであり、該当する分野の現状や動向を把握するため、そして従来の研究の知識を獲得するために必要不可欠である。本節では、サーベイのプロセスを、文献の収集、読解、読後の整理の 3 つに分け、それぞれについて利用可能なシステムや関連研究についてまとめる。

文献の収集において、現在一般的に利用されているものは Web 上での検索エンジンである。特に、Google Scholar や CiNii といった文献データベースの検索システムサービスは広く活用されている。これらのサービスは、通常の Web 検索と同様にキーワードによる検索であるため、問題点として、ユーザの知識が及ぶ範囲でしか文献サーベイが行えないこと、また、知識の少ないユーザに対しては必要以上に多数の文献を提示してしまうことが挙げられる。例えば、分野が違っていてもアルゴリズムが参考になる文献をサーベイする時は、ユーザがそのアルゴリズムにおける専門用語などの知識を適切に持っていない

*連絡先： 首都大学東京システムデザイン研究科
〒191-0065 東京都日野市旭が丘 6-6
E-mail: hachiya-masami@sd.tmu.ac.jp

ければならない。また、そのような文献を見つける際に、アルゴリズムが参考になるという事を把握するためには多数の文献の不必要な部分も精読しなければならず、ユーザの負担が大きくなる。

文献の読解に関しては、グループで行う読解として代表的なものに輪読がある。輪読のやり方は様々なものがあるが、例えば、グループの研究に関連がある文献を分割して各メンバに割り当て、担当となった部分の要約を発表し、グループ内で見解を話し合うという方法がある。輪読の利点として、分担して読解することで個人の負担が減る点、報告というアウトプットにより理解が深まり記憶に残りやすい点が考えられる。一方欠点として、参加者全員に関連のある文献を対象とせざるを得ない関係上、各個人の関心との一致度が低くなってしまう場合がある点、発表のための資料作成に時間を多く費やしてしまう点が考えられる。

輪読のようにグループ活動を中心とした学習形態は協調学習と呼ばれる。協調学習は、学習者がグループ活動の中で互いの学習を助け合い、一人一人の学習に対する責任を果たすことで、グループとしての目標を達成していく、協調的な相互依存学習である[1]。その効果として、競争状況や個人学習状況よりも学習効果が高いことや、動機づけやコミュニケーション能力、相互調整能力に効果があり、学習者相互に様々な気づきをもたらすことが指摘されている[2]。協調学習の際の情報共有の場として、近年 SNS が注目されている。SNS をグループメンバ間のコミュニケーションツールとして利用し、メンバが個々に所有する諸問題や情報をグループメンバ間で共有することで、個人の学習に間接的に影響を与えることが期待できる[3]。しかし、文献サーベイに関するソーシャルメディアを利用した協調学習としては、文献読解よりも後述する文献整理支援を対象としたものが多い。

読後の文献整理に関しては、サーベイした文献が増えるほど管理が大変になり、逆に研究効率が悪くなってしまう場合がある。研究者を支援するツールの一つとして「Mendeley」¹がある。「Mendeley」は、パソコンにインストールするアプリケーション「Mendeley Desktop」と、ウェブブラウザで利用できる「Mendeley Web」を組み合わせる。「Mendeley Desktop」では PDF ファイルの文献情報を管理することができ、「Mendeley Web」ではオンラインの研究者ネットワークを利用することができる。「Mendeley」では文献情報は「Mendeley」のサーバに保存されるが、「Mendeley Desktop」をインストー

ルしたパソコンがあれば登録内容をパソコンに同期させてローカルで利用することも可能である。さらに、研究者仲間で「グループ」を作成して、登録した文献情報に付けた注釈を共有することができる。しかし、文献の管理が目的であるため、複数文献の内容を俯瞰的に眺めたり、比較を行うと言った機能は備えていない。また、日本においては著作権上の縛りが厳しいため「Mendeley」の最大の特徴である情報共有機能は制限を受ける[4]。この他、協調学習活動を支援しながらノートの整理を行えるシステムとして、「ReCoNote」がある[5]。「ReCoNote」は Web ブラウザ上で動作するシステムで、学生同士が互いに調べた内容（ノート）を共有し、相互に関連付け、全体をまとめる作業を支援する機能を搭載している。ログイン時に上下 2 つに異なった内容のノートが表示され、自分のノートの中に考えた内容を記録したり、他人のノートを閲覧することができる。個人のノートの他にグループのノートを作成することもできる。これらのノートは学習者自身が互いの間にリンクを貼ることによって関連付けが行われる。しかし、ユーザ自身が関連性を見極めてリンクを貼るため、ユーザが意識していない観点による共通点を見つけ出すことは難しい。また、「ReCoNote」は協調学習活動の支援が目的であるため、複数ユーザで一つの成果を出す学習での利用に適している。

3 提案システムと手法

3.1 提案手法の概要

本稿では、文献のサーベイをサポートするために、グループ内での協調的な文献サーベイを想定する。グループ内での各研究者の研究テーマは互いに関連があると考えられることから、個人の文献サーベイ結果をグループメンバで共有することで、文献サーベイを効率的に行えることが期待できる。また、あくまでも個人での文献サーベイ活動を基本に置くため、輪読とは異なり各自の関心に応じ自由に文献を選ぶことができる。

本提案では、各研究者がサーベイした文献をグループメンバ限定の SNS 上で共有する。図 1 に、SNS 上でのサーベイ情報の共有の概要を示す。各研究者は 3.2 節で述べる観点に従って文献の要約を作成し、SNS 上にアップロードする。現在は、文献自体はアップロードせず、Web 上で公開されている文献の場合はその URL、されていない場合は書誌情報を引用することとしている。さらに、文献同士の関連づけには、文献要約へのコメントや関連文献の紹介を利用する。SNS 上でのやり取りで関連づけられた文献

1 Mendeley <http://www.mendeley.com/>

は、研究者個人に応じた組織化を行い、3.4 節で述べる表形式で提示する。

提案手法を用いることで、関心のある観点において類似する文献をまとめ、その要約を確認することができる。観点という一種の制約を要約作成に課すことで、2 節で述べた文献整理支援に関する既存研究よりも、関連文献組織化を効率的に行えることが期待できる。

また、自分自身が読んだ文献だけでなく、グループ内の他者が読んだ関連文献もサーベイの対象とすることができる。これは、2 節で指摘した文献収集における問題点の解決にも貢献することが期待できる。これらの特徴により、効率的なサーベイが可能になると考える。

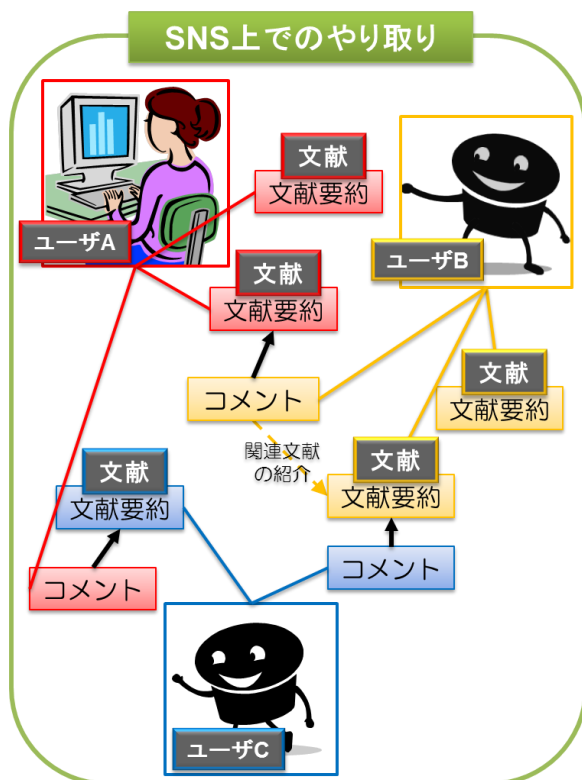


図1 SNS 上でのサーベイ共有の概要

3.2 文献要約作成

本提案ではまず、各研究者がサーベイした文献の要約を研究者自身が作成し、SNS 上に投稿する。文献要約は、指定した5種類の観点に基づき整理して作成する。5種類の観点とは、「目的」「手法・アルゴリズム」「研究環境」「評価方法」「結果」である。「目的」には、サーベイした文献の主な目的を自由記述し、研究室の研究との関連度を、かなり競合／関連あり／関連なし、の3択から選択する。「手法・アルゴリズム」については、文献で利用されている

手法・アルゴリズムの名称、あるいは概要の説明を自由記述し、それについての文献中での説明の程度を、詳細な説明あり／少し説明あり／ほとんどなし、の3択から選択する。「研究環境」はソフト・ハードに関する項目であり、プログラミング言語名、使用機器名、データベース名、ライブラリ名などを自由記述し、それが珍しいもの場合は概要の説明も記述する。「評価方法」には、その文献で用いられている評価方法についてまとめる。まず、被験者実験の有無を、行っている／行っていない、の2択から選択する。テストコレクションを使用している場合には、その名前と概要の説明を自由記述する。さらに、評価指標名を記述し、評価方法が参考になるかどうかを、かなり参考になる（しっかりされている）／多少参考になる／あまり参考にならない／評価をしていない、の4択から選択する。「結果」については、文献でまとめられている総合的な結果を自由記述形式で要約する。

観点を採用した理由として、文献を読む際に、必ずしも全ての内容が参考になるわけではないことが挙げられる。どの部分がどのように自分の研究に参考になるか、を意識しながらサーベイを行うことで、効率的なサーベイを行うスキルが見つけられると考える。また、多少異なる研究をしている研究者間であっても、例えば同じアルゴリズムが異なる対象に適用されている場合など、観点到分割して考えれば共通点が見つかる場合がある。そのため、観点をなしで議論するよりも必要なコメントが得られる可能性が高くなることが期待できる。提案システムにおいても、3.4 節で説明する組織化（表生成）を行う際に文献間の関連性の把握が容易になることが期待できる。

文献要約を SNS 上に投稿する際には、上記の5種類の観点ごとの要約の他に、文献の分野・タイトル・著者名などの情報も記載する。さらに文献の内容に関する疑問点がある場合は要約と共に質問を書くことができる。

3.3 SNS 上での共有

3.2 節で作成した文献要約は、グループメンバー限定の SNS 上に投稿し、グループ内で共有する。4 節で述べるように、本稿で利用した SNS 上では各文献要約の投稿記事にコメントをつけることが可能であるため、コミュニケーションを記録し、グループ内で公開することができる。SNS 上で全体公開せずに、研究グループ内での閉じたグループに限定する理由は、SNS 上でのグループを現実と同じ空間に絞ることで、相手の知識や研究環境を踏まえた具体的なア

SNS 上からユーザに関連する文献を収集し、3.4 節に述べた方法で、Excel のシートとして表形式にまとめた例を図 3 に示す。3.4 節で述べたインタラクティブな閲覧機能は、Excel のマクロ機能を用いて実装している。図 3 に示すシートとは別のシートに観点毎の文献間関連度の計算結果を記録しておき、それを参照することで観点・文献指定によるソートを実現している。ただし、3.4 節で提案した機能のうち、色の濃淡を付ける機能については今後実装の予定である。

1	A	B	論文001	論文002	論文003
2	ルーム名	ルーム名	論文001	論文002	論文003
3	収録巻 目付	収録巻 目付	2012年4月号01 0055	2012年4月号01 0056	2012年4月号01 01
4	タイトル	タイトル			
5	著者名	著者名			
6	■論文情報■	■論文情報■			
7	出版	出版			
8	発行年	発行年			
9	URL	URL			
10	分類	分類	アンケーザ解析クラスリング、HiGraph、MDS	時系列データ、クラスリング、詳細検索、クラス選択	クラスリング・絶対クラスリング・数値表示
11	観点	観点	【手法・アルゴリズム】	【手法・アルゴリズム】	【手法・アルゴリズム】
12	■目的■	■論文の目的	→横を行われているアンケートの形式にて、用意された回答を選択し、そこから選択された、自由に文章で回答して、そこから自由回答形式である。 自由回答形式のアンケートでは、回答者が自由に文章で回答することができ、回答者の文章の表現の自由が期待できる。 しかし、自由回答形式のアンケートは、回答者の回答の自由であること、分析に多大な時間と労力を要する。そのよまでは多量な回答者と十分な回答を、分析することができない。 したが、回答の自由が期待されている。そこで、文章を生成して、自由回答形式の回答者の文章をキーワードとして抽出する手法を提案し、適切な分類および文章の類似性の把握が可能となることを示す。	→横断研究の普及に伴い、計数値システムでのデータ・分析は非常に重要となる。 →多大な情報が必要である。その情報の中には、時系列データで表現できるものがある。 時系列データには、時系列の可視化システムがあり、時系列の可視化システムは、時系列の可視化システムであり、時系列の可視化システムは、時系列の可視化システムである。 時系列の可視化システムは、時系列の可	

Twitter 上で共感を生み出すツイートの性質に関する考察

Consideration on Characteristics of Sympathy-arousing Tweets on Twitter

大川 陽聡^{1*} 高間 康史¹

Hisato Ohkawa¹, Yasufumi Takama¹

¹ 首都大学東京

¹Tokyo Metropolitan University

Abstract: Twitter 上では様々なコミュニケーションが発生しており、その情報伝播特性や震災時に果たした役割など、多様な観点から分析がなされている。我々は、不特定の相手向けになされたツイート（つぶやき）に対し多数の人が共感を抱くケースに着目し、その発生メカニズムの解明を目的としている。本稿では、ツイート自体が持つ特徴に着目して機械学習を適用することにより、共感を生み出すツイートの性質について考察する。

1. はじめに

Web は情報の共有やコミュニケーションの場として広く活用されており、特に近年のスマートフォンの発展・普及によりその傾向はさらに強まっている。Web を活用するためのサービスとして、ブログ・SNS・ミニブログ等が広く普及しているが、各サービスにおいてやり取りされる情報やコミュニケーションに対して多様な観点から分析がなされている。特にミニブログサービスの代表的な例である Twitter¹については、平成 23 年 3 月の東日本大震災時に被災地や避難所の情報共有や被災者への励ましなどに活用されたことが知られている。ここで注目すべきは、役立つ情報の共有だけでなく、他者とつながるためのメディアとして Twitter が利用されている点である。特に、Twitter や Face book のようなソーシャルメディアでは、他者とつながる事自体を目的とした利用の割合が高いと考える。この時重要なのは、役立つ情報を含んでいるか否かより、読み手の共感を得ることが重要であると考えられる。

我々は、不特定の相手向けになされたツイート（つぶやき）に対し多数の人が共感を抱くケースに着目し、そのメカニズムを解明することを目指し研究を進めている。本稿では、ツイート自体が持つ特徴に

着目して機械学習を適用することにより、共感を生み出すツイートを識別する分類器を生成し、共感を発生させるツイートが持つ性質について考察する。

2. 関連研究

2.1. Twitter 上における情報伝播

Twitter 上では、ユーザは他のユーザをフォロー (follow) し、また他のユーザからフォローされることにより、フォロー・フォロワーというつながりが発生する。このつながりをフォローネットワークと呼び、ユーザはフォローネットワーク上で主に以下の 6 つの機能を用いてコミュニケーションを取り合い、情報共有を行うことができる。

1. ツイート (tweet)
2. 返信 (リプライ, reply)
3. 公式リツイート (RT, official retweet)
4. 非公式リツイート
(非公式 RT, unofficial retweet)
5. 言及 (mention)
6. お気に入り (favorite)

ツイート、リプライ、リツイート、言及は短文を投稿する機能であり、お気に入りは他者の投稿内容を保存しておく機能である。

ツイートは、ユーザが 140 文字以内の短文を投稿することの総称であり、特定の相手へ向けたものではないものが大部分である。リプライは、文頭に「@ユーザ名」をつけた投稿であり、特定のユーザを対象とした投稿として用いられる。リツイートは、他者の投稿を引用し、紹介や拡散を行う機能である。

¹ <http://twitter.com>

*連絡先：首都大学東京大学院システムデザイン研究科
情報通信システム学域
〒191-0065 東京都日野市旭が丘 6-6
E-mail: ohakawa-hisato@sd.tmu.ac.jp

公式リツイートは他者の投稿を改変することなくそのまま引用するが、非公式リツイートは他者の投稿を自身の投稿の一部として用いることで、自身の意見を加えた投稿をすることができる。言及は、文頭以外の文中に「@ユーザ名」を含む投稿であり、特定のユーザに関するツイートであるという意味合いを持つ。

Twitter については、今までに以下の様な研究がなされている。

1. フォローネットワークによる情報伝播に関する研究[1]
2. Twitter ユーザの特性に関する研究[2]
3. リツイートに関する研究[3]
4. Twitter 外部からの影響に関する研究[4]
5. Twitter 上での特定の話題に関する研究[5]
6. ツイートの内容に関する研究[6]

情報伝播に関する研究としては、フォローネットワーク上での情報伝播についての研究である震災による情報伝播ネットワークの変化[1]などがなされている。Twitter ユーザの特性に関する研究としては、Twitter 上でフォローすべきユーザの推薦手法に関して、コミュニケーションに着目した Twitter フォローユーザ推薦[2]に関する研究がなされている。リツイートに関する研究としては、非公式リツイートに含まれる訂正コメントに着目した、リツイート構造を用いたデマ拡散防止支援手法[3]が提案されている。また Twitter 外部からの影響に関する研究として、Twitter 上のテキストが含む URL の傾向を把握する目的で、リンクを含むつぶやきに着目した Twitter の分析[4]が行われている。Twitter 上での特定の話題に関する研究としては、Twitter 上で注目されているキーワードに関する情報を抽出し提示するシステムの構築を目的として、Twitter におけるつぶやきの関連性を考慮した改良相関ルール抽出に関する研究[5]がされている。ツイート自体の内容に関する研究としては、任意のツイートとツイートへのリプライを対象とし、リプライの内容を分析することでツイートが持つ感情表現を推定する研究[6]がなされている。

これらの研究では、Twitter 上でのツイートの伝わり方や、リツイート、リプライなどを介した二次的な情報を中心に研究が行われているが、一時的な情報であるツイートそのものが持つ性質についての分析は行われていない。

2.2. 共感とソーシャルメディア

共感という概念に関する研究は、ソーシャルメディア分野のみでなく、経営学の分野でも行われてい

る。相原らは、ビジネスを行う特定の集団における経営イノベーションの向上を目的とし、その達成のために共感を用いた経営イノベーションと共感ネットワークについて研究している[7]。この研究では、「意思の疎通を図る＝親密になる」というプロセスを踏まえて組織内でのコミュニケーションが深まっていく段階における一部分として共感を定義しており、共感が発生するための条件には一定のコミュニケーションが必要であると指摘している。

また、ソーシャルメディア分野で用いられている共感の事例として、「SIPS」が挙げられる。「SIPS」とは、ソーシャルメディアが普及した状態で考えられる消費者の消費行動プロセスの考え方であり、共感 (Sympathize) → 確認 (Identify) → 参加 (Participate) → 共有・拡散 (Share&Spread) というプロセスのことである。これは、Twitter を「情報を共有し、新たな情報に出会える場」と認識した際の考え方であり、消費行動はユーザへの共感から始まるという仮定の基に研究が行われている。

これらの研究では、ソーシャルメディアやコミュニケーションにおける共感の重要性が指摘されているが、どのような行動が共感を得ることができるかについては議論されていない。

3. 共感を生み出すツイートの識別手法

本稿では、ツイートの共感発生メカニズムの解明を目的として、共感を呼ぶツイートの性質について考察する。分析手順として、分析対象とするツイートを収集し、データセットを構築する。収集したツイートについて、ツイートが共感を生み出すか否かに関係すると考えられる属性を定義し、特徴ベクトルを求める。最後に、特徴ベクトル集合に対して機械学習を適用し、学習結果からツイートの性質について考察する。本稿では Twitter 上におけるツイートに対する共感を「そのツイートを不特定多数のユーザが閲覧したときに、投稿した背景が想像でき、それに同感できる」ことであると定義する。ツイートを収集して特徴ベクトルを求め、機械学習によって共感を生み出すツイートとそれ以外を分類することで、ツイートの持つどの性質がユーザに共感を呼び起こさせるのかを考察する。

3.1. ツイートの収集と共感のラベル付け

ツイートの収集には Twitter Streaming API と Twitter REST API を用いる。また、本研究の対象であるツイ

ートは「不特定多数のユーザが閲覧するツイート」であるため、著名人のユーザ（投稿者）に対して収集時点における最新のツイート，および非公式リツイートを含めて一定の件数分取得する．取得した各ツイートを閲覧し，ツイートに対する共感が発生しているか否かを手作業により判別しラベル付けする．本稿では，第一著者の判断によりラベル付けを行った．以下に，共感を生み出すとラベル付けされたツイート，それ以外とラベル付けされたツイートの例を表 1 に示す．共感を生み出すと判定したツイートに関しては，テレビ番組の録画に失敗しているという悲しい状況から思わずツイートをしたくなるという経験がある人は比較的多いと考えられる．また，共感を生み出さないと判定したツイートは，宣伝がメインであり，共感する人は少ないと考える．

表 1 ツイートのラベル付けの例

ラベル	ツイート本文
共感を生み出す	録画失敗！！絶望だ。。。なんて最悪な日だ。。。
共感を生み出さない	息子が初めてしゃべる言葉を「パパ～バカ～」を目指してます。「●●のピンネタ LIVE」■/■（水）開場 18：30／開演 19：00【会場】▲▲【チケット】前売 1500 円／当日 1800 円（全席自由）P コード：××-×× チケット発売中！

3.2. 属性の定義と抽出方法

本稿で用いる属性を表 2 に示す．これらの属性は共感に関係すると考えられるもののうち，Twitter 上で使用可能なサービスやツイート内のテキストから取得可能であることを考慮して決定している．表 2 に抽出した属性名と属性値の形式および取りうる値を示す．fav, rt はユーザに依存する属性，それ以外はツイートに依存する属性である．fav および rt は，それぞれ対象ツイートが持つお気に入り登録数と公式リツイートをされた回数である．これらは，Twitter Streaming API のパラメータを参照することによって取得可能である．お気に入り登録やリツイートなどはユーザに閲覧されることにより発生するため，共感を生み出す前段階である閲覧の有無に関連すると想定している．term は対象ツイートと対象ツイートの直前に投稿されたツイートの投稿間隔であり，単位は分以下を切り捨てた時間である．これらは，Twitter Streaming API のパラメータから参照した値から計算して求める．頻繁に出現するツイートとた

まにしか出現しないツイートの間に共感発生の有無に関して差異があるのではないかと考えるに基づきこの属性を定義している．characters は対象ツイートの文字数であり，twitter の仕様上 1 から 140 の値を取りうる．テキストの文字数はツイートの第一印象に関連すると考える．unofficial は対象ツイートが非公式リツイートをしているか否かであり，「RT @」をテキスト内に含む場合に yes，それ以外は no の値をとる．gladness, anger, sadness, pleasure はそれぞれ喜怒哀楽の感情表現がテキスト内に含まれているか否かであり，判定には感情語辞書を作成し用いた．ツイートに含まれる感情表現によって，閲覧したユーザへの影響があるのではないかと考えるに基づいている．agreement はツイートが肯定的か否定的か中立かということを示しており，thinking はツイートが賛成的か反対的か中立かということを示している．これらはそれぞれ賛成語辞書，肯定語辞書を作成して判断している．ツイートの持つ印象がユーザになんらかの影響を与えるとの考えに基づき定義している．

表 2 属性名と属性値の形式

属性名	属性値の形式
fav	real 型, 0～
rt	real 型, 0～
term	real 型, 0～
characters	real 型, 1～140
unofficial	{yes, no}
gladness	{yes, no}
anger	{yes, no}
sadness	{yes, no}
pleasure	{yes, no}
agreement	{agree, disagree, neutral}
thinking	{positive, negative, neutral}

3.3. 機械学習

表 2 に示す各ツイートの属性値を説明属性，共感発生の有無を目的属性と定めてデータセットを作成し，機械学習を行う．本稿の目的に適した学習器が不明であるため，ここでは，対象となるデータセットに対して複数の学習器を使用し分類学習を行った結果について考察する．交差検証法を用い，適合率 (Precision)，再現率 (Recall)，正答率 (Correctly Classified Instances) により性能を評価する．各指標の定義を式 (1) (2) (3) にそれぞれ示す．

$$\text{Precision} = \frac{TP}{TP + FP} \quad (1)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2)$$

$$\begin{aligned} \text{Correctly Classified Instances} \\ = \frac{TP + TN}{TP + FP + FN + TN} \end{aligned} \quad (3)$$

ここで、TP, FP, FN, TN は以下を意味する。

- TP : 実際の共感発生は有で、有と判定した数
- FP : 実際の共感発生は無で、有と判定した数
- FN : 実際の共感発生は有で、無と判定した数
- TN : 実際の共感発生は無で、無と判定した数

共感を生み出すツイートの数は少ないことが想定され、実際に本稿で収集したデータでも 4.1 節で示すように共感を生み出すツイートの割合は低くなっている。このようなデータの特性および本稿の目的を考えると、共感を生み出すツイートを取りこぼすこと、すなわち FN が大きくなることを避けるべきである。そこで分類学習においては、FN へのペナルティが大きくなるようにコストを設定して学習を行う。

4. 実験結果

4.1. データセット

今回の実験で収集するツイートは「不特定多数のユーザが閲覧可能であるもの」を対象とするため、著名人の Twitter ユーザによるツイート、非公式リツイート、および言及を対象としている。著名人ユーザは一般人ユーザに比べてフォロワー数が非常に多い傾向があり、またユーザ層を特定しないため不特定多数のユーザが閲覧していることが多い。著名人ユーザは芸能人・企業人・学者・作家・スポーツ選手等様々な種類のタイプが挙げられるが、ユーザの絶対数が多く、ツイートの頻度が高めなタイプである芸能人、芸人を対象とした。また、リプライは特定の相手を対象としたつぶやきであること、公式リツイートは対象ユーザが第一次情報源でない投稿であることから、本実験においては除外している。

データセットは前述の芸能人、芸人の著名人ユーザの中から 5 名を選択し、各ユーザのツイートを収集し作成した。ツイート収集の手法として、各ユーザのツイート、非公式リツイート、言及であるツイートを 2012 年 9 月 29 日において最新のツイートから時系列順に数えて 50 件を取得した。作成したデータセットの概要を表に示す。

表 3 データセットの概要

ユーザ名	タイプ	フォロワー数	フォロワー数	共感発生数
庄司智春	芸人	133	12,611	7
陣内智則	芸人	174	144,809	18
有吉弘行	芸人	204	1,699,163	13
ROLA	タレント	133	1,048,947	9
ベッキー	タレント	0	826,753	14

4.2. 機械学習結果の考察

用意したデータセットに対して、データマイニングツール weka²を用いて機械学習を行った。精度の推定には交差検証法を用い、表 2 に示す属性のデータ型が利用可能で、かつコスト設定が可能なものについて学習を行った。本稿で用意した 5 つのデータセットについて全般的に性能がよかった 6 種類の手法 : SGD, SMO, MultilayerPerceptron, LWL, AdaBoostM1, LMT について表 4~表 9 にそれぞれ示す。ここで、性能評価は再現率を重視し、多くのデータセットに対して再現率が 7/10 以上でかつ適合率が 1/3 以上となるものを選択している。後者の条件は、あまりにも多数の負例を誤検出してしまうものを除くためである。また、コストは FP に対するコストを 1 とした場合の FN に対するコストを表す。

実験結果より、共感を生み出すツイートが持つ性質について考察する。陣内智則氏のデータセットを除き、ツイートの文字数 (characters) がある値以下の時に共感を生み出すと判定する学習結果が多く見られた。ツイートの長さが冗長であればあるほどユーザがテキストの全てを閲覧することがなくなってしまい、共感の発生を妨げる要因になっているのではないかと考える。また、要点がまとまっている短めの文章の方が共感を発生させやすいと考える。一人だけこのような傾向が見られなかった陣内智則氏のデータセットでは、ツイート中の文字数が増える顔文字を使用する頻度が多かったが、顔文字は感情を表現しやすく閲覧者に伝わりやすい傾向があるため他とは異なる結果になったと考える。

一方、agreement や thinking など、ツイート自体が持つ意見の傾向は共感の発生に影響を及ぼしている想定していたが、学習結果からはそのような傾向は読み取れなかった。この点については、より大きなデータセットでさらに検討をすすめる必要があると考える。gladness, anger, sadness, pleasure の感情表現に関する属性については、特に悲しみの感情を

² <http://www.cs.waikato.ac.nz/ml/weka/>

表すツイートが他の感情表現を表すツイートより共感を生み出すツイートの属性となる場合が多く見られた。悲しい記憶は人間の心に残りやすく、そのためツイートを投稿したユーザの状況をイメージしやすくなり、共感が発生する可能性が高くなるためと考える。

表 4 SGD の評価値

ユーザ名	コスト	精度	再現率	正解率
庄司智春	1.0or5.0	0.50	0.86	0.86
陣内智則	8.0	0.74	0.94	0.86
有吉弘行	17	0.63	0.92	0.84
ROLA	4.0	0.58	0.78	0.86
ベッキー	13or20	0.56	1.0	0.78

表 5 SMO の評価値

ユーザ名	コスト	精度	再現率	正解率
庄司智春	25	0.33	0.86	0.90
陣内智則	4.0	0.69	1.0	0.84
有吉弘行	2.0	0.71	0.77	0.86
ROLA	7.0	0.86	0.78	0.94
ベッキー	8.0	0.48	0.93	0.70

表 6 MultilayerPerceptron の評価値

ユーザ名	コスト	精度	再現率	正解率
庄司智春	9.0or10	0.45	0.71	0.84
陣内智則	5.0	0.84	0.89	0.90
有吉弘行	1.0	0.77	0.77	0.88
ROLA	7.0	0.58	0.78	0.86
ベッキー	20	0.65	0.79	0.84

表 7 LWL の評価値

ユーザ名	コスト	精度	再現率	正解率
庄司智春	18~20	0.33	0.86	0.74
陣内智則	1.0	0.67	1.0	0.82
有吉弘行	5.0~10	0.55	0.92	0.78
ROLA	20	0.33	0.67	0.70
ベッキー	8.0	0.48	0.86	0.70

表 8 AdaBoostM1 の評価値

ユーザ名	コスト	精度	再現率	正解率
庄司智春	15	0.25	0.57	0.70
陣内智則	4.0~6.0	0.69	1.0	0.84
有吉弘行	15	0.57	0.92	0.80
ROLA	11	0.64	0.78	0.92
ベッキー	11	0.52	0.86	0.74

表 9 LMT の評価値

ユーザ名	コスト	精度	再現率	正解率
庄司智春	18	0.24	0.57	0.68
陣内智則	5.0	0.68	0.94	0.82
有吉弘行	5.0	0.67	0.92	0.86
ROLA	4.0	0.80	0.89	0.94
ベッキー	7.0	0.52	0.86	0.74

5. おわりに

本稿では、Twitter において不特定多数のユーザがツイートを閲覧したときにツイートに対して共感するメカニズムについて調査するため、共感を生み出すツイートを識別する分類器を学習し、得られた分類規則の考察を行った。分析結果より、ツイート文字数や悲しみの感情が共感の発生に関係ある可能性のある結果が得られた。本稿で用いたデータセットはユーザ数、ユーザあたりのツイート数も小規模であるため、今後はより大規模なデータセットを用意して分析を行う予定である。また、今回のデータセットでは第一著者がラベル付けを行ったが、複数の評価者による判定でラベル付けを行う必要があると考えている。さらに、ツイートを閲覧する側の状況や性質に関する分析も必要であると考えている。Twitter ユーザは、ツイートを閲覧した際にリプライやリツイート等の様々なサービスを用いてアクションを起こすことができる。これらのアクションが共感を生み出すツイートに対してどう起こされるのか今後検討をしていく。

参考文献

- [1] 臼井 翔平, 鳥海 不二夫, 石井 健一郎, 間瀬 健二: 震災による情報伝播ネットワークの変化, 人工知能学会全国大会, 1C3-OS-12-3, 2012
- [2] 北村 太一, 小川祐樹, 諏訪博彦, 太田敏澄: コミュニケーションに着目した Twitter フォローユーザ推薦, 人工知能学会全国大会, 3E1-R-6-5, 2012
- [3] 中原 英美, 富永 一成, 牛尼 剛聡: リツイート構造を用いたデマ拡散防止支援手法, DEIM Forum 2012, F2-3, 2012
- [4] 吉田 光男, 乾 孝司, 山本 幹雄: リンクを含むつぶやきに着目した Twitter の分析, DEIM Forum 2010, A5-1, 2010
- [5] 鈴木 啓太, 新美 礼: Twitter におけるつぶやきの関連性を考慮した改良相関ルール抽出による話題抽出, 言語処理学会 第 17 回年次大会 発表論文集, pp.468-471, 2011

- [6] 堀宮 ありさ, 坂野 遼平, 佐藤 晴彦, 小山 聡, 栗原 正: Twitter における発話者へのリプライを用いたユーザー感情推定手法, DEIM Forum 2012, F2-1, 2012
- [7] 相原 憲一, 倉島 由佳子: 経営イノベーションと共感ネットワーク, 国際 P2M 学会 研究発表大会, pp. 151-160, 2010

価値判断に基づくユーザモデリング手法を用いた 情報推薦システムの提案とその特性に関する考察

服部 俊一^{1*} 毛 中杰¹ 高間 康史¹
Shunichi Hattori¹, Zhongjie Mao¹, Yasufumi Takama¹

¹ 首都大学東京大学院システムデザイン研究科

¹ Graduate School of System Design, Tokyo Metropolitan University

Abstract: 本稿では、ユーザがどの属性を重視してアイテムへの評価を決定するかという属性毎の価値判断、いわばユーザの「こだわり」を推論しユーザモデリングを行う手法を用いた情報推薦システムを実装し、ユーザの嗜好に基づき推薦を行う従来のシステムとの特性の違いを考察する。ユーザの嗜好を用いる従来手法とユーザの価値判断を用いる提案手法を組み合わせることで、より少ない情報からユーザの特性を推論することが可能になると考える。

1 はじめに

本稿ではアイテムに対するユーザの価値判断に基づく情報推薦システムを提案し、その特性について考察を行う。情報化技術の発展による情報量の増大に伴い、利用者にとって有用な情報を見つけ出す情報推薦システムが情報フィルタリングの一手法として注目されている。しかし、代表的な手法である協調フィルタリング [1] を用いた既存の情報推薦システムでは新規に利用を始めたユーザや最近追加されたアイテムに対しての情報の少なさから、推薦の精度が低くなってしまうという cold-start 問題が指摘されている [2]。また、書籍であれば著者やジャンルなど、アイテムの属性値に関する情報を推薦に用いる手法は内容に基づくフィルタリングと呼ばれ、特定の属性値に対するユーザの嗜好情報が得られれば適切な推薦が可能である。嗜好情報はユーザが明示的に与える場合と、ユーザの行動情報からシステムが推定する場合があり、ユーザの負担の観点からは後者が望ましいとされる。しかし、ショッピングサイトなどで取り扱われている多様なジャンルに属するアイテムを推薦対象とした場合、膨大な行動情報を獲得しなければ多くの属性値に対して十分な情報を集めることは困難である。

一方、個人の嗜好や消費行動を推定するための概念として「価値観 (Personal Values)」が挙げられる。価値観は個人の嗜好や消費行動に間接的な影響を与える要素であり、マーケティングや商品開発の分野では広く活用されている。価値観はアイテムの好き嫌いや良し悪しではなく、どの要素を重視するかという、アイ

テムの属性に対する価値判断を表す要素であることから、これを用いることでより少ない情報でユーザの嗜好や特性を推論することが可能になると考える。しかし、ユーザの価値観のモデル化およびそれに基づく情報推薦システムまだ確立されていない。

本稿では価値観と繋がり深い要素としてユーザの「こだわり」に着目した情報推薦システムを提案する。ユーザがアイテムのどの要素を重視しているかを推論することによってユーザの価値観をモデリングし、それに基づいてユーザの価値判断を反映させたアイテムを推薦する。本稿で取り扱うユーザモデルでは、ユーザがどの属性にこだわりを持ち、その属性に対する価値判断がアイテム全体に対する評価にどれだけ強く影響を与えるのかを推論する。このような推論手法を、ジャンル・著者などの属性値を好むかどうかを元に推論を行っていた内容ベースフィルタリングに適用することで、より少ない情報から適切な推薦アイテムを推論することが可能になると考える。本稿では予備実験として価格.com のレビュー (クチコミ) を対象としてユーザモデルを作成し、提案手法の特性や推薦への適用可能性について考察する。

2 関連研究

2.1 情報推薦手法

情報推薦を行うためにはユーザやアイテムの特性をモデリングし、その結果に基づき推薦対象となるアイテムをフィルタリングする必要がある。既存の情報推薦手法の多くは協調フィルタリング (Collaborative Filtering) と内容ベースフィルタリング (Content-Based

*連絡先: 首都大学東京大学院
システムデザイン研究科情報通信システム学域
〒191-0065 東京都日野市旭が丘 6-6
E-mail: shattori@krectmt3.sd.tmu.ac.jp

Filtering)に分類することができる [3]. それぞれの手法の特徴について以下に述べる.

2.1.1 協調フィルタリング

協調フィルタリングは多くのユーザの嗜好情報を過去の行動という形で記録し, そのユーザと嗜好の類似した他のユーザの嗜好情報を用いてユーザの嗜好を推測する手法である [1]. 協調フィルタリングの利点は, アイテムの属性情報がなくても推薦が行えること, および処理が手軽であることであり, これらの理由からショッピングサイトなどで現在最も広く利用されている手法である. 一般に, 協調フィルタリングにより精度の高い推薦を行うためには, 多数のユーザに情報推薦システムが利用され, 多くのアイテムに関する行動情報が収集可能であることが必要となる. そのため, 推薦システムを新たに利用し始めたユーザや新規に追加されたアイテムに対しては行動情報の少なさから類似する嗜好を持つユーザを発見できず, 推薦の精度が低くなってしまうという欠点がある. これは cold-start 問題 [2] と呼ばれる. また, 推薦対象として膨大なアイテムが存在し, その多くにおいてユーザとの関係が希薄である場合, ユーザに関する情報が十分確保されていても精度の高い推薦を行うことは困難である. これは sparsity 問題 [4] と呼ばれ, cold-start 問題と併せて協調フィルタリングを用いた情報推薦手法全般に共通する課題とされている.

協調フィルタリングにおいて, cold-start 問題および sparsity 問題に関する研究は広く行われている. 代表的な手法として, 嗜好パターンが類似するユーザをモデル化することで精度の向上を実現するモデルベース法が挙げられ, Breese らはクラスタモデルを用いてモデルベース法を実装している [5]. モデルベース法以外のアプローチもいくつか行われており, Park らはユーザに加えて filterbot と呼ばれるロボットがアイテムへの評価を行うことで, ユーザ・アイテムの情報が少ない状態でも推薦に必要な情報を収集可能にする手法を提案している [6]. Lee らはユーザ同士の友人関係を類似度として利用することで, ユーザのアイテムの関係が希薄になる状態の改善を試みている [4]. また, Yildirim らはランダムウォークを用いてユーザ-アイテム間の類似度を求め, アイテムの推薦を行う手法を提案している [7]. これらの手法のように, 推薦システムにおけるユーザの行動情報の少なさを他ユーザの情報や仮想的な情報で補うことにより cold-start 問題および sparsity 問題の解決を目指すアプローチが主流となっている.

2.1.2 内容ベースフィルタリング

内容ベースフィルタリングはアイテムの内容とユーザの嗜好情報を比較し, その関連度に基づいてフィルタリングを行う手法である [8]. アイテムの内容には評価のポイントとなる属性の値が用いられ, 本稿ではこれを属性値と呼ぶ. 例えば映画では監督や俳優の名前, ジャンル (アクションやコメディなど) が属性値となる. 内容に基づくフィルタリングは協調フィルタリングと比べ, システムを使い始めたばかりのユーザでも特定の属性値に対する嗜好情報が得られれば精度の高い推薦が可能であるという利点があり, 楽曲推薦などで活用されている [9]. しかし, ショッピングサイトなど多様なジャンル・アイテムを取り扱う際には, アイテムの属性値は膨大なパターンが存在するため, ユーザとアイテムの属性値との関係が希薄になってしまうケースも多いと考える. そのため, 多様なジャンル・アイテムを推薦対象とした場合, 多くの属性値に対して推論を行うのに十分な情報を集めることは困難であり, 協調フィルタリング同様に cold-start 問題および sparsity 問題が課題となる.

内容ベースフィルタリングにおいてこれらの問題に取り組んでいる研究として, 関らのコンテキストを考慮した飲食店推薦システムが挙げられる [10]. 情報推薦のための行動情報取得はシステムの利用ログなどから自動的に行動情報を収集する暗黙的手法と, アイテムや属性値に対する好き嫌いをユーザに直接回答してもらう明示的手法の2つに分類することができる [11]. 暗黙的手法はユーザの負担が低いというメリットがあるが, 対象ユーザの行動情報を十分収集する必要がある. これに対し, 関らは事前に蓄積されたアイテムのコンテキスト情報から, 求めるコンテキストをユーザが明示的に指定することで新規ユーザでも適切な推薦結果が得られるとしている. しかし, アイテムに関しては事前に十分な量の行動情報を獲得する必要があることから, 新規アイテムを適切な推薦対象として扱うことは難しいと考える. ユーザ・アイテム双方に対して cold-start 問題・sparsity 問題を解決するためには, 少量の行動情報から推薦アイテムを獲得する手法が必要と考える. そこで本稿では, 次節に述べる価値観に着目する.

2.2 価値観に基づく嗜好・消費行動の推論

価値観は消費者の嗜好や行動に強く影響を及ぼすと考えられており, マーケティングの分野では古くから利用されている. Rokeach は消費者の嗜好に関わる価値観を 18 の要素に分類した Rokeach Value Survey [12] と呼ばれる調査方法を提案し, 多くの調査で利用されている. Vinson らは, 保守的な価値観を持つ大学と革

新たな価値観を持つ大学、それぞれに所属する学生の間に有意な嗜好の差があることをアンケート調査により明らかにしている [13]。近年でも、Holbrook が消費・購買行動に影響を与える価値観を 8 つに分類する [14] など、消費者の嗜好と価値観は関連の深いテーマとして研究および調査が進められている。Web インテリジェンスの分野においても、価値観はユーザの嗜好と関連の深い要素として利用されている。宮尾らはアンケートによりユーザの価値観を調査し、ユーザが持つ価値観に最適化された機能を持つ SNS プロトタイプを構築している [15]。また、Hattori らは、ユーザの嗜好と価値観の関連をアンケートにより調査し、情報推薦への適用可能性について考察している [16]。

従来の内容ベースフィルタリングでは、例えば映画であればジャンルや出演俳優など、アイテムの属性値に対する好き嫌いを購買情報などから暗黙的に取得し推薦アイテムを決定するアプローチが一般的である。しかし、ショッピングサイトなど多様なアイテムを推薦対象として取り扱う場合は、前節に述べたように新規ユーザに対して推薦を行うのに十分な情報を集めることは困難である。また、あるアイテムを好むからといって、アイテムの内容を示す属性値を全て好むとは限らず、従来手法では属性値に対する評価とアイテムに対する評価の関係は考慮されていない。そこで本稿では、価値観に基づいてアイテムへの評価に強い影響を与える属性を推論することで、ユーザの価値判断に合致するアイテムをより短期間で取得することを目指す。

3 価値観に基づく情報推薦システム

本節ではユーザのこだわりに着目した情報推薦システムの概要について述べる。3.1 節ではユーザのこだわりをモデリングする手法について、3.2 節では情報推薦システムの構成についてそれぞれ述べる。

3.1 ユーザモデリング手法

本節では、価値観と繋がり深い要素としてユーザの「こだわり」に着目したモデリング手法について述べる。本稿では、情報推薦における価値観は各属性に対する「こだわり」の強さとして表れると考える。そこで、ユーザの属性に対するこだわりを、図 1 に示すように「どの属性を重視してアイテムの評価を決定するか」を判断するための基準として定義する。

提案手法が従来の内容ベースフィルタリングと異なる点は、以下に示す 3 点に分類できる。

属性の利用： 属性値である著者の名前やアクションなどのジャンル名の代わりに、「著者」「ジャンル」

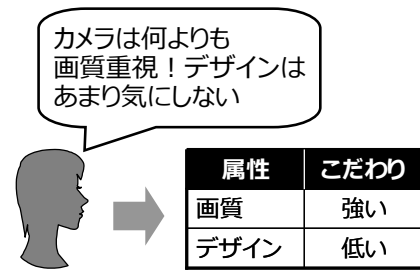


図 1: 属性に対するこだわりの例

といった属性をモデリングに用いる。内容ベースフィルタリングで多様なアイテムを取り扱う場合、大量に存在する属性値に対して十分な量の評価を収集することは困難であり、特に新規ユーザにおいて属性値との関係が希薄になってしまうと考えられる。提案手法では属性を推論に用いることで、新規ユーザに対しても各属性に対して推薦に必要な量の評価を収集することができると考える。

評価に与える影響度を推論： アイテムの内容に関する好き嫌いではなくアイテムの評価に与える影響度を推論する。内容ベースフィルタリングでは暗黙的に評価を収集することが多いが、あるアイテムを好むからといってそのアイテムの属性値を全て好むとは限らず、従来手法では属性値に対する評価とアイテムに対する評価の関係は考慮されていない。提案手法では、価値観に基づいてアイテムへの評価に強い影響を与える属性を推論することで、ユーザの価値判断に合致するアイテムをより短期間で取得可能になると考える。

評価の明示的な収集： 属性に対する評価を明示的に収集する。従来の内容ベースフィルタリングでは上記で述べたように暗黙的手法を用いることから属性値とアイテムへの評価の関係は考慮されないが、従来手法で明示的に情報収集を行う場合、大量の属性値が存在することからユーザにとって負担が大きい。提案手法では属性値に対して少数の属性を用いることから、ユーザに大きな負担をかけることなく情報を収集することが可能になると考える。

以上に述べたように、提案手法ではユーザがどのアイテムや属性値を好むかどうかではなく、属性に対する評価とアイテムに対する評価の関係を明示的に収集し推論に利用する。このような手法によって、どの属性がアイテムへの評価に強く影響を与えたか、つまりユーザはどの属性に強いこだわりを持っているかを推論することができると考える。

表 1: アイテムへの評価例

(1) アイテムAへの評価 (2) アイテムBへの評価

属性	極性	属性	極性
総合評価	好評	総合評価	不評
デザイン	不評	デザイン	好評
画質	不評	画質	不評
操作性	好評	操作性	不評
バッテリー	好評	バッテリー	不評

表 2: 評価一致率の計算例

属性	一致	不一致	評価一致率
デザイン	0	2	0.00
画質	1	1	0.50
操作性	2	0	1.00
バッテリー	2	0	1.00

本稿では、アイテムに対する評価極性（好評または不評）に加えて各属性に対する評価極性を抽出して分析を行う。アイテムに対する評価極性および各属性に対する評価極性を用いて、属性毎にアイテムの評価に与える影響度を推論しユーザモデルを作成する。提案手法ではこの影響度を評価一致率と呼ぶ指標を定義して表す。ユーザ u がアイテム i に対して行った評価 $e_{ui} \in E_u$ において、あるアイテム i の極性 $p_{item}(u, i)$ 、および属性 j の極性 $p_{attr}(u, i, j)$ が一致するかどうかを調べ、一致する評価の回数（アイテムの個数）を $O(u, j)$ 、一致しない回数を $Q(u, j)$ とする。この時、アイテム i における属性 j の評価一致率 $P(u, j)$ は式 (1) で算出される。これにより、ユーザのこだわりを表すユーザモデルは属性数を m とすると m 次元のベクトルとして表される。

$$P(u, j) = \frac{O(u, j)}{O(u, j) + Q(u, j)} \quad (1)$$

提案するユーザモデルでは、あるユーザが行った評価からそれぞれの属性に対する評価一致率を計算し、属性ごとに保持する。例として、ユーザがある 2 種類のデジタルカメラに対して評価を行った結果を表 1 に示す。また、表 1 の評価に基づいて属性毎に評価一致率を計算した例を表 2 に示す。表 2 の計算例に示す属性「操作性」「バッテリー」のような評価一致率が高い属性はユーザが強いこだわりをもっており、アイテムの評価に影響を与える「推薦時に重要度の高い属性」とであると推論される。一方で、表 2 の「デザイン」「画質」のような評価一致率の低い属性はアイテムの評価にそ

れほど影響を及ぼさず、「推薦時に重要度の低い属性」とであると推論される。

3.2 情報推薦システム概要

3.1 節で述べた手法を用いて推論したユーザのこだわりに基づく情報推薦システムの構成図を図 2 に示す。本システムは、以下に示す 3 つのモジュールから構成される。

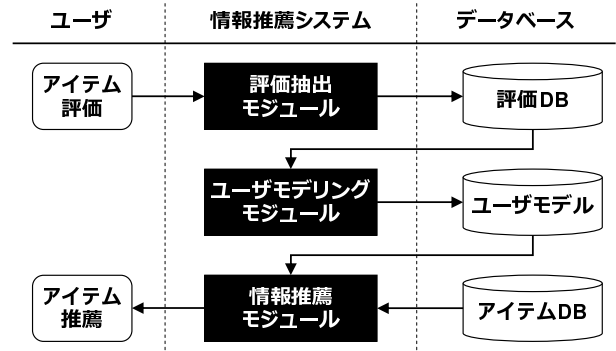


図 2: 情報推薦システム構成図

評価抽出モジュール： システム上でユーザが行った評価（アイテムの総合評価および各属性への評価）を取得し、評価 DB に保存する。評価 DB はユーザ毎に作成され、アイテム毎に付与された属性に対する評価極性を保持する。

ユーザモデリングモジュール： 評価抽出モジュールが抽出したユーザのアイテムに対する評価情報を用いて 3.1 節で述べた手法に基づきユーザのこだわりをモデリングし、各属性毎に評価一致率を求める。その結果をユーザモデルとして DB 上に保存する。

情報推薦モジュール： ユーザモデリングモジュールで作成されたユーザモデルを用いて、推薦するアイテムをユーザに提示する。アイテム DB からユーザが高いこだわりを持つ属性において高い評価が付けられているアイテムを検索し、推薦候補として抽出する。

推薦候補となるアイテムおよびシステムで対象とする属性はレビューサイト「価格.com¹」などの既存サービスを情報源とし、取得した商品情報をアイテム DB としてシステム上に保持する。価格.com の場合は図 3 に示すように、商品ジャンル毎にあらかじめ決められた属性を対象としてレビュー投稿が行われており、こ

¹<http://kakaku.com/>

れを用いることでアイテムの属性を容易に取得できる。しかし、ユーザが持つこだわりは多様であり、ここであらかじめ用意されている属性だけに留まらない可能性が高い。この場合、商品レビューからも属性を取得することで、より多様な属性を取り扱うことが可能と考える [17]。

また、既存の推薦手法の多くは推薦精度の向上を目的としており、ユーザの嗜好に近いアイテムや似たような嗜好を持つユーザが好むアイテムを推薦対象として扱っている。しかし、推薦されたアイテムはユーザにとって既知のものであることが多く、満足な推薦結果を得ることができない場合があることが指摘されている [18]。提案手法ではユーザが持つこだわりに基づいて推論を行うが、特定の属性に対して強いこだわりを持つユーザは新たな着眼点に触れることによって思考に飛躍が生じ、従来とは異なる評価によって意思決定が行われる可能性がある。この現象は認知科学分野において“mental leap”と呼ばれる [19]。特定の属性に強いこだわりを持っている状態は「思考のはまりこみ」状態にあるとも言え、そのようなユーザに対して新たな着眼点を提示することは、ユーザがその状態から脱する機会となり得る。このような新たな観点を提案手法を用いて導入することにより、ユーザにとって意外性があり、かつ真に価値のある推薦戦略の実現が期待できると考える。

デザイン	★★★★★ 4
発色・明るさ	★★★★★ 4
シャープさ	★★★★★ 5
調整機能	★★★★★ 5
応答性能	★★★★★ 3
視野角	★★★★★ 4
サイズ	★★★★★ 5
満足度	☆☆☆☆☆ 4

図 3: 価格.com の属性別レビュー

4 予備実験：評判情報からのユーザモデル作成

4.1 概要

提案システムの主要コンポーネントの一つであるユーザモデリングモジュールについて、予備実験を行った結果について示す。本実験では価格.com にレビュー（本サイトではクチコミと表記）を投稿しているユーザを対象とし、式 (1) を用いて各属性における評価一致率

を求めユーザモデルを作成する。価格.com ではユーザ投票により、レビューを投稿しているユーザをランキング形式でジャンル毎にまとめている。ジャンル「デジタル一眼カメラ」および「デジタルカメラ」においてそれぞれ上位 50 位にランキングされているユーザのうち、2012 年 5 月 13 日時点で過去 1 年間に 5 件以上のレビューを書いているユーザ 78 名、計 942 件のレビューをモデリングの対象とする。

提案手法ではアイテム・属性に対する評価極性を用いるが、価格.com において評価は星による 5 段階で表される。そこで、本実験では 5 段階からなる評価値およびその平均値との差を用いて評価極性を判定する。あるユーザ u が評価した n 個のアイテムに対する評価値の平均値 $\mu_{item}(u)$ を下記に示す式 (1) により求める。ここで、 E_{ui} はアイテム i に対するユーザ u の評価値を表す。あるアイテム i において、評価値 E_{ui} が $\mu_{item}(u)$ より高ければその評価極性 $p_{item}(u, i)$ を好評、低ければ不評と判定する。

$$\mu_{item}(u) = \frac{1}{n} \sum_{i=1}^n E_{ui} \quad (2)$$

また、アイテム i において、全ての属性（総数 m ）に対しユーザ u が与えた評価値の平均値 $\mu_{attr}(u, i)$ を、下記に示す式 (2) により求める。 $\mu_{item}(u)$ と異なり、 $\mu_{attr}(u, i)$ についてはアイテム単位で判断する。あるアイテム i の属性 $attr_j$ に対する、評価値 E_{uij} が $\mu_{attr}(u, i)$ より高い値を持つ属性の評価極性 $p_{attr}(u, i, j)$ を好評、低ければ不評と判定する。

$$\mu_{attr}(u, i) = \frac{1}{m} \sum_{j=1}^m E_{uij} \quad (3)$$

4.2 ユーザモデル作成

ジャンル「デジタル一眼カメラ」および「デジタルカメラ」における各ユーザについて、それぞれ属性毎の評価一致率を求め、ユーザモデルを作成した。作成されたユーザモデルの一部を例として表 3 に示す。表 4 は、評価一致率が 0.70 または 0.80 以上である属性が、全ユーザモデル（78 名）中にどれだけ存在するかをまとめたものである。この表から、「操作性」「携帯性」に強いこだわりを持つユーザは比較的少ないが、それ以外の属性はアイテムの評価に影響を与えることが多いことがわかる。また、ユーザモデル中にどれだけ強いこだわりを持つ属性が存在するかをヒストグラムとしてまとめたものが図 4 である。0.7 以上の評価一致率を持つ属性の数は 1 ユーザあたり平均で 2.60、0.8 以上の場合は平均で 1.68 であった。表 4 において、評価一致

表 3: 作成されたユーザモデル

ユーザ	属性							
	デザイン	画質	操作性	バッテリー	携帯性	機能性	液晶	ホールド感
user1	0.61	0.61	0.68	0.35	0.58	0.58	0.42	0.71
user2	0.70	0.70	0.50	0.70	0.80	0.90	0.70	0.80
user3	1.00	1.00	0.43	0.71	0.86	0.57	1.00	0.14
...	...	...	...	...	...	...	...	...
user78	0.69	0.63	0.88	0.38	0.75	0.75	0.69	0.56

率の高いユーザ数が同様の属性が多いことと併せて考察すると、ユーザ毎にこだわりを持つ属性はそれぞれ異なっていることがわかる。この結果より、ユーザそれぞれ異なるこだわりを反映したユーザモデルの作成が可能であると考えられる。

また、各ユーザがこだわりをもつ属性を推論するだけでなく、全ての属性をバランス良く評価するユーザや特定の属性に強いこだわりを持つユーザなど、ユーザのこだわりをいくつかのタイプに分類することも可能と考える。こだわりの推論に加えてその評価傾向を踏まえたモデリングを行うことができれば、推薦アイテムの選択だけでなくその推薦戦略決定に対しても提案手法を適用可能になり、ユーザの価値判断をより反映させた推薦が期待できる。

表 4: 各属性に対して強いこだわりを持つユーザ数

属性	0.70以上	0.80以上
デザイン	29	19
画質	41	25
操作性	19	10
バッテリー	22	14
携帯性	18	12
機能性	22	19
液晶	20	13
ホールド感	32	19

5 おわりに

本稿では、ユーザがどの属性を重視してアイテムへの評価を決定するかという属性毎の価値判断、いわばユーザの「こだわり」を推論するユーザモデリング手法、およびそれを用いた情報推薦システムを提案し、従来手法との特性の違いについて考察した。また、予備実験として価格.com のレビュー（クチコミ）を対象と

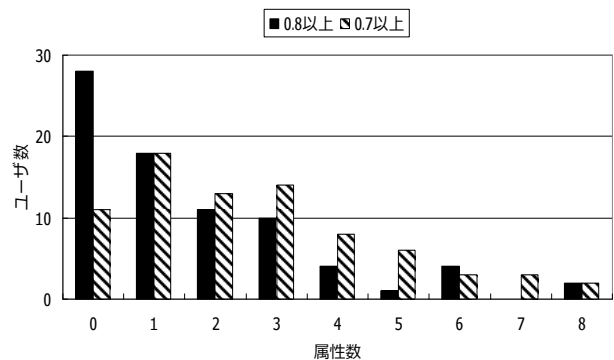


図 4: ユーザが強いこだわりを持つ属性数の分布

してユーザモデルを作成し、その実現可能性および推薦戦略決定への適用可能性について考察した。

今後は提案した情報推薦システムの実装を進め、従来手法との比較実験により提案手法の有用性を検証する。本稿では価格.com においてあらかじめ用意された属性を用いてユーザモデリングを行ったが、ユーザが持つこだわりは多様であり、対象とした属性だけに留まらない可能性が高い。楽天データ公開²ではユーザが投稿した商品に対するレビュー文を公開しており、これを情報源として自然言語処理技術によりレビュー文から属性およびその評価を抽出することで、より多様な属性をモデリングの対象として扱えることが期待できる。加えて、ユーザの価値観をモデリングする手法として、ユーザのこだわりに加えてその評価傾向についても分析を行う。評価傾向に基づきユーザをいくつかのタイプに分類することで、提案手法をアイテムの推薦戦略決定などにも適用していくことを目指す。

参考文献

- [1] P. Resnick, N. Iacovou, M. Suchak, P. Bergstrom, and J. Riedl, “GroupLens: An Open Architec-

²<http://rit.rakuten.co.jp/rdr/>

- ture for Collaborative Filtering of Netnews,” Proceedings of ACM 1994 Conference on Computer Supported Cooperative Work, pp.175-186, 1994.
- [2] A. I. Schein, A. Popescul, L. H. Ungar and D. M. Pennock, “Methods and metrics for cold-start recommendations,” 25th Annual ACM SIGIR Conference, pp. 253-260, 2002.
- [3] D. Riecken, “Personalized Views of Personalization,” Communications of the ACM, Vol. 43, No. 8, pp. 26-28, 2000.
- [4] S. Lee, J. Yang and S. Y. Park, “Discovery of Hidden Similarity on Collaborative Filtering to Overcome Sparsity Problem,” Proceedings of discovery science: 7th international conference, DS 2004, (LNAI 3245), pp. 396-402, 2004.
- [5] J. S. Breese, D. Heckerman, and C. Kadie, “Analysis of Predictive Algorithms for Collaborative Filtering,” Uncertainty in Artificial Intelligence 14, pp. 437-52, 1998.
- [6] S. T. Park, D. Pennock, O. Madani, N. Good and D. DeCoste, “Naive filterbots for robust cold-start recommendations,” KDD '06 Proceedings of the 12th ACM SIGKDD international conference on Knowledge discovery and data mining, pp. 699-705, 2006.
- [7] H. Yildirim and M. S. Krishnamoorthy, “A random walk method for alleviating the sparsity problem in collaborative filtering,” RecSys '08 Proceedings of the 2008 ACM conference on Recommender systems, pp. 131-138, 2008.
- [8] F. Pachet, P. Roy nad D. Cazaly, “A Combinatorial Approach to Content-based Music Selection,” Proceedings of IEEE Multimedia Computing and Systems International Conference 1999, pp. 457-462, 1999. IEEE Multimedia, vol. 7, pp. 44-51, 2000.
- [9] 吉井和佳, 後藤真孝, 音楽推薦システム, 情報処理, 50 巻 8 号, pp. 751-755, 2009.
- [10] 関匠吾, 中島伸介, 張建偉, アイテム利用時のユーザコンテキストを考慮した情報推薦システムの提案, 第 3 回データ工学と情報マネジメントに関するフォーラム (DEIM 2011) , B1-1, 2011.
- [11] 神島 敏弘, 推薦システムのアルゴリズム (2), 人工知能学会誌 23 巻 1 号, pp.89-103, 2008.
- [12] M. Rokeach, “The Nature of Human Values,” New York: The Free Press, 1973.
- [13] D. E. Vinson, J. E. Scott, and L. M. Lamont, “The role of personal values in marketing and consumer behavior,” The Journal of Marketing, Vol. 41, No. 2, pp. 44-50, 1977.
- [14] M. B. Holbrook, “Consumer value: a framework for analysis and research,” Routledge, 1999.
- [15] 宮尾 和樹, 原田 利宣, SNS サイトの分類とユーザの価値観に基づくプロトタイプの構築, デザイン学研究, 55 巻 1 号, pp.81-90, 2008.
- [16] S. Hattori and Y. Takama, “Investigation about Applicability of Personal Values for Recommender System,” Journal of Advanced Computational Intelligence and Intelligent Informatics (JACIII), Vol. 16, No. 3, pp.404-411, 2012.
- [17] 服部 俊一, 毛 中杰, 高間 康史, 価値観に基づく情報推薦のためのレビュー分析手法の提案, 第 26 回人工知能学会全国大会, 3K2-NFC-3-6, 2012.
- [18] C. N. Ziegler, S. M. McNee, J. A. Konstan, and G. Lausen, “Improving Recommendation Lists Through Topic Diversification,” WWW '05 Proceedings of the 14th international conference on World Wide Web, pp.22-32, 2005.
- [19] 藤本 和則, 松下 光範, 本村 陽一, 庄司 裕子, 意思決定支援とネットビジネス, オーム社, pp.134-140, 2005.

電子書籍を対象とした視覚的スタイル自動付与システムの検討

A study of automatic visual styling system for digital books

渡邊祥太^{1*} 大野 将樹¹ 沼尾 雅之¹
Shouta Watanabe¹ Masaki Oono¹ Masayuki Numao¹

¹ 電気通信大学大学院 情報理工学研究科

¹ Graduate School of Informatics and Engineering, The University of Electro-Communications

Abstract: In the advancement of the digital book technology, the desire that we want to read books more happily and convenient has risen. Actually, the digital book reader that we can freely specify the background color etc. appears. However, a current visual style values only the easiness of visibility, and doesn't consider the content of the novel. They do not help us to understand the contents. In the present study, we propose the automatic visual styling system used for the digital book. The system performs (1)Scene division, (2)Clustering of similar scene and (3)Mood presumption of each cluster. The result is used for visual style of novel. In this paper, we evaluated the first (1) Scene division based on time, place, and person information. By the experiment, the result of F value 0.216(0.413 correct answer in one sentence before and after) was obtained.

1 はじめに

近年、ニュースなどの紙媒体の電子化、電子書籍の普及が進みつつある。iPad 等の電子リーダーの登場によりそれはよりパーソナルなものとなり、個人が楽しむ書籍である小説の電子化も進んでいる。

その中で、小説の電子データを利用した研究も盛んにおこなわれている。馬場らは、登場人物に基づいた小説のモデル化を行っている [1]。ストーリーや登場人物に基づいて検索や分類される小説に対し、索引語に基づくモデル化だけでは限界があると考え、テキストのモデル化を目指している。英文学の推理小説を対象に、テキストから自動的に人物相関図を作成し、それを可視化している。星川は、形態素解析システム「茶筌」を用い、若い作者と文学史上の作者の文章を比較し、その中で若い作者の文章の特徴を捉える研究を行っている [2]。

電子化により、人が小説を読むスタイル自体も変化してきている。電子書籍では、テキストのみの編集だけでなく、そのフォントや段落なども自由に変えることができる。そのため読者は「より楽しく」「よりわかりやすい」読書を楽しみたいという欲求を持ち、電子書籍の視覚的スタイルへの興味も高まっている。実際に、文字のフォントを変更する事ができたり、背景色を変えたりできる機能が付いたリーダーなども既に登場してきている。この読者側の欲求に対し、編集者、

作者側でも今後さらに増加していく電子書籍に対して「より手軽に視覚的スタイルを作成したい」という欲求が高まっていると考える。しかし現状の視覚的スタイルは視認のし易さのみを重視しており、内容理解を補助しているわけではない。小説の内容自体を考慮していないため、それぞれの小説独自の視覚的スタイルを提供しているわけではない。

そこで本研究は、電子書籍を対象とした視覚的スタイル自動付与システムを提案する。提案システムとしては (1) 場面分割、(2) 類似場面のクラスタリング、(3) 各クラスタのムード推定を行う。この3ステップから、小説の内部情報を考慮した視覚的スタイルを作成する。

2 関連研究

小説を対象とした場面分割の研究として、時・場・人に基づいた場面分割を行っている研究がある [3]。この研究では、場所、時間、登場人物に着目し、それが変化した点を分割点と定義し、物語のシーン分割を試みている。形態素解析を用い、場所、時間、登場人物の候補となる語句の抽出を行い、抽出した語句を用いてシーン境界推定実験を行っている。シーンの開始文から場所候補、時間候補、登場人物候補の語句を各々のプール(二文分)に蓄積し、新しい入力文毎に、各プールに含まれる語句との異なり数に基づきペナルティを計算している。このペナルティが閾値より大きくなった時、新しい入力文がシーンの境界であると推定して

*連絡先: 電気通信大学大学院 情報理工学研究科
〒182-0021 東京都調布市調布ヶ丘1丁目5-1
E-mail: w1131119@edu.cc.uec.ac.jp

いる。ペナルティは時間単語、場所単語、登場人物単語の異なり数を用い、重みをつけ、総合的に決定されている。本研究では、場面区切り候補点として [3] におけるシーン分割点である「登場人物が変化した点」「場所が変化した点」「時間が変化した点」を利用する。そこで場面区切り候補として、時間、場所、人変化点をそれぞれ独立して抽出する。プールする文の数や、閾値、単語抽出の際のルールなどを変化させ実験を行う。

テキストを分割する手法としては、TextTiling 法 [4] がある。基準点 t から同数の語を含むように左右に窓を設け、その窓内での語彙の重なり尺度から類似度を算出し、それにより分割を行っている。

語彙的結束性に基づき、テキストを場面に分割する研究に [5] がある。LCP という統計的な指標を提案し、場面分割を行っている。LCP は、テキストの各位置について、その近傍の単語列の結束度を記録したものである。LCP をグラフとして表し、その極小値が分割点としている。

テキストと視覚的メディアを融合させる研究に [6] がある。俳句のテキストを分析し、特徴付ける語を抽出し、そこから Web 画像検索、画像解析を行い、俳句と同時に画像を表示するシステムに関する研究である。テキストデータのムードに基づき、他のメディアを表示させるという点で関係性が高い。本研究では、よりストーリー変化が多い小説を題材とし、場面分割、クラスタリングを経て、画像を含んだ視覚的スタイルを同時表示させる。

3 提案システム

本研究では小説の三大要素として「時間、場所、人」があると考え、それを小説の内部情報とし、内部情報に基づいた視覚的スタイル作成を行う。「時間、場所、人」情報とは、どの時間帯なのか、場所はどこなのか、その場にいる登場人物は誰なのか、という情報である。具体的には図 1 のような、場面分割、クラスタリング、ムード推定、視覚的スタイル作成というステップで行われる。

本研究では、(1) 場面が変化した事がわかる、(2) 類似場面が認識できる、(3) 場面のムードがわかる、という三つの目的を満たす視覚的スタイル作成を目指す。場面分割ステップでは、人が小説を読んでいる中で場面が変化したと感ずる点を見つける。小説の場面と区切りを定義し、場面区切り候補の抽出、そしてそれを用いた場面分割を行う。クラスタリングステップでは、場面分割ステップにより分割された場面を用い、類似場面のクラスタリングを行う。視覚的スタイルにおいて、類似場面には同じ「挿絵」が挿入される。ムード推定ステップでは、場面ムードを表す視覚的スタイル

に向け、各場面のムード推定を行う。ムードは喜怒哀楽の 4 指標で判定され、その値に応じ、視覚的スタイルにおける「背景色」が決定される。



図 1: システム

3.1 場面分割

小説の場面分割を行う。まず、小説の場面から場面が変わる瞬間とはどのようなものを定義する。次に、場面分割のための場面分割候補点として、場面区切り候補を抽出する。そしてその区切り候補を用い、場面分割点を抽出する。

3.1.1 場面の定義

場面は「時間」「場所」「人」という三大要素で構成される。その中で、場面分割点を以下のように定義した。

- 「時間」、「場所」が同時に変化したとき
- 「時間」、「人」が同時に変化したとき
- 主要登場人物が変化した時

本研究では、この 3 つのいずれかの変化が起こった際に、その変化点を場面分割点であると定義する。

小説において、「時間」は大きな役割を持っている。「時間」という要素の上に、「人」「場所」という要素がある事で場面は構成されている。東条は、時間軸上にあるスナップショット、いわゆる時点を可能世界と考え、時間軸とはそのような可能世界の連鎖であると考えている [7]。小説においては、時点の連続的固まりが場面である。小説で、違う時間帯へ飛ぶ時や、過去への回想に変化する時、大きく時間が進む時等は、時点の連続的固まりが非連続的に別の固まりへと移動した事を意味する。そこで「時間」の変化は、場面変化に大

きな意味を持っていると考えた。しかし、小説においては小さい「時間」変化も発生する。「3 時間後、気づくと夜が明けていた」等は、「時間」変化の一種であるが、小さすぎて小説の場面変化とは感じられない。時間が変化しても、その上に存在する「人」「場所」要素が変化しないと、小説において大きな変化だと感じる事はできない。そこで、本研究では「時間」とともに「人」が変化した場合、あるいは「場所」が変化した場合に、そこを場面分割点と定義した。

また「人」も小説における重要な要素である。小説には大きく分けて限定視点、全知視点、客観視点がある[8]。限定視点の小説は登場人物の視点からストーリーが進められ、全知視点、客観視点では神の視点(全員の心理がわかる視点)、客観視点ではカメラの視点(起こっている出来事のみが淡々と描かれる)からストーリーが進められる。限定視点において視点主が変化する事は、読者の視線を変化させることと同等であり、大きな変化といえる。また全知視点、客観視点においても、主人公が変化する事は、読者が注目していた人物が変わる事と同等で、大きな変化といえる。そのため視点主の変化や、物語で語られる主人公が変化するなど大きい「人」の変化も場面変化として考えた。しかし、小説には途中で誰かが登場するなどの、微小な人変化も存在する。そこで各場面において、主要登場人物を抽出し、その主要登場人物が変化する時に場面分割が起こると定義した。

3.1.2 場面区切り候補抽出

まずは小説テキスト全体から場面区切り候補を切り出す。視覚的スタイルを作成するうえで、この場面が最小単位となり、場面区切りによって視覚的スタイルが変更される。

- 時間：時間が不連続変化した点
- 場所：場所が変更される点
- 人：登場人物の増減が起こる文

を小説の場面区切り候補と考え、三種類の区切り候補の抽出を行う。小説テキストを改行ごとに区切り、1 文目、2 文目、3 文目とそれぞれの文章に番号を付ける。この文章が小説内の区切り候補の最小単位となる。次に予め定めた窓幅に従い、 t 文目付近の要素単語(時間、場所、人単語)をそれぞれ抽出し、単語ベクトル $V(t)$ を作成する。時間単語としては、日本語形態素解析システム JU-MAN により、カテゴリが時間と判定されるもの、場所単語としては、同じく JUMAN によりカテゴリが場所と判定されるもの、人単語としては、JUMAN によりカテゴリが人と判定されるもの、そし

て形態素解析 Mecab により人名として判定されるものを利用した。すべての t に対し、 $V(t)$ を作成し、隣り合う単語ベクトル ($V(t), V(t+1)$) のユークリッド距離を求め、その距離が閾値を超えた時 $t+1$ を区切り候補とする。時間単語により作成された $V(t)$ から得られた区切りを時間区切り候補、場所単語により作成された $V(t)$ から得られた区切りを場所区切り候補、人単語により作成された $V(t)$ から得られた区切りを人区切り候補とする。これにより、時間、場所、人がそれぞれ変化する場面区切り候補を得る事ができる。

3.1.3 場面候補の結合

上記で得られた、場面区切り候補を用い、定義された場面分割を行う。

- 「時間」「場所」が同時に変化したとき
- 「時間」「人」が同時に変化したとき

の切り出しにはそれぞれ時間区切り + 場所区切り、時間区切り + 人区切りを利用する。時間区切りの前後 t 文内に場所区切りがあった場合、その時間区切りを場面分割点とする。同様に、時間区切りの前後 t 文内に人区切りがあった場合、その時間区切りを場面分割点とする。

- 主要登場人物が変化した時

に対しては、区切り候補で区切られた各場面候補 n に対し、主要登場人物を推定する。

場面候補内において形態素解析を行い、人物、人名として判定されるものの中で、素文内に最も多く登場している単語を主要登場人物とする。もし場面候補 $n+1$ 内に、主要登場人物が推定できない場合は、その場面候補 $n+1$ の主要登場人物は、場面候補 n の主要登場人物と同一であると判断する。場面候補 n と場面候補 $n+1$ を比較し、主要登場人物が変化していた場合、場面候補 $n+1$ の最初の文を場面分割点とする。

3.2 クラスタリング

次に抽出した場面のクラスタリングを行う。本研究での視覚的スタイルでは、似ている場面には共通の視覚的スタイルを与える。これにより、読者が視覚的にわかりやすく読書を進めることができる。類似の指標としては、場面分割で利用した「時間」「場所」「人」要素を利用する。場面分割ステップにおいて、分割された場面に対しクラスタリングを行う。

方法としては、階層的クラスタリングを行う。分割された場面 t における、時間単語ベクトル $T(t)$ 、場所

単語ベクトル $S(t)$ 、人単語ベクトル $H(t)$ をそれぞれ作成する。全場面对し、 $T(t)$ 、 $S(t)$ 、 $H(t)$ からクラスタリングを行い似ている場面を一つのクラスタに入れる。これにより、類似している場面グループ群を得る事ができる。同一場面グループ群は、視覚的スタイルにおいて同一の挿絵が挿入される。

3.3 ムード推定

次に、視覚的スタイル作成に向け、場面のムード推定を行う。各場面对し、視覚的スタイルに利用するムードを付与する。ここでのムードとは各場面の感情属性であり、喜怒哀楽の4つの値で表現される。

ムードの判定には形容詞を利用した、喜怒哀楽辞書を利用する。各場面内の形容詞を、辞書を用い喜怒哀楽のどれに属するか判定する。これにより、各場面の喜怒哀楽に利用する4つの値が得られる。それぞれの値を、場面グループ群内の全文章数 N で割る事で、喜怒哀楽の4つの値が生成される。この値は各場面对し付与され、これが場面のムードとなる。

3.4 視覚的スタイル作成

最後に、類似場面クラスタリングの情報、ムード情報から視覚的スタイルの作成を行う。

クラスタリングで求められた、各場面グループ群に対し、挿絵を割り当てる。場面グループ群 t 内の全名詞から、最も多く登場する名詞をキーワードとして抽出し、Google AJAX Search API を利用し、挿絵画像を得る。場面グループ群を単位として行うため、類似場面には同一の挿絵が挿入される。

次にムード推定で求められた、喜怒哀楽4つの値を用い、背景色を決定する。喜、怒、哀、楽に対しそれぞれ、ピンク、赤、青、黄を割り当ててそれを基本色とする [12]。喜怒哀楽の4値から、それぞれの色を合成し、背景色を決定する。

ファイル形式としては ePub を出力する。ePub とは、HTML5, CSS3 によって構成される、電子書籍用ファイル・フォーマットである。Android、iOS ソニー・リーダー、楽天 kobo Touch 等の電子書籍リーダーや、Google Chrome、Mozilla Firefox などのブラウザで利用可能となっており、世界標準の規格化を進めている形式である。

入力として与えられた小説テキストと、ここまでのステップで得られた場面分割情報、ムード情報を、テキストを保持する HTML5、そのテキストのフォントや配置情報などを保持する CSS3 を介し、ePub へと変換する。

4 実験

システムのステップの一つである、場面分割における場面区切り候補の実験を行った。抽出する区切りは、時間区切り、場所区切り、人区切りの三種類である。各小説を改行ごとに区切り、区切られた文 t それぞれに対し、ベクトル $V(t)$ を作成し、 $V(t)$ と $V(t+1)$ のユークリッド距離を算出し、それが閾値を超えた時、 $t+1$ 文目を区切りとする。実験では、正解データと比較して算出された、 F 値の他、ベクトルや閾値などを変化させ、その中で最適な組み合わせも見つける。

4.1 準備

準備として、ベクトル作成のための窓幅、区切り判定の閾値、単語ベクトル、出力形式を以下のように設定した。全てのパターンで実験を行い、この中で最適な組み合わせを見つけた。

- 窓幅：7 種類
7、5、4(1)、4(2)、3、2(1)、2(2)
- 閾値：4 種類
平均値、4 分位数、8 分位数、10 分位数
- 単語ベクトル (時間、場所)：4 種類
全単語ベクトル、素文単語ベクトル、全単語次元削減ベクトル、素文単語次元削減ベクトル
- 単語ベクトル (人)：12 種類
人物 (上記 4 種類)、人名 (上記 4 種類)、人名 + 人物 (上記 4 種類)
- 出力形式：2 種類
通常、二分連続出力禁止

窓幅としては、奇数値 k のとき、 t 文目とその前後 $(k-1)/2$ 文を一つの窓とした。偶数値に対しては、4(1) は t 文目と前 1 文と後 2 文、4(2) は t 文目と前 2 文と後 2 文、2(1) は t 文目と $t+1$ 文目、2(2) は t 文目と $t-1$ 文目をそれぞれ窓とした。閾値はすべての文 t における、 $V(t)$ と $V(t+1)$ のユークリッド距離における、平均値、分位数を利用した。全単語ベクトルとは全要素単語から作成したベクトルであり、素文単語ベクトルとは鍵括弧内には不確かな情報が多いとの仮定から素文のみの要素単語を利用し作成したベクトルである。次元削減ベクトルとは、単語ベクトルの次元数が多くなりすぎてしまうという問題点から、一度しか登場していない要素単語を削除したベクトルとなっている。人ベクトルの人物とは、Juman において人を意味すると判定

された単語(医者、彼など)であり、人名とは Mecab に
おいて人名(前田、太郎など)であると判定された単語
である。また出力として、変化が大きい時に、連続し
て区切りを出力してしまうという問題点があったため、
通常出力の他に、二分連続出力を禁止した場合での実
験も行った。

4.2 実験データ

実験データとしては、電子文藝館 [9] から 10 個の小
説を用いた(表 1)。各小説に対し、時間区切り(場所が
不連続変化する文)、場所区切り(場所が変化する文)、
人区切り(登場人物が増減する文)を人手で付与し、こ
れを正解データとした。

各小説に対し、システムが出力した区切りと、人手
での正解データを比較し結果とした。「完全一致」の正
解率、「前後 1 文」内での正解率をそれぞれ F 値で求め
て評価した。「前後 1 文」で評価を行ったのは、視覚的
スタイル作成において区切りが完全一致する必要性は
そこまで高くないためである。

表 1: 実験に使用したデータ

No	タイトル	作者
1	把手のない扉	谷本多美子
2	潮流	穂高健一
3	幸福の彼方	林芙美子
4	紫の記憶	水樹 涼子
5	ロボットとベッドの重量	直木三十五
6	スターダスト・レビュー	浅田次郎
7	藪を這う	中山孝太郎
8	ドミノのお告げ	久坂 葉子
9	殺意の造型	森村誠一
10	鬼灯の女	小笠原 幹夫

4.3 実験結果

実験結果は表 2、表 3 ようになった。実験結果は表
2、表 3 ようになった。

5 考察

全体を通して適合率が低く、再現率が高いという結
果になった。これは全体的に出力が多く、過剰分割に
なっているからだと考えられる。過剰分割し
てしまった現段階の分割点から、更に正解を絞ってい
く必要があると考えられる。閾値は平均値が最も良
い傾向にあり、窓幅は 2(前 1 文) のものが良い結果に
なった。

表 2: 全条件で最もよかった F 値
完全一致 前後 1

小説	時間	場所	人	時間	場所	人
1	0.104	0.270	0.172	0.122	0.460	0.296
2	0.189	0.273	0.162	0.436	0.407	0.386
3	0.258	0.261	0.191	0.444	0.400	0.452
4	0.316	0.235	0.400	0.462	0.462	0.667
5	0.000	0.222	0.154	0.182	0.160	0.278
6	0.161	0.233	0.091	0.231	0.418	0.303
7	0.357	0.379	0.233	0.571	0.780	0.571
8	0.207	0.353	0.275	0.500	0.636	0.629
9	0.109	0.114	0.204	0.261	0.233	0.366
10	0.095	0.250	0.205	0.278	0.491	0.523
平均	0.180	0.259	0.209	0.349	0.445	0.447

表 3: F 値が高かった組み合わせ

完全一致	特徴ベクトル	閾値	出力	窓幅
時間	時間(素)	平均値	通常方式	2(2)
場所	場所 2(素)	平均値	2 連禁止	2(2)
人	人 2(素)	十分位	2 連禁止	2(2)
前後 1 文	特徴ベクトル	閾値	出力	窓幅
時間	時間(素)	四分位	2 連禁止	3
場所	場所(素)	平均値	2 連禁止	2(1)
人	人 2(素)	四分位	2 連禁止	3

5.1 小説

小説によって、結果に大きく差が出た。結果の悪い
小説を見てみると、時間、場所、人の要素単語が極端に
少ない傾向にあった。また「～は～へ行った。」といっ
た状況説明文が少ない事がわかった。逆に結果良い小説は

私は立ち上りました。そして自分の部屋へはいると
急に信二郎がかわいそうになって来ました。信二郎は
どんな風に生きるのか。私はやっぱり黙っているのが
いいのでしょうか(「ドミノのお告げ」より)
などの状況描写が細かい傾向にあった。状況描写が細
かくなると、場所単語や人物単語をシステムが拾いや
すくなるため、結果が総じて良くなったのだと考えら
れる。

5.2 時間、場所、人

時間、場所、人という分割点別でみると「時間」区切
りが見つけにくい事がわかった。これは時間単語の種
類が多く、直接時間と関係ないもの(「瞬間」「先」な
ど)も多く含まれているためだと考える。「その瞬間」

などの時間単語は、概念的には時間を意味しているものの、区切り変更には関係ないため、ゴミとなってしまった。逆に「場所」単語は、「学校」「公園」など単純な場所を示すものが多く、区切りの精度が高くなった。

5.3 特徴ベクトル

特徴ベクトルとしては、素文のみの単語ベクトルを使用した方が良い結果となった。これは

「いや、ちがう。私のかつて部下だった女性が、出身地の奥尻島で結婚して三年目で亡くなった。有能なひとだった。その墓参りなんだ」「大地震で犠牲に？」「そうじゃなくて、骨髄の病気で亡くなった。若いのに、血液のガンで。彼女が函館の病院に入院しているとき、見舞いに一度きたきりで、亡くなったとき、私はすでに海外勤務で、葬儀に参列できなかった。前々から、墓参りしたいとおもいながら、もう五年近くが経ってしまった」「(「潮流」より)

のような会話で過去の女性「彼女」の事を話しているシーンなどで、その場にはいない「彼女」に反応して分割してしまうためである。また会話文の途中で、区切りが変化する事が少ないことも関係していると考えられる。次元削減(一度のみの単語削除)をした特徴ベクトルは、人の区切りに効果があった。これは人に関連する単語の種類が多すぎたためと思われる。登場人物とは言えない人単語は、システムにおいてゴミになってしまう。一度しか登場していない人単語は、まさにゴミである可能性が高いため、それを削除できる次元削減ベクトルが特に効果的であったと考えられる。逆に時間単語は、

とうとう一睡も出来なかったらしくて夜が明けた。壁に掛かっている時計を見ると六時であった。男の顔を見るとぐっすり眠っている有様であった。(「藪を這う」)より)

のように一度しか登場しなくても重要な単語が多いため、次元削減における効果は低くなった。

6 まとめと今後

本研究では、電子書籍を対象とした視覚的スタイル自動付与システムの検討を行った。視覚的スタイルの作成に必要なステップとして、(1) 場面分割、(2) 類似場面のクラスタリング、(3) 各クラスターのムード推定を検討し、その中で場面分割のための時間区切り、場所区切り、人区切りを抽出する実験を行った。実験結果としては、完全一致のF値として時間: 0.180、場所: 0.259、人: 0.209、前後1文内での正解のF値として、時間: 0.349、場所: 0.445、人: 0.447の結果を得た。

今後は、場面区切り候補に対して、時間単語の選別やさらなる次元削減などを行うことで、より高い精度を目指していく。そして、得られた区切り候補から場面分割点を得るための実験を行っていく。また分割した場面を用い、次のステップである、類似場面をグループ化するクラスタリング、そしてベクトル情報を用いたムード推定、そして実際に視覚的スタイルをつけた電子ePub形式の書籍の作製を行っていきたい。

参考文献

- [1] 馬場こづえ, 藤井敦:「小説テキストを対象とした人物情報の抽出と体系化」言語処理学会第13回年次大会発表論文集 2007
- [2] 星川法子「形態素解析による若い作家の小説の特徴の研究」, SONODA 情報コミュニケーション, 2005
- [3] 小林聡:「場・時・人に着目した物語のシーン分割手法」情報処理学会自然言語処理研究会報告 2007
- [4] M. A. Hearst, "TextTiling: Segmenting Text into Multiparagraph Subtopic Passages", Computational Linguistics, vol. 23, no. 1, pp. 33-64 (1997).
- [5] 小嶋秀樹, 古郡廷治:「単語の結束性にもとづいてテキストを場面に分割する試み」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, 1993
- [6] 湊匡平, 尾内理紀夫:「俳句に適合する合成画像を生成するシステムの検討」映像情報メディア学会技術報告 2009
- [7] 東条 敏「言語・知識・信念の論理」
- [8] 「小説の人称と視点」,
http://www.cre.ne.jp/user/tenmyo/writing/person_and_viewpoint.html
- [9] 日本ペンクラブ電子文藝館
<http://bungeikan.org/domestic/>
- [10] 日本語形態素解析システム JU-MAN,
<http://nlp.kuee.kyoto-u.ac.jp/nl-resource/juman.html>
- [11] 形態素解析エンジン Mecab,
<http://mecab.sourceforge.net/>
- [12] 箕田妃希, 長幾朗「「空メール」色彩メールによる感情表現の試み」インタラクショナル 2003

検索対象としてのデスクトップイメージ画像

Desktop Images as Target of Information Retrieval

梅村恭司¹ 武並佳則² 折原幸治² 熊谷摩美子¹ 杉浦遼一¹ 若松翔¹ 月田晴貴¹

Kyoji Umemura¹, Yoshinori Takenami², Orihara Koji², Mamiko Kumagai¹, Roichi Sugiura¹,
Sho Wkamatsu¹ and Tukita Haruki¹

¹ 豊橋技術科学大学 情報・知能工学系

¹ Toyohashi University of Technology, Dept. Computer Science

² 住友電気情報システム ビジネスソリューション事業本部

² Sumitomo Electric Information Systems Co., Ltd.

Abstract: Recently, major information for human comes and go through computer, and recording the activity on computer can be regarded as a kind of life log of human. Unlike life log in video form, the life log of computer records are rich in text information. This makes it possible to handle the record more effectively than video. This report presents a prototype system for retrieval of the captured image of desktop of computer. This report also discusses various characteristics of captured image collection compared with text collection. Finally, this report presents some situation where the proposed system will be useful.

概要

人間の活動に関わる検索を考えたときに、コンピュータ画面を通じて情報を取り入れる比率は大きくなっており、人間の活動記録の検索の一つの例としてデスクトップイメージの検索を考えることができる。一方で、一般的なビデオ画像に比べて画像のなかに文字が多く含まれ、それによって検索できる利点がある。われわれは、これを具体的に示すために、検索のプロトタイプを作成している。本発表では、コンピュータの操作画面画像（デスクトップイメージ画像）の検索のプロトタイプシステムを通じて、検索対象としての特色と検索の意義を具体的に報告する。

1. はじめに

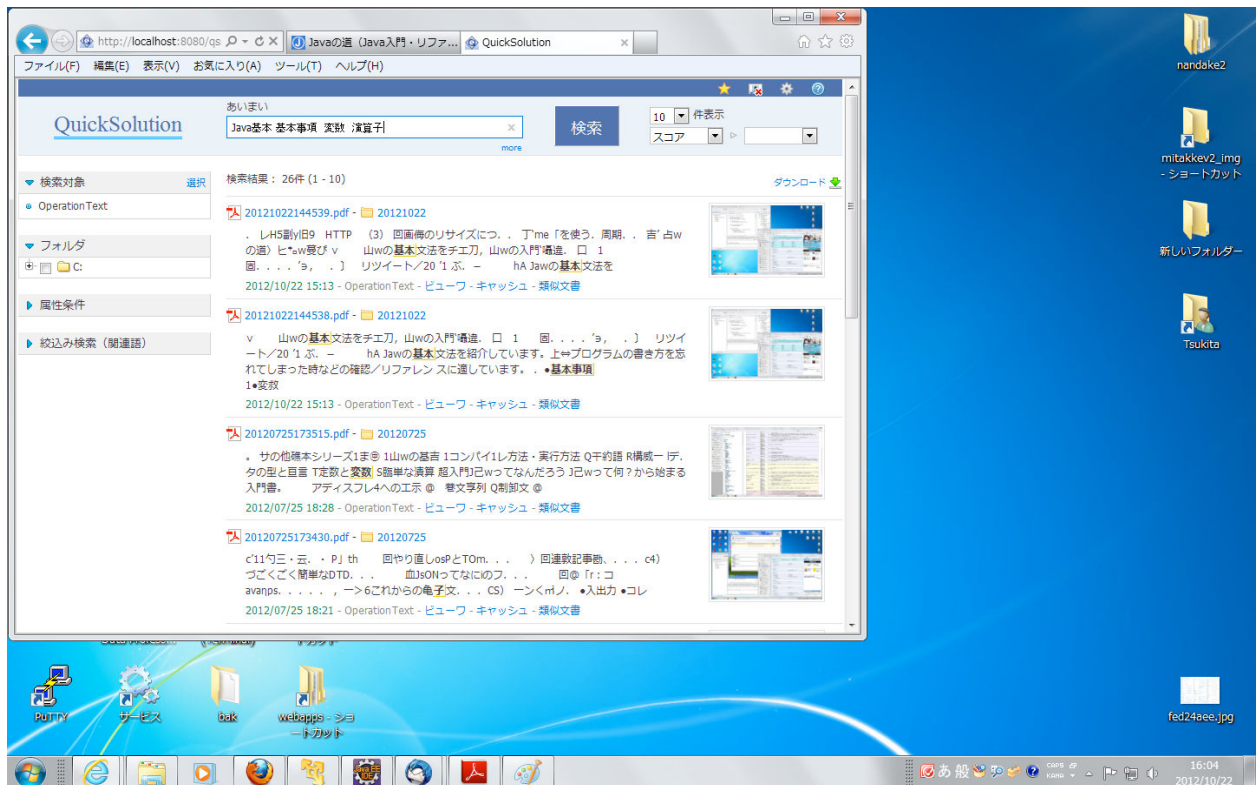
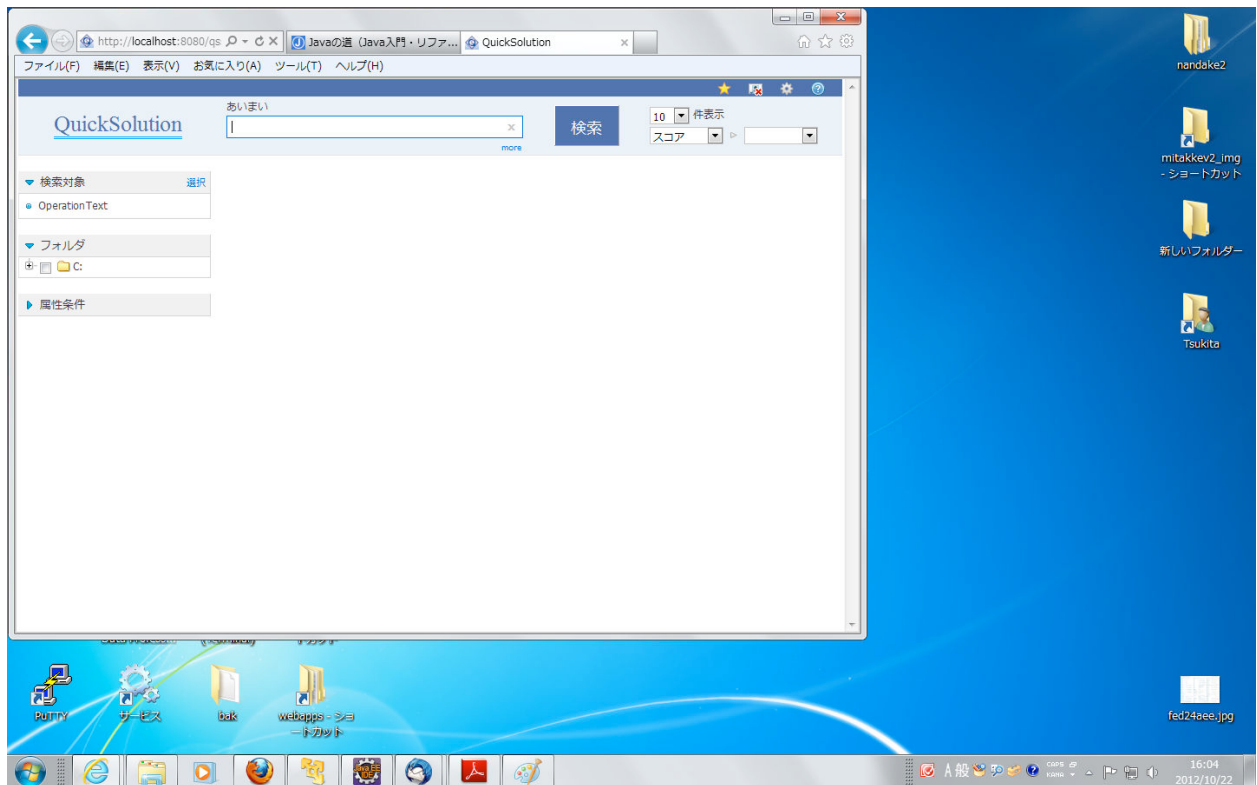
ポータブルデバイスとネットワーク発展により、人間が取得する全ての情報を記録して、活用することに注目が集まっている[1]。ここにおいて取得されるものは、人間の目がとらえる映像が主たる要素を占める。この情報は先頭から順番に調べていくのには大き過ぎるものであるので、検索のシステムが必要であるが、映像の検索には課題が多い。

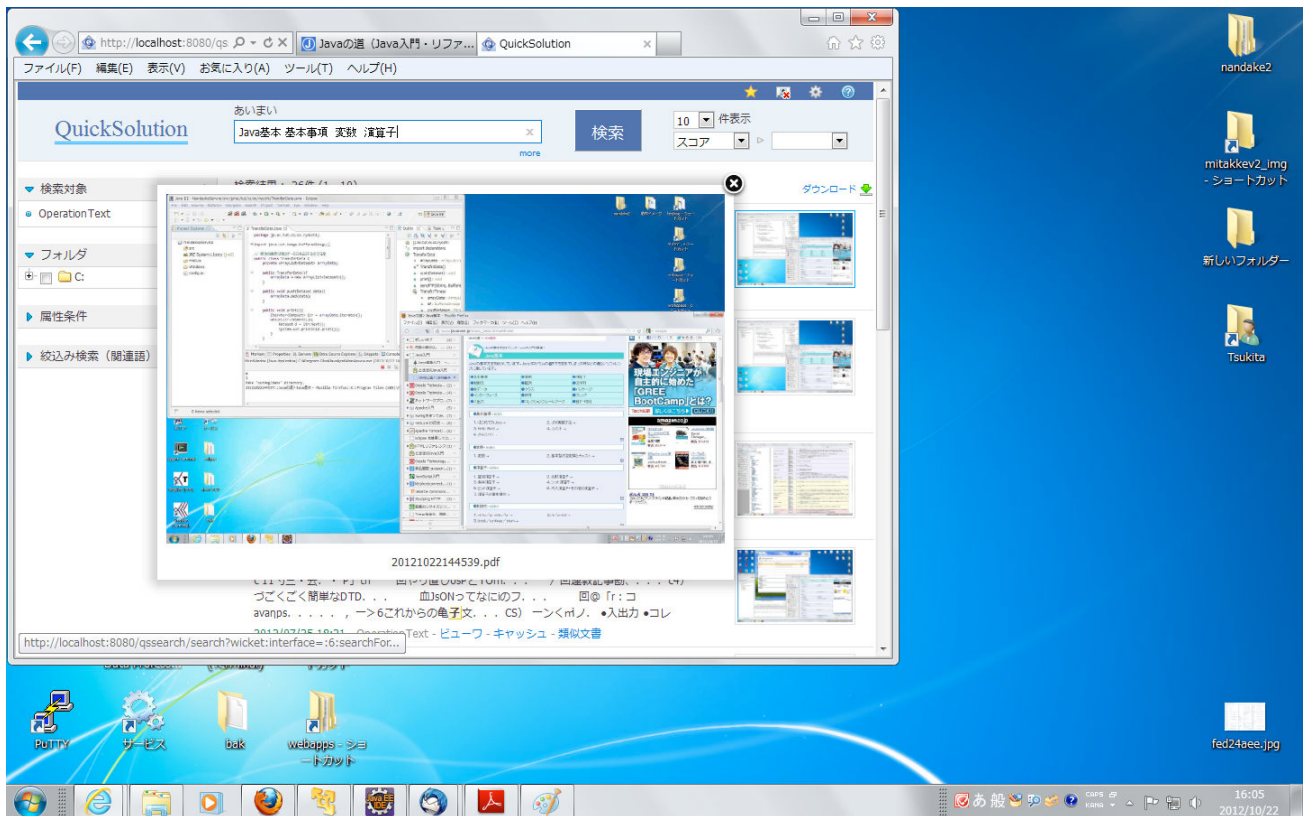
近年、書籍、テレビ、新聞、会議などの多くの情報源がコンピュータを経由して利用できるようになり、その比率は大きくなる傾向がある。ここで、装

着型の記録デバイスの代わりに、コンピュータの画面を記録することを考えた。コンピュータ経由以外の人間の情報活動は対象にならないかわりに、携帯デバイスを装着、管理する人への負担がなくなると同時に、画像としては鮮明な画像が得られ、また、画像のなかにも索引の対象となりうる文字が多く含まれる特徴がある。

コンピュータの活動を記録し検索する提案はなされている[2]が、ここでの対象は利用したメールやドキュメントそのものであり、それが画面に現れている画像とは異なっている。デスクトップ画像であれば、ある文章を入力していたとき、その編集中的イメージだけでなく、参考としてひらいていた Web ブラウザのイメージが同時に検索できる。このような「同時にひらいている」という情報は、ユーザの経験を構成する重要な要素であるが、メールやドキュメントだけを対象とすると、その情報が失われてしまう。

このような背景から、デスクトップイメージの検索を考えることができる。一方で、一般的なビデオ画像に比べて画像のなかに文字が多く含まれ、それによって検索できる利点をもちつつ、画像という人間が見た状況を検索対象としているという特徴をもつ。本稿では、これを具体的に示すために、検索のプロトタイプを作成し判明した知見を報告する。





2. プロトタイプシステムの操作手順

まず、デスクトップ画像はタイマーによって定期的に（プロトタイプシステムでは 30 秒に 1 回、設定で変更可能）検索サーバに転送しておく。サーバは、それを定期的（プロトタイプシステムでは、1 日に 1 回、設定で変更可能）に索引付けをする。

その準備が済んでいる状態でのプロトタイプシステムの操作状態を図 1 から図 3 に示す。図 1 においては、検索のために検索サーバの検索要求用の url を Web ブラウザでひらいたところである。これは市販の検索システム[5]の検索画面、そのままである。

ここから、検索語として、「Java 基本 基本事項 変数 演算子」を入力し、検索ボタンをマウスでクリックすると図 2 の画面になる。検索の意図は、以前に Java の言語機能をまとめてある Web ページをみたことがあるので、そのときのデスクトップイメージを想定している。図 2 においては、検索画面に含まれている文字列をスニペットとし、検索画面全体を小さくしたサムネイルがリストで表示される。画面は、1 日毎に日付に相当するフォルダに分かれているので、そのフォルダの情報と、時間情報に従ったファイル名の検索対象が表示される。検索のオ

プションで、時間指定もできるほか、ファイル名に日付と時間があることを利用して、検索入力にその時間いれることもできる。

サムネイルで、目的のものと思われる画面があった場合には、その画面のさらに拡大したイメージを見る事ができる。その操作を行っているのが図 3 である。これが正しいということにあれば、Acrobat Reader などの PDF 文書と関連つけられたアプリケーションで PDF に変換されたデスクトップ画像をひらくことができ、印刷やクリップボードでの利用ができる。

3. プロトタイプシステムの構造

プロトタイプシステムの構造は、すでに既に発表したシステム[3]と同一であるが、検索エンジン[5]の機能強化のために、既に発表のシステムにくらべ実用的なものになっている。

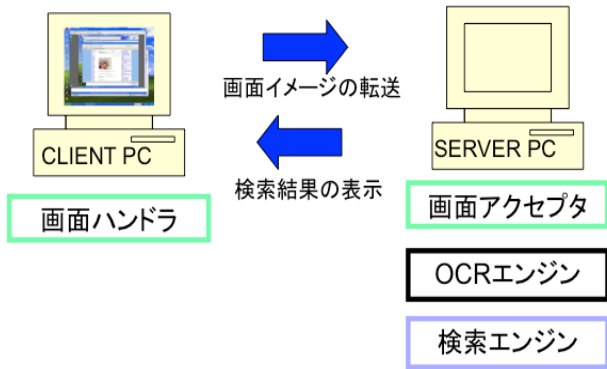


図4 プロトタイプシステムの構成

プロトタイプシステムは検索サーバとクライアントからなるサーバクライアントモデルで作成した。一つの検索サーバで複数のクライアントのデスクトップの検索が行える。図4は、[3]と同一のものであるが、サーバとクライアントに配置されているコンポーネントの名前を示す。以下に、それぞれのコンポーネントの機能を説明する。

画面ハンドラは、設定された一定間隔ごとに画面をバックグラウンドでキャプチャし、それをサーバに転送する。クライアントに追加するコンポーネントであるので、処理は最小限にした。画面をキャプチャして、サーバに転送する程度であれば、マルチコアが一般的となった近年のコンピュータでは、特に存在を感じさせない程度の負荷のコンポーネントである。

画面アクセプタは、ネットワークから画像イメージを取得したのち、相手のIPアドレスと日付から適切な格納ディレクトリを決め、それを画像jpeg形式で格納する。

OCRエンジンは、設定されたディレクトリ／フォルダに格納された画像ファイル进行处理し、同じ名前のPDFファイルに変換する。そのときに、OCR処理した文字列は透明のテキストとして、結果のPDFファイルに埋め込む動作をする。ここは、市販のOCRソフト[4]を使っている。

OCRで処理されたデスクトップイメージは、コンピュータのプログラムからするとテキストを含むPDFファイルの形式であり、処理をするときには通常のテキストのPDFファイルと同様に操作できる。それゆえ、PDFファイルを生成するフレームワークを作成したことで、市販の検索システムに手を加えるなく、プロトタイプシステムが動作した。

検索エンジンは、設定されたディレクトリ／フォルダにあるPDFを検索対象にし、一定時間ごとに索引を更新する。そして、Webブラウザを経由して、検索質問を受付、Webブラウザで結果を返す。この

検索システムは製品[5]をそのまま使用している。この応用のためにカスタマイズは行っていない。

4. OCR アプローチ

デスクトップイメージ画像には文字情報が含まれることが多いということが本システムの着眼点であるが、この文字情報を画像から取り出す操作が必要である。

もともと、デスクトップに文字を書くときには、もとの文字列をコンピュータが表示しているので、そのレベルでのオペレーティングシステムから情報を取り出すというアプローチも考えられるが、われわれは、ここにOCRを使うことにした。

OCRはコンピュータに負荷のかかる操作であり、また、読み取った結果に誤りが含まれる率も高くなる。正確な検索のためには、OCRを利用することは問題が生じるかもしれないのは事実である。

しかしながら、Webブラウザで表示されている情報、たとえば、スクリーンショットの一部を利用した操作説明などでは、人間に読める文字であってもコンピュータは文字出力でないこともある。また、OSの機能を使おうとするために、OSのバージョンの違いに影響を受けやすかったり、複数のOSで画面ハンドラを作成するのが難しくなったりする。

以上、まとめると、OCRを利用することで、処理装置への負荷や検索精度の低下が心配されるが、一方で、文字が取り出せる範囲が広がることと、プロトタイプシステム全体の機種依存性が下がることとなる。

この応用では、OCRで読まれた文字列について検索をかけるので、表記の誤りに強い検索エンジンが必要である。使用した検索エンジンは、文字バイグラムをベースに検索する機能があり、キーワードの一部がOCRの読み取りの誤りで失われても、検索はできるという性質があり、OCRの誤りで検索精度が低下する可能性はあるが、まったく検索できないということはない。

5. サムネイル機能

初期のプロトタイプシステム[3]で使用したシステムと現在のシステムとを比較すると、検索を実行した直後の図2で現れるウィンドウのサムネイル機構の有無が大きな違いとなる。

この機能は、検索エンジンの機能のバージョンアップによるものである。画面ハンドラや画面アクセプタについて本質的に[3]で作成したものと同一であり、OCRもバージョンアップにより読み取り精度

が向上したものの、ユーザのインタフェースという観点からは OCR による変更はない。

サムネイルは、検索対象を選ぶときに、スニペットのテキストを補佐し、全体的なレイアウトから対象を選べる機能として一般的なものであるが、自分が過去に見たデスクトップ画面を検索するときには、ウィンドウの色やレイアウトがシーンを選ぶときのとりかかりとして、特に有用な機能となる。

6. ビューワ機能

初期のプロトタイプシステムと比較して、もう一つ重要な検索エンジンの機能拡張はビューワ機能である。ビューワ機能は、サムネイルで見込みがありそうなデスクトップ画像について、検索ワードをハイライトした状態で、内容を確認できる機能である。

どのような画面化を確認することは、デスクトップ画像を検索するときの主要な目的と考えられ、ビューワで検索ワードがどこに表示されるかがわかることによって操作性の改善がなされた。

7. システムの負荷

現状のシステムはコンセプトの実現性の検討が目的であり、プロトタイプであって、製品ではないのでシステムの負荷は第一義ときには考慮せずに設計したが、大規模な実験にも耐えるシステムになっている。

まず、クライアントシステムについては、操作画面をネットワークに送る操作であり、その間隔として、画面がかなり変化するであろうという時間間隔、具体的には数秒から数十秒の間を想定しており、操作の妨げになるような負荷ではないし、実際に画面ハンドラが動作していることをコンピュータ操作で意識することはない。

処理速度のネックはコンピュータ画像の OCR 処理である。この処理は、ディレクトリの監視などのタイミングもあるが、1画面に十数秒かかる。しかしながら、OCR 処理は画像取得に間に合うようにおこなう必要はなく、操作がなにもないときに、一度に行うことができる。また、将来、OCR 処理がネックになったとしても画像毎に独立の処理であるので、必要に応じて、複数のコンピュータで分担する方法も容易に実装できると考えられる。

プロトタイプシステムでは固定の時間間隔でデスクトップ画像を転送するため、ハードディスク領域を圧迫することが心配されるが、消費するディスク容量を計算すると、実験のためのプロトタイプシ

テムとしては、大規模な実験にも耐えられるものであることがわかる。具体的に、1日8時間、1000日の記録をとることを考える。10秒に1枚、デスクトップ画像を取得するとすると、1分に6枚、1時間に360枚、1日に8,640枚、これを約10,000枚と考えて、100日で1,000,000枚となる。1920×1080ドットの解像度の画面から生成した一枚のPDF画像は、200キロバイト程度なので、これは、200ギガバイトのファイル容量で保持できる量となる。

将来的には、デスクトップ画像であることを想定した簡略な OCR 処理で OCR の負荷を下げることもできそうであるし、ファイル容量についても、ほぼ同じ画像は処理しないなどの工夫ができると考えられる。

8. 多くの類似文書が存在する問題

実際に検索を行ってみると、デスクトップイメージの性質に由来する問題がいくつか生じたので、それについて報告する。まず、この検索対象には類似の文書が数多く存在することがある。

図2において、複数の画面で、ほぼ同一のサムネイルとスニペットが表示されていることが分かる。これは、30秒おきに画面を取得しているためであり、画面上は、まったく変化がなくても新しい画像が新しいドキュメントとして生成される。これは、文書ファイルにおいて、自動セーブ機能により、多くのバージョンが格納されている状況での文書検索に相当する問題といえる。

この性質により、2つの課題が生じる。一つ目は、検索結果を提示するときに、類似している文書(画像)から代表の文書(画像)を選び出し、それを提示して、多様な候補を検索結果候補として選ぶ課題である。通常の実験でも、同様な課題があり、クラスタリングなどの手法が応用されるが、この場合には、取得時間あるいはバージョンが利用できる、より精度よく代表の文書(画像)を選ぶ事ができる可能性があるが、それは面白い課題である。

もう一つの課題は、検索における「検索語」に対する重みの問題である。ここで「検索語」と括弧をつけて記述したのは、使用したシステムでは文字バイグラムも「検索語」として扱っており、いわゆる単語でないものも含まれるためである。「検索語」の重みの決定には、文書頻度をもとに多くの式が考案されているが、文書の生成モデルとしては、それぞれの文書が独立の確率モデルから生成されると仮定されることが多く、使用した検索システムも、その仮定にしたがって語の重みを計算して、検索結果の候補の順位を決定している。今回の状況では、前提

となる文書生成の仮定が崩れており、そのもとでは適切な「検索語」の重みをあたりに検討する必要がある、これも重要な課題である。

9. 時間軸の扱いの問題

検索の対象となる PDF ファイルは、定まった時間で取得された画面イメージであるので、全体として操作画面を微速度撮影した映像とみなすことができる。これは、通常の実検索とは異なる性質である。通常の実検索では、個々は独立のファイルの集合と考えるので、検索結果について、前後を見る機能がない。

このような機能はビデオのフレームを対象とする検索のように、検索結果から前後を表示したり、ビデオのキーフレームに対応する画面を特定したりすることが必要となる。

ここでも、自動保存があるようなアプリケーションにおいて、すべての中間状態を含めて文書検索をする状況とデスクトップイメージの実検索は類似性がある。この問題において、デスクトップイメージの実検索は1ストリームのデータであるのに対し、文書情報は、ファイル上は1ストリームであっても、カットペーストにより枝分かれやマージが生じうるので、より複雑な検索対象といえる。

以上のことより、バージョンや名前の変更などを含めた文書の時間軸変化を考慮した検索研究を行うための準備として、デスクトップ画像を実検索対象にして検索機能の検討をおこなうことは価値があるように思える。

10. 画像検索の実使用のシナリオ

検索における新技術の開発のためだけでなく、実際的な状況でも画面検索が有効と考えられるケースがある。

最初のシナリオは、コンピュータを用いた演習を行っているケースであり、教師が生徒ごとの進捗状況をまとめたいというシナリオである。課題の区切りごとに、操作画面に特徴のある文字列、たとえば、課題を遂行するプログラムに含まれる文字列などを手掛かりに、検索システムを利用して進捗を一覧にして、課題の難度を確認することや、進捗に問題のある生徒を発見するというシナリオは存在しうる。

次のシナリオは、一連の操作をしているときに、エラーがでた状況を検索し、そこから画面を巻き戻して見て、エラーの原因をさぐるようなシナリオである。

どちらの場合も、操作画面をビデオで取得して、早送り、巻き戻しなどで場所を特定するより、キー

ワードを指定して検索するほうが効果があることが期待できる。現状のプロトタイプでもこのようなシナリオのために利用できるが、より操作性を高めるように検討すること、検索システムのユーザインタフェースを改良していくこともできると考えられる。

まとめ

本稿では、デスクトップ画像を実検索するという問題に対し、既存の OCR システムと検索システムを利用してプロトタイプの作成をし、そのプロトタイプの操作を報告した。そして、プロトタイプシステムがコンピュータにかかる負荷の大きさを示し、現在のコンピュータでは、十分に実験できる対象であることを述べた。

そして、この課題は、自動保存のある文書と類似性があることを述べ、その準備として問題が単純化されており、一方で、ビデオ検索との接点もあることを述べた。

最後に、デスクトップ画像の実検索が、それだけで応用の価値があると考えられるシナリオを示し、プロトタイプから次のステップに行くときの方向性を示唆した。

参考文献：

- [1] 堀 鉄郎, 相澤 清晴: “ライフログビデオのためのコンテンツ推定”, 映像情報メディア学会技術報告 Vol. 27, No. 72 pp. 67-72, (2003)
- [2] Susan Dumais: Stuff I've Seen: a system for personal information retrieval and reuse, ACM SIGIR 2003, pp. 72-79, (2003)
- [3] 熊谷摩美子, 梅村恭司, 岡部正幸, 阿部洋丈: 操作画面を対象とする検索システムの構築, 情報処理学会 72 回全国大会, 4R-2, pp767-768, (2010)
- [4] 株式会社ハイパーギア: HG/PscanServPlus 製品ページ, <http://www.hypergear.com/index.html>, (2012)
- [5] 住友電工情報システム株式会社: Quick Solution, <http://www.sei-info.co.jp/quicksolution/index.html>, (2012)

データの楽曲としての可聴化 —心電図データのドラム演奏化—

Sonification of Data as Music —Drum Play Based on Electrocardiogram—

柳田拓人^{1*} 沖田善光¹ 中村晴信² 杉浦敏文¹ 三村秀典¹

Takuto Yanagida¹, Yoshimitsu Okita¹, Harunobu Nakamura², Toshifumi Sugiura¹,
and Hidenori Mimura¹

¹ 静岡大学

¹ Shizuoka University

² 神戸大学

² Kobe University

Abstract: There are many methods for analyzing heavyweight and composited numeric data corresponding to various statistical matrix and types of data. Conventionally, for representation of these analytical results, charts and their extensions by three-dimensional view, animation, and interactivity are used. As one of the representations, we aim at sonification of analytical results, especially representation as music, for stimulating perceiving the difference of the target data. In this paper, we report software that generates drum play from arbitrary numeric data.

1 はじめに

長大で複合的な数値データ、例えば、長時間に渡る様々な生体データや、近年、ビッグ・データと呼ばれる大規模データ・セットに対して、それらを解析するために、様々な統計指標（指標値セット）やデータ種別に合わせた処理手法が提案されている。これまで、その指標値セットの表現方法として、従来から存在する各種グラフや、三次元化、アニメーション化、インタラクティブ化などによるその発展版といった、データの可視化手法が用いられてきた。しかしながら、対象となるデータが複合的であればあるほど、また、その性質が未知であればあるほど、データの解析に必要な指標値セットのサイズは大きくなり、可視化手法による表現そのものや、表現されたものに対する解釈が困難になることが予想される。

そこで筆者らは、性質が未知である解析対象データにおいて、複数のデータ間の差異に対する気付きを促すことを目的として、指標値セットの可聴化、特にメロディ、リズム、パートなどによって元来、複合的に構成される楽曲としての再構築と表現、すなわち可聴化を目指している。本稿では、その目標に向けた試みの一つとして、著者らがこれまでに未病分野の研究とし

て心電図・脈波解析システムを開発してきたことを踏まえ、心電図出力などの時系列数値データをドラム演奏に変換するソフトウェアを開発したので、報告する。

2 従来研究

筆者らはこれまで、生体データである心電図と脈波の波形データ（時系列数値データ）を解析し、両波形に特有の形状から導き出される各種指標値を算出するソフトウェア、並びにその表示用ソフトウェアを開発してきた。また、指定の一拍波形（心臓の鼓動一拍と対応した心電図波形）に対応する各種指標値の、正常範囲からの逸脱の有無にしたがってドラム演奏パターンを構築し、演奏させるシステムも提案し、筆者らのソフトウェアに実装している [1]。

しかしながら、これは心電図という特定のデータとその指標値に特化したものであり、しかも、その指標値が正常範囲に入っているかどうかという二値データの組み合わせをドラム・セットにおける各楽器の打音の有無に対応付けたものであるため、一般的な統計指標に対する応用が難しい。一方で、このような生体データを解析しようとする中で、多種多様な統計指標は得られるものの、それらが何を意味しているのかを解釈する上で必要となる、データ同士の関係や差異を直感的に捉え難いという課題があった。

*連絡先：静岡大学電子工学研究所
〒432-8011 浜松市中区城北 3 丁目 5-1
E-mail: takty@rie.shizuoka.ac.jp

3 手法

ドラム演奏を幾つかの離散的なパラメータによって表現し、それと表現対象のデータの統計的な指標値とを対応付けることによって、データをドラム演奏に変換する。コンピュータ上で演奏データを扱う手法として、一般に、MIDI 規格が広く用いられており、プラットフォームや言語環境によっては、プログラム上から扱う手段が提供されている。MIDI データにおいて、ドラム演奏は基本的に、どの楽器をどのタイミングでどの強さ（ベロシティ）で演奏するのかという情報として表現される。このように MIDI データは演奏データとしては非常にプリミティブなため、単純にこれをドラム演奏のパラメータとすることは出来ない。そこで、本稿では、一般的に使われるリズム・パターン、並びにそれに付随するドラム演奏上に必要な要素をパラメータとして取り扱う。

3.1 ドラム演奏のパラメータ化

本稿では、ドラム・セットとして、バス・ドラム（キック）、スネア・ドラム、ハイハット、ライド・シンバル、クラッシュ・シンバルを扱い、参考文献 [2] に掲載されている 4 拍子または 3 拍子で 2 小節分ずつある 40 のリズム・パターン（内上記以外の楽器を使用したものを除く）を基本リズム・パターンとして採用する。また、ハイハット、ライド・シンバル、クラッシュ・シンバルはリズム・パターンにおける役割が似ていることからひとまとまりとして考え、基本リズム・パターンを、バス・ドラム、スネア・ドラム、ハイハット他の 3 パートに分ける。さらに、パートごとに、ベロシティの違い、スウィング（リズムのハネ）の違いはそれぞれ後述するパラメータによって表現することとし、パターン自体を共通化する。そして、他の基本となるベロシティや拍子数、また、その分割数などを含めて、22 のパラメータにまとめる（表 1）。

各パラメータは、数値が大きくなると演奏の複雑さが増すように設定する。そこで、パートごとに分けた基本リズム・パターンを、次の複雑さの定義に従って並び替える。パターンは基本的に 4 拍子で 2 小節分なので、パターンを 8 拍分に分割し、各 1 拍パターンの出現頻度が大きなものから順に並ぶように順位を付ける。さらにこの 8 拍を、偶数拍、奇数拍では奇数拍を、1 小節目、2 小節目では 1 小節目を優先するように順位を付ける。これによって、パートごとに基本リズム・パターンを、その中で一般的であるものが先頭になるように並び替える（図 1）。

「ビート強調: 小節強調ベロシティ」、「拍強調ベロシティ」、「半拍強調ベロシティ」の各パラメータは、それぞれ各小節の最初の音、各拍の最初の音、各半拍の

表 1: ドラム演奏のパラメータ.

パターン	種類 拍子数 分割数
ベロシティ	主楽器 副楽器
揺らぎ	タイミング ベロシティ
スウィング	半拍% 種類（表、裏、両方） 4 分の 1 拍%
ビート強調	小節強調ベロシティ 拍強調ベロシティ 半拍強調ベロシティ

最初の音について、基本的なベロシティからどれだけベロシティを増加させるのかを表す（図 2）。例えば、拍強調ベロシティを 10 にすると、リズム・パターンのうち、各拍の表のタイミングでの音のベロシティが他よりも 10 増加し、その結果、4 ビート感が強調される。このパラメータとリズム・パターンの組み合わせによって、いわゆる 8 ビートや、16 ビートといったリズムを表現する事が可能である。

スウィングは、リズムの「ハネ」を表現するパラメータである（図 3）。「スウィング: 半拍%」は、1 拍を 4 分音符で表す時、その 4 分音符を 2 分割する割合を表し、均等に分割する、すなわちハネていないリズムの時は、50%となる。同様に「スウィング: 4 分の 1 拍%」は、8 分音符を 2 分割する割合を表す。なお、半拍%に関しては、偶数拍のみ、奇数拍のみ、あるいは両方をハネさせるのかを設定可能とする。スウィングのパラメータによって、ジャズのようなジャンルにおいて用いられるリズム・パターンを再現することが出来る。

その他のパラメータについては以下のとおりである。「パターン: 拍子数」は、1 小節における拍の数を表し、4 拍子のリズム、3 拍子のリズムといったものを表現する。「パターン: 分割数」はリズムの最小時間単位をその 1 拍の何分の 1 にするのかを表す。例えば、分割数が 4 の時、拍子が 4 分音符で表される時、その最小単位は 16 分音符となる。「ベロシティ: 主楽器」はパートの演奏の強さを表し、「ベロシティ: 副楽器」は複数の楽器を含む基本リズム・パターンにおける補助的な楽器の強さを表す。「揺らぎ: タイミング」はリズム・パターン規定のタイミングからの変動の上下最大幅を、「揺らぎ: ベロシティ」は上記ベロシティとビート強調によって指定された規定のベロシティからの変動の上下最大幅を表す。これによって、演奏の不正確さを表現する事が出来るため、ドラム演奏にリアリティを持たせることが可能となる。

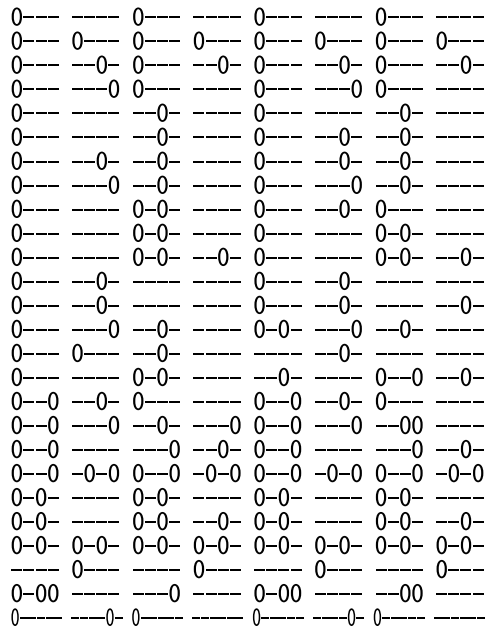


図 1: パス・ドラムのリズム・パターン．二つある小節のそれぞれを 16 分割した時に音のある位置を「0」で表した．パス・ドラムを含む 36 パターン中，重複していない 26 パターンを一般的なものから順に並べた．

3.2 指標値セットの取得

汎用的な解析データに対する指標値として，一般的な統計値を求める．本稿では，与えられた時系列数値データの，サンプル数，最小，最大，合計，平均，分散，標準偏差を統計値として扱う．また，データのパワー・スペクトル，位相スペクトルについても同様に統計値を求める．さらに，第 1 から第 3 までの「フォルマント」についても，その周波数とパワーの他に同様の統計値を求める．ここでのフォルマントは，解析データを周波数分析してもとめたパワー・スペクトルにおける幾つかのピークである¹．フォルマントを検出し，その大きなものから順に三つを選び，それぞれ，その周波数帯域のみを通過させるバンド・パス・フィルタ（実際にはゼロ・パディングと逆 FFT）を適用する．このように取り出した三つデータそれぞれに対して，統計値を計算する．以上により，任意の時系列データに対して，43 の統計値が得られる．

指標値と演奏パラメータとの対応付け，ならびに，データの比較を容易にするため，複数のデータから算出した 43 の統計値それぞれに対して，複数データ内での最小値と最大値を求め，それにしたがって値を 0 から 1 に正規化する．これによって，任意の統計値をドラム演奏パラメータと対応づけ，複数のデータ間での

¹ フォルマントを検出したのは，筆者らのこれまでの研究において，生体信号である心電図や脈波を扱った経験による．

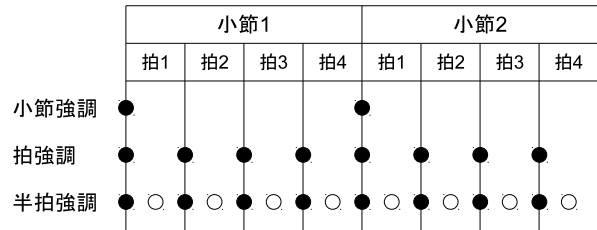


図 2: ビート強調，小節強調，拍強調，半拍強調，それぞれにおいて，対応する丸印の位置にある音のペロシティが増加する．なお，半拍強調において白丸で示したのは，裏拍と呼ばれるタイミングである．

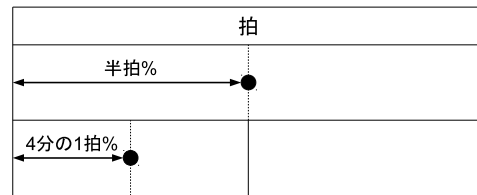


図 3: スウィング．半拍%は 1 拍を 2 分割する時の前半の割合を表す．同様に，4 分の 1 拍%は，半拍を更に 2 分割する際の前半の割合を表す．いずれも，50%のときに等分，すなわちスウィングしていない（ハネていない）リズムを表す．

違いを演奏パラメータの違いに置き換えることが容易となる．なお，指標値の複数データ間の最大値，最小値によらずとも，対象となるデータの性質がある程度わかっている場合などは，目的に応じて，任意の範囲を正規化することも可能である．

3.3 実装

Java 言語により，任意の時系列数値データを CSV 形式で複数読み取り，その統計値に基づいてドラム演奏をするプログラムを実装した（図 4）．ドラム演奏のパラメータについては，幾つかを統合し，テンポ（BPM）を加えることによって 22 にまとめた．統計値とドラム演奏パラメータとの対応付けは，ユーザが画面上で「パッチ」を繋げることによって任意に設定可能とした．

4 課題

今後の課題は次のとおりである．まず，ドラム演奏だけではない他の楽器を含めた楽曲を実現するために，演奏のパラメータ化が必要である．本稿ではドラム演奏に限定したが，それでも前述のとおり参考文献に存在するリズム・パターンの組み合わせに限定しても 22 パラメータ存在する．今後，例えば，ロック，ジャズ，

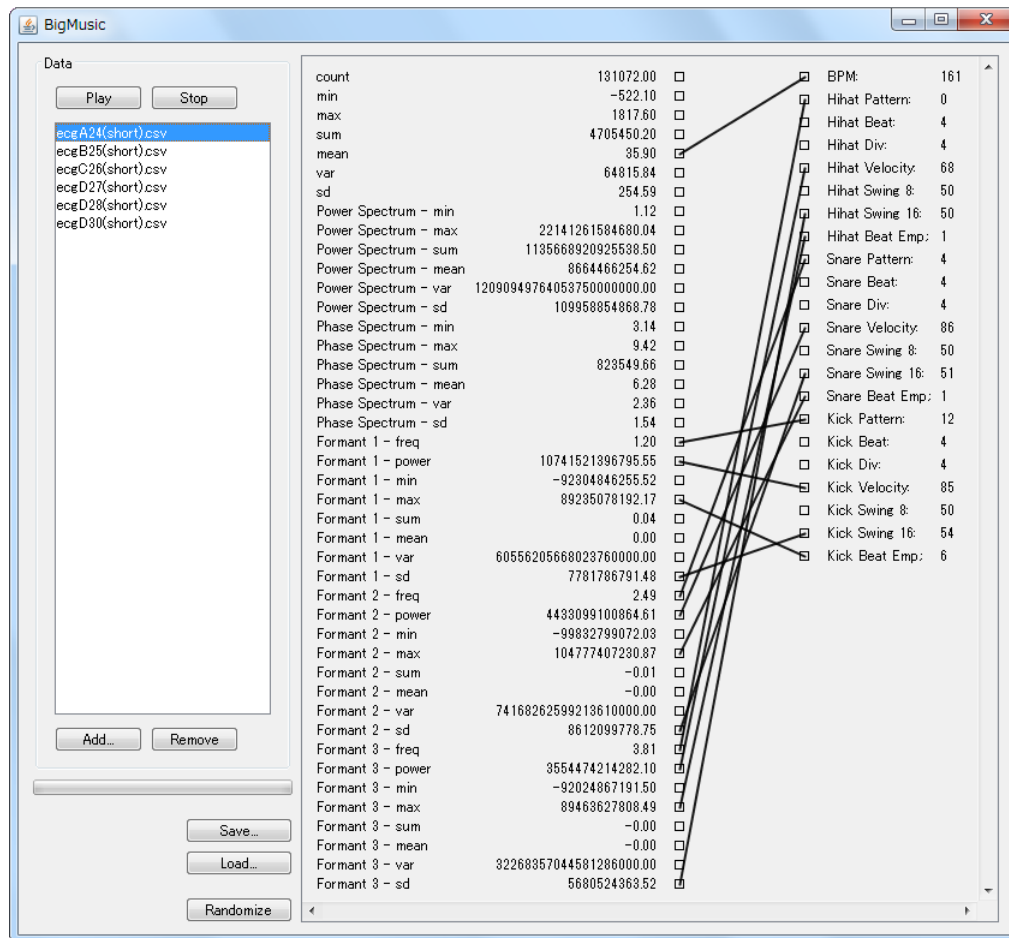


図 4: ドラム演奏プログラムの画面。

R&B といった一般的な楽曲の実現は困難であるとしても、例えばミニマル・ミュージックのような形式の楽曲を実現するためには、ドラム以外の楽器の演奏も考慮する必要がある。そういった他の楽器の演奏をある程度限定されたパラメータによって、かつ自由度を持って表現する手法を今後検討する。

次に、指標値セットと演奏パラメータとの対応付けをどのように決定するかが、情報の表示という観点から重要である。今後、指標値セットのサイズ、ならびに演奏パラメータ数が大きくなった場合、その適切な組み合わせを解析対象に合わせて人間が決定するのは容易ではない。したがって、指標値セットと演奏パラメータとの組み合わせ問題として、自動化、半自動化手法が必要とされる。ここでは、差異の大きな指標値から順に、人間にとって気づきやすい演奏の違いを生み出す演奏パラメータと対応づけることが重要である。

さらに、そもそも人間が演奏の差異をどれだけ聞き取れるのか、気づくことができるのかという問題もある。先に述べたように、楽曲は本質的に（それが容易に取り出されるかどうかは別にして）パラメータ化され

構造化されたデータの複合的な表現であり、これを複合的な指標値セットと対応づけることは、自然な試みであると考えられる。しかしながら、その複数の楽曲から違いを見出すことができるかどうかは、聞く側の音楽的能力に依存する可能性がある。楽曲として表現された指標値セットの分かりやすさの評価、あるいは、単純な数値としての列挙や従来の視覚的な表現手法との比較も、今後の課題である。

参考文献

- [1] 柳田拓人, 沖田善光, 中村晴信, 甲田勝康, 杉浦敏文, 三村秀典, “Heart beat drum: 未病診断のための心電図聴覚化,” 電子情報通信学会技術研究報告 (HIP, ヒューマン情報処理) HCS2012-8, HIP2012-8, 第112 巻, pp.43-46, 那覇, 05-22 to 23 2012.
- [2] 船橋識圭, “リズムパターンマスター 即戦力の MIDI グルーヴ 200,” DTM Magazine 6 月号, 2012.

モバイルファースト戦略に基づく Android 対応システム WATATUMI の電子カルテについて

Electronic Medical Record System for WATATUMI Android Compatible Mobile Strategy Based on Mutual Agreement First

串間宗夫¹ 田之上光一¹ 酒田拓也¹ 荒木賢二¹ 鈴木斎王¹ 荒木早苗¹ 山崎友義¹
曾根原登²

Muneo Kushima¹, Kouichi Tanoue¹, Takuya Sakata¹, Kenji Araki¹, Muneou Suzuki¹, Sanae Araki¹,
Tomoyoshi Yamazaki¹, and Noboru Sonehara²

¹ 宮崎大学医学部附属病院医療情報部

¹Section on Medical Information, Faculty of Medicine, University of Miyazaki

² 国立情報学研究所情報社会相関研究系

² Information and Society Research Division, National Institute of Informatics

Abstract: WATATUMI is a mobile electronic medical record that runs on Android. WATATUMI has access to the same database(Cache') as PC version electronic medical records (IZANAMI).

In addition, it uses a different interface with the same database. One month prior to the introduction of WATATUMI in May 2011, We operated in parallel with the nursing PDA by using smartphones in some traditional hospital rooms. The time study conducted in that period, surveyed nursing workload input. We report on the results.

1. はじめに

WATATUMI は、Android 上で稼働するモバイル用電子カルテである。PC 版電子カルテ (IZANAMI) [1]と同じデータベース (Cache') にアクセスし、両者は同じデータベース上の異なるユーザーインターフェイスである。

2011 年 5 月の WATATUMI 導入前の 1 か月間に、一部の病棟で従来の PDA[2-3]とスマートフォンによる看護業務を併行運用した。その期間でタイムスタディを行い、看護入力業務負荷の調査を行い、その結果について報告する。

2. モバイルファースト戦略

モバイルファースト戦略とは、ソフトウェアの開発、導入において、PC 版よりモバイル版を優先することである。表 1 に宮崎大学病院のモバイルファースト戦略についてまとめた。

近年、Apple 社やマイクロソフト社は、明確なモバイルファースト戦略の下に OS 等の開発を行なっている。今後、宮崎大学病院の電子カルテにおけるモバイル端末の導入は、モバイルファースト戦略に

表 1 宮崎大学病院のモバイルファースト戦略

コンセプト	・ モバイルファースト戦略とは、ソフトウェアの開発、導入において、PC版よりモバイル版を優先することである。近年、Apple社やマイクロソフト社は、明確なモバイルファースト戦略の下にOS等の開発を行なっている。今後、宮崎大学病院の電子カルテにおけるモバイル端末の導入は、モバイルファースト戦略に基づいて進めていくこととしている。電子カルテにおけるモバイルファースト戦略のメリットは、業務の迅速化である。医師の判断や指示が遅れると、時として患者の生命に関わることもある。
特徴	・ PC版電子カルテの機能をモバイル版に移植 ・ ユビキタス環境で利用価値の高い機能はPC版よりモバイル版を優先して開発 ①コミュニケーション機能の強化 ②電子カルテと電話の連携 ③リアルタイム動画配信(手術場ビデオ、生体モニター)
効果	・ 業務の迅速化

基づいて進めていくこととしている。電子カルテにおけるモバイルファースト戦略のメリットは、業務の迅速化である。医師の判断や指示が遅れると、時として患者の生命に関わることもある。

3. モバイル導入の経緯

宮崎大学病院では、2006 年 5 月の電子カルテ導入時に、すでに、業務用 PDA (NEC インフロンティア製 Pocket@iEX) をベッドサイド看護業務のために 250 台導入し、観察項目の入力、注射・輸血の実施入力、手術場の本人確認、等に活用し、運用に乗っ

ていた。

2011 年 5 月の電子カルテ更新を控えて、PDA の機器更新について検討を行ったが、業務用 PDA は高価なために、十分な台数を確保するのが難しく、急速に普及が進んでいるスマートフォンを代替機とすることとした。1 台の価格は 3 分の 1 となった。過去のベッドサイド携帯端末 PDA の画面を図 1,2 に示す。2011 年 5 月に、従来行なっていたベッドサイド看護業務をモバイル端末 (GalaxyS230 台) に置き換え、稼働後、順次機能を追加していった。2012 年 5 月には、医師のほぼ全員にモバイル端末 (GalaxySII350 台) を配布し、医師のユビキタス環境構築を図った。表 2 に WATATUMI 開発の背景をまとめた。更に、表 3 に WATATUMI 開発の経緯を示す。また、回診用として、タブレット端末も各科に 2 台 (10 インチと 7 インチ) 配布した。2012 年 12 月には、PHS の後継としてモバイル端末に SIP フォン機能を追加する予定である。

4. システム概要

WATATUMI は、Android 上で稼働するモバイル用電子カルテである。図 3 に WATATUMI システムの構成を示す。PC 版電子カルテ (IZANAMI) と同じデータベース (Cache) にアクセスし、両者は同じデータベース上の異なるユーザーインターフェイスである。データ連携を行うのではない。図 4 に WATATUMI の画面遷移を示す。

現在稼働している機能は、患者リストバンドのバーコード読み取りによる患者認証、入院患者一覧、患者検索、全オーダーの参照、全オーダーのステータス変更 (実施入力等)、観察項目の参照と入力、検歴参照、全カルテ文書参照・承認、指示簿参照、画像参照、写真撮影、顔写真撮影、コミュニケーション機能、である。

今年度中の開発として、SIP フォンと電子カルテの連携、カルテ文書入力、処方・注射・検査オーダー入力、熱型表・重症熱型表参照、生体モニタリアルタイム参照、である。一部の機能は、PC 版電子カルテにないか、モバイル版を先に稼働させており、モバイルファースト戦略の現れと言える。図 5 にシステムの機能拡張について示す。(白色は提供済み。緑色は 2012 年 12 月以降。)

5. 評価

2011 年 5 月の WATATUMI 導入前の 1 か月間に、一部の病棟で従来の PDA とスマートフォンによる看護業務を併行運用した。その期間でタイムスタデ



図 1 ベッドサイド携帯端末 PDA



図 2 PDA の画面

表 2 に WATATUMI 開発の背景

	過去の端末の問題点	WATATUMIでの対応
端末(台数)	NECインフロンティア Pocket@i EX (250台)	Androidスマートフォン 看護師 GALAXY S SC-02B SIMなし (250台) 医師 GALAXY S II LTE SIMなし (400台)
携帯性	ポケットに入らないので、折角の携帯端末なのに台車に乗せている。首から下げると肩が凝る。	看護師のポケットに入る。首から下げて問題なし。
操作性	画面展開は遅い。	画面展開早い。 医師の電子カルテ参照端末としても申し分ない速度。
バーコード	赤外線ボタンを押しながら対象に合わせるのが難しい。	CCDカメラ 操作は簡単。対象にカメラを合わせるのは、慣れが必要。
機能の拡張	業務目的以外なし	スマートフォンなので、一般のアプリが豊富。 インターネット接続も可。 院内電話やグループウェアとしての利用も可。
価格	価格が高い。	半額以下に低下。

ィを行い、看護入力業務負荷の調査を行った。調査の結果、入力時間は短縮し、PDA と比べて入力業務負荷が減少する効果を得た。

全病棟導入後の 7 月に行った利用者アンケート調査 (対象は看護師 421 名、回収率 : 88.6%) より、携

帯性、操作性、視認性の項目で利用者より高い評価を得た。導入後の”評判”として、看護師からは「ベッドサイド業務には必須」「患者から急に検査結果を聞かれてもすぐに答えられるので便利」「オーダーの変更が手元で分かるので便利」などがある。医師からは、「回診で画像が見られるのが便利」とのプラスの意見以外に、「PHS と 2 つ持つのは面倒なので早く PHS 機能を追加してほしい」「看護師からオーダー切れを伝えられても、オーダー入力ができないのは不便」などがあり、これらの意見を踏まえて、今年度中に対応する。

6. 今後の展望

医療従事者がユビキタス環境を備えることにより、情報共有の拡大、迅速化、高度化を図ることができる。このメリットを電子カルテに限定する必要はない。全職員が、メール、カレンダー、ファイル（スプレッドシート、文書）等を迅速に共有すれば、通知の徹底、会議の効率化、災害時の対応等で大きな効果が期待できる。このような考えの下で、宮崎大学医学部では、Google のグループウェアである GoogleApps を 2011 年 8 月に導入し、2012 年 4 月に、メールアカウントの Gmail への移行を完了した。今後は、診療に限定せず、あらゆる業務を対象にモバイルファースト戦略を進めていく予定である。

参考文献

- [1] Muneo Kushima, Kenji Araki, Muneou Suzuki, Sanae Araki and Terue Nikama : Text Data Mining of In-patient Nursing Records Within Electronic Medical Records Using KeyGraph, The International Association of Engineers (IAENG) International Journal of Computer Science, Vol.38, Issue.3, pp.215-224, (2011).
- [2] 山崎友義, 飯塚真人, 有田憲司, 串間宗夫, 鈴木斎王, 荒木賢二, 奥村智子, 高橋歌子, 甲斐由紀子, 梅本勝博: アンドロイド端末運用による電子カルテ入力業務負荷軽減効果の検討, 第 31 回医療情報学連合大会, Vol. 2-G-1, No.3, pp. 630-631, (2011).
- [3] 串間宗夫, 荒木賢二, 鈴木斎王, 荒木早苗, 仁鎌照絵, 有田憲司, 山崎友義, 甲斐由紀子: Android 端末による電子カルテシステムのバーコードスキャン機能について, 第 31 回医療情報学連合大会, Vol. 1-I-3, No.4, pp. 868-869, (2011).

表 3 WATATUMI 開発の経緯

年月	経緯
2006年5月	旧システムのベッドサイド携帯端末稼働。運用には乗っているが、「大きい」「高価」等の問題もあり。
2009年～	次期に向けて検討を本格的に開始。
2010年7月	ワイヤレスジャパン2010にて、Android端末の普及を確信
2010年9月	Androidを前提に次期携帯端末システムの設計開始。
2011年1月	端末機種選定
2011年2月	看護師に対しソフトウェアレビュー実施
2011年3月	操作訓練開始
2011年4月	一部の病棟で、並行運用開始
2011年5月	ANDROIDスマートフォンによる携帯端末WATATUMI本稼働
2011年9月	医師向けのサービス(検歴、画像)開始予定
2012年3月	医師向けのサービス第2弾(文書、連絡、写真)
2012年5月	医師全員に端末配布 (GALAXY S II LTE約400台)

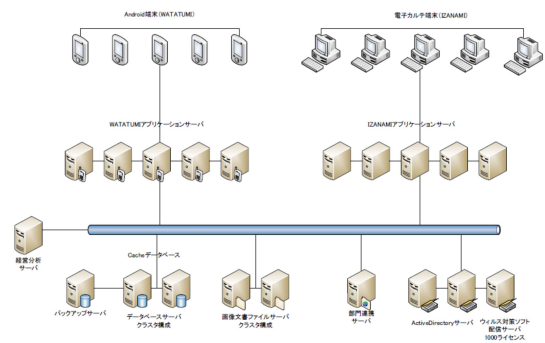


図 3 WATATUMI システムの構成



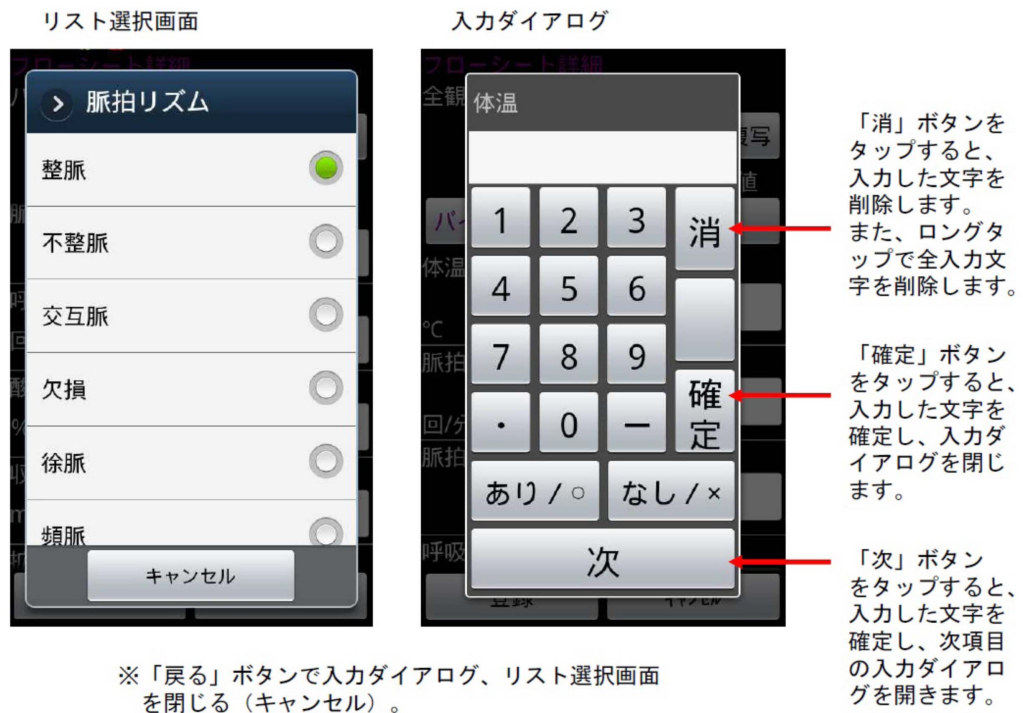
(a)



(b)



(c)



(d)

図 4(a),(b),(c),(d) WATAUMI の画面遷移

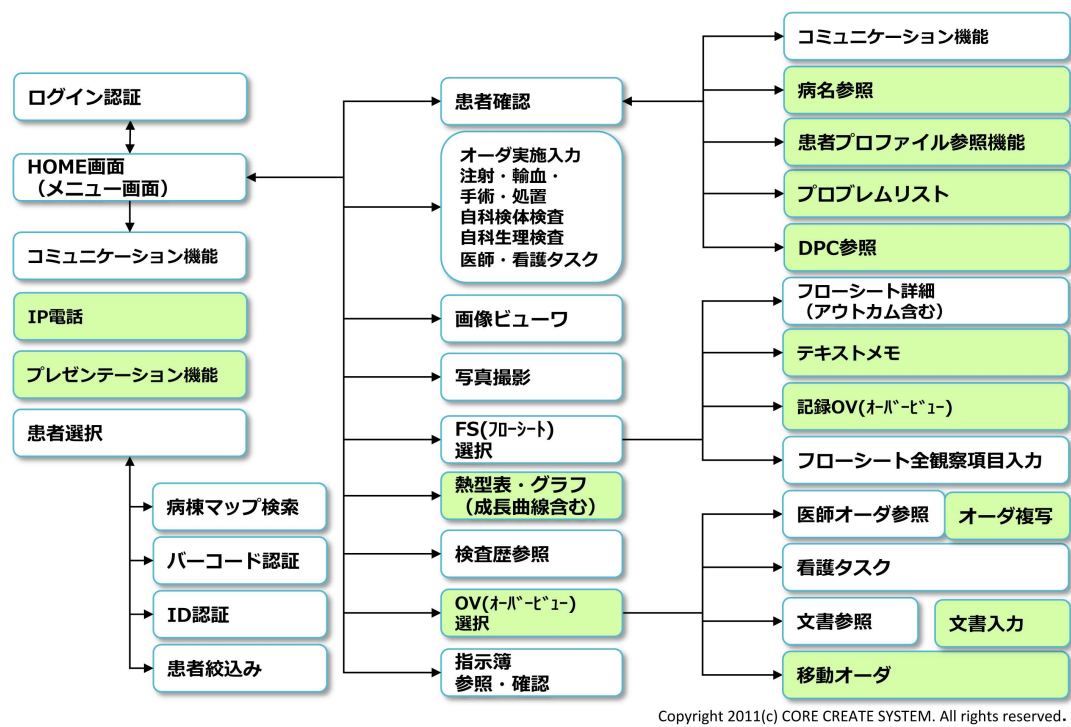


図5 システムの機能拡張について示す。(白色は提供済み。緑色は2012年12月以降。)

写真属性と画像特徴を用いたホット撮影スポット・ アノテーション

Hot Photo-Spot Annotation using photo attributes and image features

小関 基徳¹ 熊野 雅仁^{2*} 亀井 貴行¹ 小野 景子² 木村 昌弘²
Motonori Koseki¹ Masahito Kumano² Takayuki Kamei¹
Keiko Ono² Masahiro Kimura²

¹ 龍谷大学大学院理工学研究科電子情報学専攻

¹ Division of Electronics and Informatics, Graduate School of Science and Technology, Ryukoku University

² 龍谷大学 理工学部 電子情報学科

² Department of Electronics and Informatics, Faculty of Science and Technology, Ryukoku University

Abstract: In our previous work, we presented a method of extracting a pair of a major photo-spot and its hot-period, which is called a *hot photo-spot*, from a large number of geotagged photographs with timestamps that many people have taken. However, as for explaining the hot photo-spots extracted, it was in general difficult to annotate them clearly since each of them can have a variety of photos. In this paper, we propose a method of explaining each hot photo-spot by classifying the photos in it based on their image features and attributes such as their geotags and timestamps. Using real data from “Flickr data”, we experimentally demonstrate the effectiveness of the proposed method.

1 はじめに

近年、写真を撮影する際、どこで撮影したか(地理情報)を、写真に付与できるデジカメやカメラ付き携帯端末が一般化し始めている。また、Flickr¹など、多くの写真共有サイトが賑わいを見せており、写真と共に地理情報を登録できる機能を備えているため、Web空間に共有化された、膨大な地理情報付き写真データが蓄積され続けている。写真は、撮影者の心をつかむ対象に遭遇したとき撮影されることが多いことから、写真が単なる記録ではなく、撮影者の何らかの意見を内在化させていると考えることができる。つまり、大量の写真群は、意見の集合と見なすことができ、写真そのものから得られる情報や、付随する情報をうまく集約すれば、集合知が得られる可能性がある[味八木 10]。

一方、Web空間に電子化された観光情報が溢れるに従い、計算機科学の領域では「観光」が注目されており[川村 10, 松原 11]、近年、観光や旅行支援への応用も期待できる新しいアプローチとして、地理情報付き写真群を用いる研究が脚光を浴びている。Crandallらは、

大量の地理情報付き写真と、写真の画像特徴を用いて、空間的なクラスタリングを行い、多くの人が訪れる人気スポットや、ランドマークのある主要地域が得られることを示した[Crandall 09]。この地理情報付き写真群を用いる空間に着目した研究は、魅力的な地域を抽出し[Kisilevich 10b]、視覚的に探索する研究[Kisilevich 10a]や、写真に付与された文書情報も利用することで、地域ごとの地理的トピックを抽出し、地域間の文化を比較して新たな知識を発見する研究[Yin 11]や、観光マップを生成する研究[王 11]などに派生している。

また、Crandallらは、同一の撮影者が同日に複数の写真を撮影した場合、写真が撮影された時間情報を追跡し、地理情報と併用することで、撮影地点の軌跡が得られることも示した。この地理情報および時間情報を用いる研究は、旅行する人々の写真撮影行動から旅行行動をマイニングする研究[Arase 10]や、旅行の計画を支援する研究[Yin 10]、旅行計画の経路を生成する研究[Lu 10]などに派生している。

Crandallらが抽出した人気スポットは、実空間に局在する地域に写真群が密集することを重点に置くため、年間を通じて人々が訪れ写真を撮る地域が優先的に抽出される傾向がある。しかし、写真の撮影スポットを推薦する問題を考える場合、「どこで」という地理情

*連絡先： 龍谷大学
滋賀県大津市瀬田大江町横谷 1-5
E-mail:{kumano,kono,kimura}@rins.ryukoku.ac.jp

¹<http://www.flickr.com/>

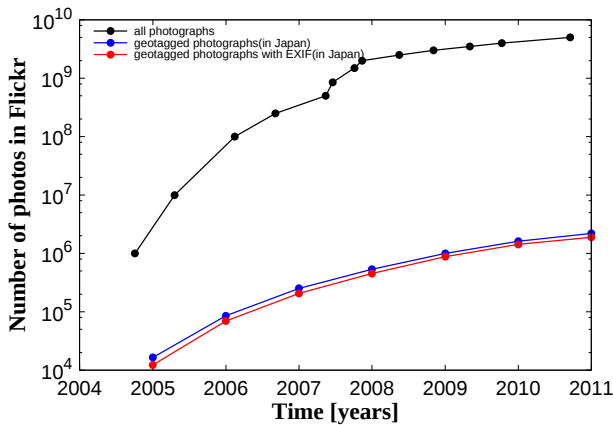


図 1: Flickr の登録写真数変遷

報だけでなく、「いつ」という時期の情報が欠けていると、旬のある撮影期間を逃しかねない。我々は、多数の撮影者が生み出した地理情報および時間情報付き大規模写真データから、撮影者の意見が反映され、他の地域と比較して普段とは逸脱して顕著に撮影数が増える、集合知的観点を背景とした格別な地域と期間のペアをホット撮影スポットと呼んで、その自動抽出問題に取り組み、ホット撮影スポットが観光スポットとして推薦できる可能性を示した [熊野 12]。

ところで、ホット撮影スポットは、空間的クラスタリングを行う上で、mean-shift 法 [Comaniciu 02] を用いているが、そのパラメータとして、カーネル関数のカーネル幅 h をあらかじめ与える必要があり、ホット撮影スポットには、 h の設定により、多数の写真が含まれることがある。ホット撮影スポット内に多数の写真が含まれる場合、一つ一つの写真に Geo-tag 情報が付随しているため、より詳細に調べれば、サブ撮影スポットが存在する可能性がある。しかし、ホット撮影スポットを観光スポットとして推薦することを考えた場合、その地域と期間に含まれる写真群が混沌とした未整理状態であるとわかりにくいという問題がある。

例えば、撮影地点がほぼ同じでも、撮影対象が異なる場合、複数の撮影対象が混在する。また、イベントなどにおいても、撮影地点はほぼ同じであるが、時間帯により異なる対象が興味の対象になる場合が考えられる。また、熊野らは、旬のある時期について、日を基準とした期間に関する結果を示したが、日の出は朝の時間帯、夜景などでは夜の時間帯が旬となる場合が考えられ、時間帯の違いによる撮影対象が混在する。つまり、撮影場所や時間帯、撮影対象の違いによる混沌とした状態を整頓することで、ホット撮影スポット内のサブ撮影スポットが明瞭になることが期待される。

Naaman らは、ひとりの撮影者による地理および時間情報付き写真データを地域およびイベントの観点から

整理する研究を行っている [Naaman 04]。しかし、ホット撮影スポットでは、複数の異なる嗜好を持つ撮影者による多様な写真が混在しているため、ホット撮影スポットを明瞭に説明することが困難であった。本研究では、ホット撮影スポットの写真群について、写真の撮影位置、撮影時間と、画像特徴に基づいてホット撮影スポットを明瞭に分類することで、ホット撮影スポットを説明するアノテーション付与法を提案する。

そこで、まず 2 章で写真共有サイトについて説明し、3 章でホット撮影スポット抽出法について述べる。また、4 章で提案法を述べ、5 章で Flickr から収集した日本全土を含む写真の実データを用いた実験と考察を行い、提案法の有効性を示す。そして、6 章でまとめる。

2 写真共有サイトと写真属性

写真共有サイトは、ユーザどうしがコミュニケーションを取る意味で、ソーシャルネットワーキングサイト (SNS) の一つであるが、主にユーザが撮影した写真を実世界の事象を捉えた情報源と見なせば、ユーザが生み出した写真を通じて情報を発信することから、ソーシャルメディアであるとも言える。2012 年現在、世界中に、数多くの写真共有サイトが存在するため、Web 空間には膨大な写真が蓄積され続けている。一方、デジタル写真には、EXIF (Exchangeable image file format) 情報が埋め込まれており、撮影時刻などの基本的な情報の他に、GPS 情報や撮影条件などの多数の情報が含まれている。Flickr は、主要写真共有サイトの一つであるが、写真の EXIF 情報が自動登録されたり、ユーザが付与したタグ情報など、多数の属性情報も共有化して閲覧できる機能を持つ。

図 1 は、Flickr に登録された各年ごとの写真総数である。図 1 の黒線は、登録総数の変遷であり、青線は、日本列島が含まれる地域の位置を指定して得られる写真総数である。ここで、Flickr から得られる位置に、少なくとも二つの観点による位置情報が含まれることを説明しておく。一つ目は GPS に基づいた撮影位置である。デジタル写真に埋め込まれたメタ情報に GPS に基づく撮影位置情報がある場合、Flickr では位置情報の登録拒否を設定していないかぎり、写真登録時にジオタグとして自動記録され、変更できなくなる。二つ目は、Flickr ユーザが地図ツール上の位置を任意に指定して登録した位置であり、登録された写真に一度も Geo-tag が記録されていない場合に設定が可能となる。二つ目の場合でも撮影位置を登録することはできるが、ユーザに一任されるため、全く関係のない位置を登録することもできる。つまり、Flickr の写真に付随する位置情報は、必ずしも撮影位置を示さないこともあるが、EXIF 情報に記録される GPS 情報は、撮影時点で記録される

ため、情報を改ざんしないかぎり、撮影位置を示すと言える。図 1 の赤線は、日本列島が含まれる地域で、EXIF 情報を持つ写真総数の変遷である。本研究では、写真の属性情報として、撮影位置と撮影時間を用いる。

3 ホット撮影スポット抽出

3.1 ホット撮影スポットとホットスポット写真

正の整数 T に対して、 T 日の期間 $[1, T]$ 内に撮影された写真データ全体の集合を、

$$\mathcal{D}_T = \{d_n; n = 1, \dots, N\}$$

とする。ここに、各写真データ d_n には、Geo-tag 情報 x_n 、時間情報 t_n が付随しており、そのことを明記するために、

$$d_n = (x_n, t_n), \quad (n = 1, \dots, N)$$

と記述する。ただし、 $x_n = (x_{n,1}, x_{n,2})$ であり、 $x_{n,1}$ と $x_{n,2}$ はそれぞれ写真 d_n が撮影された緯度と経度、 t_n は d_n が撮影された日、 N は写真データの総数である。

緯度と経度の情報を用いれば、地球表面上の点は 2 次元 Euclid 空間 $\mathbf{R}^2$ 内の領域

$$\Omega = [-\pi/2, \pi/2] \times [-\pi, \pi] \subset \mathbf{R}^2$$

上の点と同一視される。写真データ集合 $\mathcal{D}_T$ から、多くの写真が撮影される人気撮影スポットが近接して存在する地域 $R_m \subset \Omega$, ($m = 1, \dots, M$) を抽出し、その地域において格別の期間 $I_m = [T_{m,0}, T_{m,1}]$, ($m = 1, \dots, M$)、すなわち、他の地域と比較して顕著に人々がその地域で写真を撮影している期間を検出する。各 R_m を主要撮影地域、 I_m を R_m のホット撮影期間と呼ぶ。ここに、 M は抽出した主要撮影地域の総数であり、 R_m は半径 h_0 のある円板近傍に含まれる Ω 内の領域、 $1 \leq T_{m,0} < T_{m,1} \leq T$, ($m = 1, \dots, M$) である。ただし、 $h_0 (> 0)$ は、主要撮影地域のサイズを規定するパラメータである。ここで、 R_m と I_m のペア (R_m, I_m) をホット撮影スポットと呼び、与えられた T 日間の写真データ集合 $\mathcal{D}_T$ からホット撮影スポット群 $\{(R_m, I_m); m = 1, \dots, M\}$ を抽出する。

また、($m = 1, \dots, M$) に対し、地域 R_m 内で期間 I_m に撮影された写真群、つまり、 D_m に属する写真をホット撮影スポット m のホットスポット写真と呼ぶ。

$$D_m = \{d_n = (x_n, t_n) \in \mathcal{D}_T; x_n \in R_m, t_n \in I_m\},$$

3.2 ホット撮影スポットの数理モデル

次に、ホット撮影スポットの数理モデルを示す。写真データ集合 $\mathcal{D}_T$ から主要撮影地域 R_m , ($m = 1, \dots, M$)

を抽出する手法において、人々が写真をどの場所で撮影するのかに関する確率分布に対して、その確率密度関数を極大にする点の近傍が主要撮影地域であるとモデル化する。ただし、極大値が比較的小さいものについては、主要撮影地域とは考えないことにする。そのような確率密度関数の推定に対して、ノンパラメトリックアプローチであるカーネル密度推定

$$\hat{p}(x) = \frac{1}{Nh^2} \sum_{n=1}^N G(\| (x - x_n) / h \|^2), \quad (x \in \mathbf{R}^2) \quad (1)$$

を考える。ここに、 $\|\cdot\|$ は $\mathbf{R}^2$ の Euclid ノルム、 $G(s)$ はカーネル関数であり、Epanechnikov カーネルや Gaussian カーネルなどを利用する。また、 $h (> 0)$ は、主要撮影地域のサイズを規定するパラメータとして、対象とする問題のスケール（解像度）に応じてユーザが事前に指定するものとする。

我々は、Crandall らの研究 [Crandall 09] に従い、 $\mathcal{D}_T$ に属する各写真の撮影場所 x_n , ($n = 1, \dots, N$) を初期値としてミーンシフト法を適用し、式 (1) の確率密度関数 $\hat{p}(x)$ の極大値を与える点を推定するとともに、 $\mathcal{D}_T$ に属する写真のクラスタリングを行う。 $\hat{p}(x)$ の極大値を与える点として推定されたもの全体を $\{\hat{c}_m; m = 1, \dots, M'\}$ とし、各 m に対して $\hat{c}_m$ に収束した x_n , ($n = 1, \dots, N$) の全体を、

$$X_m = \{x_{n(m,j)}; j = 1, \dots, N_m\}, \quad (m = 1, \dots, M')$$

とする。ただし、 $|X_1| \geq \dots \geq |X_{M'}|$ とする。ここで、 $|X_m| \geq \mu_0$ を満たす $m \in \{1, \dots, M'\}$ の最大値 M を求める。ここに、 μ_0 はユーザが指定するパラメータである。次に、各 $m \in \{1, \dots, M\}$ に対して、 $\hat{c}_m$ を中心とし X_m を含む最小の円板近傍と領域 Ω との共通部分 R_m を求める。そして、 $\{R_1, \dots, R_M\}$ を主要撮影地域として出力する。抽出された各主要撮影地域 R_m に対して、そのホット撮影期間 $I_m = [T_{m,0}, T_{m,1}]$ を検出する手法を提案する。ここに、 $T_{m,0}$ と $T_{m,1}$ は $T_{m,0} < T_{m,1}$ なる T 以下の自然数である。

任意の $m \in \{1, \dots, M\}$ に対して、 $q_m(t)$ を R_m 内で第 t 日に撮影された写真の数とし、各 $q_m(t)$ が

$$q_k(t) = q_k^*(t) + q_0(t) \quad (2)$$

のように分解されるとモデル化する。ここに、 $q_0(t)$ は m に依存しない正整数で、地域によらず一般的に第 t 日に撮影される写真数を表す確率変数である。また、 $q_m^*(t)$ は、地域 R_m に特徴的な撮影動向を表すもので、通常の日には m によって異なる正定数値 $w_{m,0}$ をとり、ホット撮影期間 I_m において $w_{m,0}$ より大きい正定数値をとる階段関数である。ただし、各 R_m に対して、ホット撮影期間 I_m は複数個（例えば、 $I_{m,1}, I_{m,2}, \dots$ ）存在し得るが、それらの任意の 2 つの交わりは空集合である。ま

た、 $m \neq m'$ ならば、 R_m と $R_{m'}$ のホット撮影期間は一致しないとする。

任意の主要撮影地域に対して、そのホット撮影期間の候補全体は $\mathcal{J} = \{J = [T_0, T_1]; T_0, T_1 \in \mathbb{Z}, 1 \leq T_0 < T_1 \leq T\}$ であり、それらを、

$$\mathcal{J} = \{J_i; i = 1, \dots, T(T-1)/2\}$$

と番号づけする。ここで、各 R_m におけるホット撮影期間（すなわち、他の地域と比較して顕著に多数の写真が撮影された期間）を効率的に検出するために、撮影された写真の数に関して、地域 R_m , ($m = 1, \dots, M$) と期間 J_i , ($i = 1, \dots, T(T-1)/2$) の独立性を検定する。具体的には、まず Fisher 直接確率検定に従って、 R_m と独立性が低い（すなわち、Fisher 直接確率の値が小さい）期間を候補 $\mathcal{J}$ から探索する。ところで、 R_m に対する Fisher 直接確率の値が小さい期間は、他の地域と比較して顕著に少数の写真が撮影された期間という場合もあり得るので、Fisher 直接確率検定で検出された期間に対して、さらにその期間で撮影された写真数をも考慮し、 R_m におけるホット撮影期間を抽出する。

3.3 ホット撮影期間の抽出

まず、Fisher 直接確率検定に従って、地域 R_m , ($m = 1, \dots, M$) と期間 J_i , ($i = 1, \dots, T(T-1)/2$) の独立性を検定する。表 1 のような R_m と J_i に関する 2×2 分割表を考えよう。ここで、 N は写真の総数、 p_m は領域 R_m に属する写真の数、 p'_i は期間 J_i に含まれる写真の数、 $p_{m,i}$ は R_m に属し J_i に含まれる写真の数、 $p_{\bar{m},i}$ は R_k に属し J_i に含まれない写真の数、 $p_{\bar{m},i}$ は R_k に属さず J_i に含まれる写真の数、 $p_{\bar{m},i}$ は R_k に属さず J_i に含まれない写真の数を、それぞれ表す。このとき、

$$p_{k,i} + p_{\bar{m},i} = p_m, \quad p_{\bar{m},i} + p_{\bar{k},i} = N - p_m,$$

$$p_{m,i} + p_{\bar{m},i} = p'_i, \quad p_{k,i} + p_{\bar{k},i} = N - p'_i$$

である。Fisher 直接確率検定では、Fisher 直接確率

$$F_{m,i} = \sum_{j=p_{m,i}}^{\min(p_m, p'_i)} \frac{\binom{p_m}{j} \binom{N-p_m}{p'_i-j}}{\binom{N}{p'_i}} \quad (3)$$

表 1: 2×2 分割表

	J_i	$\bar{J}_i$	
R_k	$m_{k,i}$	$m_{\bar{k},i}$	m_k
$\bar{R}_k$	$m_{\bar{k},i}$	$m_{\bar{k},i}$	$N - m_k$
	m'_i	$N - m'_i$	N

が大きいほど、 R_m と J_i の独立性が高いと検定される。我々は、各 R_m に対して $p_{m,i} \geq \phi_m$ なる J_i を、Fisher 直接確率 $F_{m,i}$ の小さい順に「 $I_{m,1}, I_{m,2}, \dots$ 」とランキングし、「 $I_{m,1}$ を R_m の第 1 ホット撮影期間、 $I_{m,2}$ を R_m の第 2 ホット撮影期間、 $\dots$ 」として抽出する。ここに、 $\phi_m (> 0)$ はユーザが指定するパラメータである。

Fisher 直接確率 $F_{m,i}$, ($m = 1, \dots, M; i = 1, \dots, T(T-1)/2$) は、原理的には式 (3) に従ってナイーブに直接計算することにより求めることが可能だが、 N と T が大きくなると膨大な計算量が必要になると考えられる。そこで我々は、

$$f(\ell, j) = \log \left(\frac{\ell}{j} \right), \quad (\ell = 1, \dots, N; j = 0, 1, \dots, \ell)$$

を、漸化式

$$f(\ell, j) = \begin{cases} 0 & (j = 0) \\ f(\ell, j-1) + \log(\ell - j + 1) - \log(j) & (j \geq 1) \end{cases} \quad (4)$$

を用いて計算し、これらから Fisher 直接確率 $F_{k,i}$ を、

$$F_{m,i} = \sum_{j=p_{m,i}}^{\min(p_m, p'_i)} \exp(f(p_m, j) + f(N - p_m, p'_i - j) - f(N, p'_i)) \quad (5)$$

と計算することを提案する。式 (5) では指数値を計算する前に対数値 $f(\ell, j)$ の加算および減算を実行しているので、提案法が計算精度の劣化を抑制することも期待できる。以下に、ホット撮影スポットの抽出アルゴリズムを示す。

Algorithm 提案法

- 1: p_m を求める;
- 2: 式 (4) から $\{f(p_m, i); 1 \leq i \leq p_m\}, \{f(N - p_m, i); 1 \leq i \leq N - p_m\}, \{f(N, i); 1 \leq i \leq N\}$ を求める;
- 3: $i := 1$;
- 4: **while**($i \leq T(T-1)/2$) **do** /*期間 J_i の $F_{m,i}$ を計算*/
- 5: R_m と J_i に関する 2×2 分割表を構築 (表 1 を参照);
- 6: **if** $p_{m,i} < \phi_m$ **then**
- 7: **goto** step10;
- 8: **end if**
- 9: 式 (5) から $F_{m,i}$ を求める;
- 10: $i := i + 1$;
- 11: **end while**
- 12: $F_{m,i}$ の小さい順に J_i をランキングし、 R_m におけるホット撮影期間 $I_{m,1}, I_{m,2}, \dots$ を出力;

さらに我々は、Fisher 直接確率 $F_{m,i}$ に基づいて主要撮影地域とホット撮影期間候補のペア (R_m, J_i) をランキングすることにより、与えられた T 日間の写真データ集合 $\mathcal{D}_T$ から、格別なホット撮影スポット群を抽出する。

4 提案法

4.1 概要

本研究では、写真群を明瞭に分類する上で、各写真に付随する Geo-tag、画像データ、撮影時間を用いて、1. Geo-tag 情報、2. 画像特徴情報、3. 時間差情報の三つの情報に着目し、ホット撮影スポット m ごとに、未整理のホットスポット写真群 D_m を三情報の組み合わせによる類似度に基づいてクラスタリングすることで、ホット撮影スポットを説明するアノテーションを付与する手法を提案する。

4.2 アノテーション

Geo-tag による実距離の近さのみによって写真群をクラスタリングし、色分けラベルを付けて地図上に配置すると、サブ撮影スポットの存在が明瞭になる。また、画像特徴の近さを用いてクラスタリングし、色分けラベルを付け、Geo-tag 情報を用いて地図上に配置すると、異なる場所に同じような写真が撮影可能なサブ撮影スポットの存在を示すことができる。さらに、撮影時間の近さを用いれば、朝の撮影スポット、夜の撮影スポットなどが異なるクラスとして分類できる。

一方、Geo-tag 情報と画像特徴情報の両方の近さを複合してラベル付けをすれば、ほぼ同じ撮影位置でも、複数のクラスに分割される可能性があり、撮影対象の違いによる分類が可能になる場合が考えられる。さらに、Geo-tag 情報、画像特徴情報、撮影時間差情報の三つを統合すれば、同じ場所でも、異なる対象、異なる時間帯のサブ撮影スポットであるというラベル付けが可能になると思われる。

つまり、撮影位置の近さ、写真に映る画像の近さ、撮影時間の近さの指標の組み合わせに基づいて、分類を行った写真群が、うまく撮影スポットの特徴を表しているならば、クラスに付与されたラベルが、写真群を説明することから、本研究では、そのラベルをアノテーションと見なす。以下に、三つの情報による類似度の算出法と組み合わせ法を示す。

4.2.1 Geo-tag 情報に基づく類似度

ホット撮影スポット m 内の写真群 D_m について、各写真間の実距離に基づく類似度 S_G が $0 \sim 1$ 内になるよう正規化を行う。そこで、 $dist_G(u, v)$ ただし $(u \neq v)$ を d_u と d_v との Geo-tag に基づく距離とし、その最大値を $dist_G^{max}$ 、最小値を $dist_G^{min}$ としたとき、 S_G を次のように定義する。

$$S_G(u, v) = 1 - (dist_G(u, v) - dist_G^{min}) / (dist_G^{max} - dist_G^{min})$$

4.2.2 画像特徴に基づく類似度

近年、画像の特徴を表現する手法として、Visual words (keypoints) が注目されている [Csurka 04]。本研究では、計算速度の速い SURF (Speeded-Up Robust Features) 特徴量 [Bay 08] をホット撮影スポット内の写真群 D_m から抽出し、 K -means クラスタリングを適用し、得られたセントロイドを Visual Word とすることで、個々の写真を K 次元の Visual words ヒストグラムで表現する。ここで、 $sim_I(u, v)$ ただし $(u \neq v)$ を d_u と d_v との $\cos$ 類似度とし、その最大値を sim_I^{max} 、最小値を sim_I^{min} としたとき、画像の類似度 S_I を次のように定義する。

$$S_I(u, v) = (sim_I(u, v) - sim_I^{min}) / (sim_I^{max} - sim_I^{min})$$

4.2.3 撮影時間差に基づく類似度

異なる二枚の写真 u と v ただし $(u \neq v)$ が撮影された時間を PT_u 、 PT_v とするとき、撮影時間の近さを表すため、24 時間を円周に割り当て、 PT_u と PT_v の時間差を、円周の時計回りと、反時計回りでそれぞれ求め、差が小さい方の絶対値を撮影時間差 $diff_T(u, v)$ とする。ここで、その最大値を $diff_T^{max}$ 、最小値を $diff_T^{min}$ としたとき、画像の類似度 S_T を次のように定義する。

$$S_T(u, v) = 1 - (diff_T(u, v) - diff_T^{min}) / (diff_T^{max} - diff_T^{min})$$

4.2.4 類似度の組み合わせ

本研究では、三つの類似度を組み合わせた合成類似度 S' を以下の組み合わせで統合した。

1. $S'_G = S_G / w$
2. $S'_I = S_I / w$
3. $S'_T = S_T / w$
4. $S'_{G,I} = S_G / w + S_I / w$
5. $S'_{G,T} = S_G / w + S_T / w$
6. $S'_{I,T} = S_I / w + S_T / w$
7. $S'_{G,I,T} = S_G / w + S_I / w + S_T / w$

ただし、類似度が一つの場合 $w = 1$ 、二つ統合する場合 $w = 2$ 、三つ統合する場合 $w = 3$ とする。

4.3 クラスタリング

本研究では、 D_m からクラスを抽出するクラスタリング手法として、クラス数を自動決定できる Newman クラスタリング [Clauset 04] を用いる。この手法は、大規模な複雑ネットワークのグラフに内在するコミュニティ構造を高速に抽出する手法であるため、グラフ表現されたネットワークを構築する必要がある。そこ

で、 D_m 内の写真をノードとし、異なる二つの写真の類似度 S' を重みとするリンクを張り、完全グラフを構築して、閾値 S_h 以下のリンク削除する。得られたグラフ構造のうち、最大連結成分 G をクラスタリングの対象として、サブ撮影スポットを抽出する。

Newman クラスタリングでは、グラフ内に潜在的に存在するコミュニティ i 内のリンク密度が高く、コミュニティ間のリンク密度が低い状態を良いクラスタ構造と見なす。クラスタリングの精度を測るモジュール性指標 (modularity) $Q > 0$ を導入し、 Q が 0 になるランダムネットワークに対し、コミュニティ構造の存在を反映した Q を最大にするクラスタ数を自動決定できる。

$$Q = \sum_i (e_{ii} - a_i^2)$$

ここで、 e_{ii} は「コミュニティ i 内のノードと i 内の別のノード間のリンク総和」の「総リンク数」に対する割合を表し、 a_i は「コミュニティ i 内のノードから出ているリンク総和」の「総リンク数」に対する割合である。また、 $\Delta Q_{ij} = 2(e_{ij} - a_i a_j)$ を計算することで、高速化も行われている。本研究では、重み付きネットワーク G に Newman クラスタリング [Zhou 10] を適用する。

5 実験

写真共有サイト Flickr から収集した大量の地理情報および時間情報の付随する写真データを用いて、ホット撮影スポットの抽出を行い、ホット撮影スポットごとにクラスタの抽出実験を行った。

5.1 実験データ

日本国内のホット撮影スポットでアノテーション付与を行うため、写真共有サイト Flickr から、日本列島が含まれる矩形領域 (緯度:25.8 ~ 45.8, 経度:126.2 ~ 146.8) に含まれる 2010 年 1 月 1 日から 2010 年 12 月 31 日までの 1 年間の撮影位置・撮影時間付き写真データを収集した。ただし、日本国内に焦点を当てるため、矩形領域に入り込む他国の写真データを除いた。その結果、548,922 枚の写真データが得られた。また、本研究では、空間クラスタリングにおいて、Yin[Yin 11] らに従い、最小領域で一人の撮影者が何度撮影しても 1 度として数えた。本研究では、最小領域を 1 辺 10m の矩形領域とした。その結果、写真数 162,933 枚のデータセットとなった。

実験の前段階として、(ホット撮影スポット) の項で説明した手順に従って h を 100m として空間クラスタリングを行った。その結果、 $K'=24,954$ 箇所の主要撮影地域が得られた。 $\mu_0=100$ としたところ、 $K=205$ となっ

た。続いて、fisher 検定で各地域のホット撮影期間を検出してホット撮影スポットを得た。全ホット撮影スポットの写真数 $|D_T|$ は 10,185 枚となった。本研究では、この D_T を実験用のデータセットとする。

5.2 クラスタリングの比較

Visual Word の語数を決める $K=1000$ として Visual Words ヒストグラムを写真 d_n ごとに算出した。また、三つの類似度を用いて組み合わせた七つの類似度 S' をそれぞれ用いて、七種類の完全グラフを構築し、 $S' < S_h$ となるリンクを切断し、それぞれ最大連結成分をクラスタリングの対象とした。本研究では、リンクを切断する閾値 $S_h=0.8$ として実験を行った。

5.2.1 評価法

まず、本研究では、三つの情報による $S'_{G,I,T}$ を用いたクラスタリングによるクラスタへのラベル付けを提案法と定め、クラスタ C_k , ($k = 1, \dots, N_k$) を正解データとする。ここで、三つの情報を用いることの効果を評価するため、他の S' を用いた場合と相違があるか否かを調べる。尚、 S' の異なるクラスタリングでは、それぞれ得られるクラスタ数が異なることが予想されるが、本研究では、Micro average precision を用いることで、クラスタ数の違いによる問題を解消する。一方、他の組み合わせの特徴量を使ったクラスタリングの結果を比較対象として C'_l , ($l = 1, \dots, L_l$) とする。このとき、比較するデータのクラスタ C'_l に対して交わり具合が最大となる C_k との適合率を C'_l の得点 $f(l)$ とする。

$$f(l) = \max |C'_l \cap C_k|$$

そして、次式で比較するデータの正解データに対する Micro average precision P_{ma} を算出する。

$$P_{ma} = \sum_{n=1}^L \frac{f(l)}{N}$$

5.2.2 クラスタリング結果の評価

$S'_{G,I,T}$ に対する他の S' についての P_{ma} について、205 のホット撮影スポットで算出した値の平均値を表 2 に示す。表 2 より、最もクラスタの傾向が近いのは「画像特徴」+「時間情報」を組み合わせた $S_{I,T}$ であるが、それでも 0.62 程度の値であり、他の S' を含め、得られたクラスタ群はいずれとも異なっている独自性があることがわかる。つまり、三つを組み合わせた提案法は、部分的に共通する情報を用いていても、他の S' を用いた場合と異なるクラスタを形成していると言える。

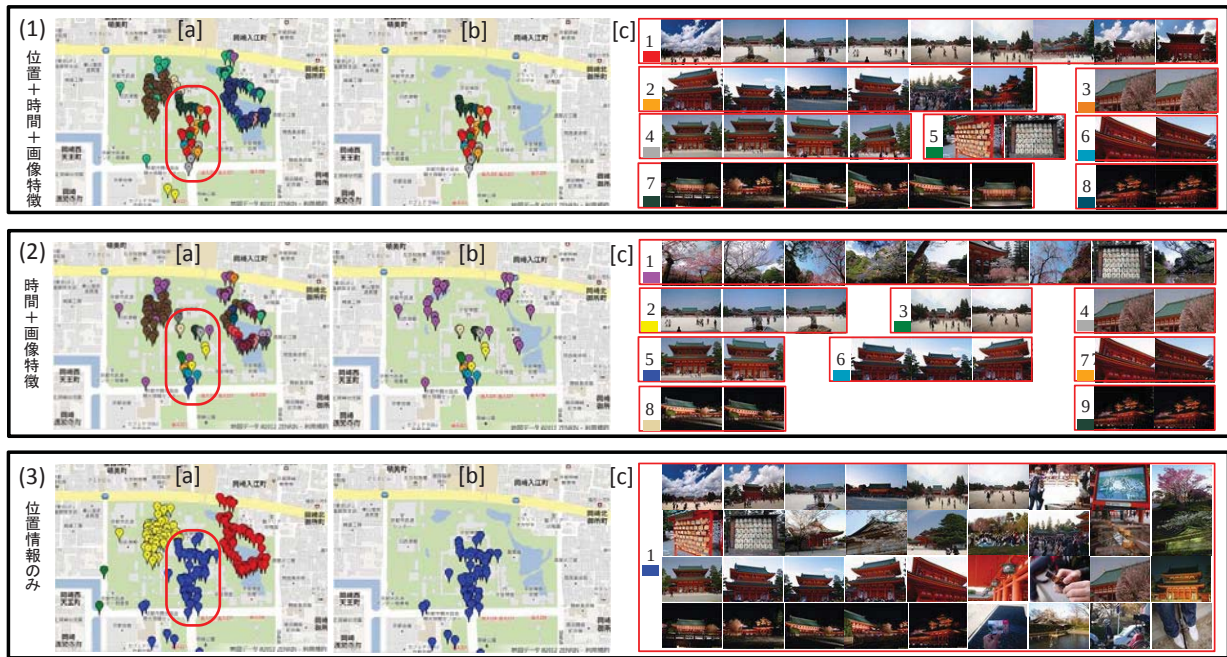


図 2: 提案法 (1) と比較法 (2)(3) (ホット撮影スポット：京都、平安神宮)

表 2: micro average precision の平均値ランキング

	画像 + 位置 + 時間情報に対する値
画像 + 時間情報	0.628234731
位置情報のみ	0.538383589
位置 + 時間情報	0.537968227
画像 + 位置情報	0.526942926
時間情報のみ	0.441309438
画像情報のみ	0.409779561

5.3 実験結果

図 2(1) は、三つの情報を用いてクラスタリングした提案法の結果を抜粋した一例である。図 2(1)[a] は、クラスタリングの結果、ホット撮影スポットの一つである、京都、平安神宮周辺のホットスポット写真群を Geo-tag を用いて Google Map に配置したものである。説明上、図 2(1)[a] の赤枠の地域に着目する。図 2(1)[b] は、その赤枠の中に全部もしくは一部が含まれる 8 クラスタであり、図 2(1)[c] は、その 8 クラスタの写真を抜粋したものである。平安神宮は、日中に撮影されることも多いものの、ライトアップがあり、夜景も数多く撮影される観光スポットである。位置情報により、クラスタは近い位置関係のものが抽出されているが、ほぼ同じ撮影地域でも、画像特徴の効果から異なる撮影対象が別のクラスタに分かれ、写真の印象が近いものが同じクラスタに含まれる様子がわかる。また、ほぼ同じ撮影地域でも、昼と夜の写真はよく分離されており、同

じ印象の写真が同クラスタに含まれることから、提案法は、うまく分類され、何が撮影できるかが明瞭に理解しやすいクラスタを形成していることがわかる。

一方、図 2(2) は、提案法に最も近かった「撮影時間」と「画像特徴」の近さを用いたクラスタリングの結果である。図 2(2)[a] では、図 2(1)[a] とほぼ同じ地域に赤枠を設定し、その赤枠の中に全部もしくは一部が含まれる 9 クラスタを示したものが図 2(2)[b] である。図 2(2)[b] 中、紫色のクラスタが、赤枠内に一部存在しているが、紫色のクラスタは、平安神宮のあちこちに分散している。これは「位置」の近さが反映されていないことが原因であると思われる。この紫色のクラスタは、図 2(2)[c] の 1 の赤枠に対応しており、同クラスタ内に含まれる写真は、位置が異なっている、同じような印象の写真が集まっていることがわかる。これは、紫色クラスタのような写真を好む者からすれば、平安神宮にいくつも同じような写真が撮影できるスポットがあることを示しているとも考えることも可能である。しかし、図 2(2)[c] の (2) や (3) は、位置の情報がないためか、実際のところ極めて近い撮影位置であるにもかかわらず、画像特徴がうまく働かなかったためか、過分割された例である。この点においては、提案法では、図 2(1)[c] の 1 クラスタに対応しているが、位置情報の効果からか、過分割が行われず、一つのクラスタとなっており同様の印象の写真が整理されている印象がある。

また、図 2(3) は、提案法に対し、2 番目に近かった位置情報のみを使った比較法である。これまでと同様に、図 2(3)[a] で、他とほぼ同じ地域を抽出した結果が

図 2(3)[b] であるが、これは図 2(3)[c] の 1 クラスタのみが抽出された。このクラスタは、時間帯の違い、撮影対象の違いに関係なく、多様な写真が混在して、どのようなクラスタなのか説明しにくい印象を受ける。以上から、提案法の有効性が伺える。

6 まとめ

本研究では、写真データに付随する Geo-tag 情報と撮影時間情報、さらに画像特徴に着目して三つの情報を統合することで、ホット撮影スポットをうまく分類し、提案法の有効性を示した。今後は、より効果的な整理法を探索するとともに、観光スポットを推薦する上で、デジタル写真に付随する数多くの EXIF 情報などを用いて、より有効な整頓法を探索する予定である。

参考文献

- [Arase 10] Arase, Y., Xie, X., Hara, T., and Nishio, S.: Mining People's Trips from Large Scale Geo-tagged Photos, in *Proceedings of the 18th International Conference on Multimedia*, pp. 133–142 (2010)
- [Bay 08] Bay, H., Ess, A., Tuytelaars, T., and Van Gool, L.: Speeded-Up Robust Features (SURF), *Comput. Vis. Image Underst.*, Vol. 110, No. 3, pp. 346–359 (2008)
- [Clauset 04] Clauset, A., Newman, M. E. J., and Moore, C.: Finding community structure in very large networks, *Physical Review E*, pp. 1–6 (2004)
- [Comaniciu 02] Comaniciu, D. and Meer, P.: Mean shift: a robust approach toward feature space analysis, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 24, No. 5, pp. 603–619 (2002)
- [Crandall 09] Crandall, D. J., Backstrom, L., Huttenlocher, D., and Kleinberg, J.: Mapping the world's photos, in *Proceedings of the 18th International Conference on World Wide Web*, pp. 761–770 (2009)
- [Csurka 04] Csurka, G., Dance, C. R., Fan, L., Willamowski, J., and Bray, C.: Visual categorization with bags of keypoints, in *In Workshop on Statistical Learning in Computer Vision, ECCV*, pp. 1–22 (2004)
- [Kisilevich 10a] Kisilevich, S., Mansmann, F., Bak, P., Keim, D. A., and Tchaikin, A.: Where Would You Go on Your Next Vacation? - A Framework for Visual Exploration of Attractive Places, in *GeoProcessing 2010*, pp. 21–26 (2010)
- [Kisilevich 10b] Kisilevich, S., Mansmann, F., and Keim, D. A.: P-DBSCAN: A density based clustering algorithm for exploration and analysis of attractive areas using collections of geo-tagged photos, in *1st International Conference on Computing for Geospatial Research & Application* (2010)
- [Lu 10] Lu, X., Wang, C., Yang, J.-M., Pang, Y., and Zhang, L.: Photo2Trip: generating travel routes from geo-tagged photos for trip planning, in *Proceedings of the 18th International Conference on Multimedia*, pp. 143–152 (2010)
- [Naaman 04] Naaman, M., Song, Y. J., Paepcke, A., and Garcia-Molina, H.: Automatic Organization for Digital Photographs with Geographic Coordinates, in *Proceedings of ACM/IEEE-CS JCDL Joint Conference on Digital Libraries*, pp. 53–62 (2004)
- [Yin 10] Yin, H., Lu, X., Wang, C., Yu, N., and Zhang, L.: Photo2Trip: an interactive trip planning system based on geo-tagged photos, in *Proceedings of the 18th International Conference on Multimedia*, pp. 1579–1582 (2010)
- [Yin 11] Yin, Z., Cao, L., Han, J., Zhai, C., and Huang, T.: Geographical Topic Discovery and Comparison, in *Proceedings of the 20th International Conference on World Wide Web*, pp. 247–256 (2011)
- [Zhou 10] Zhou, T. C., Ma, H., Lyu, M. R., and King, I.: UserRec: A User Recommendation Framework in Social Tagging Systems, in Fox, M. and Poole, D. eds., *AAAI*, AAAI Press (2010)
- [王 11] 王 佳な, 野田 雅文, 高橋 友和, 出口 大輔, 井手 一郎, 村瀬 洋: Web 上の大量の写真に対する画像分類による観光マップの作成, *情報処理学会論文誌*, Vol. 52, No. 12, pp. 3588–3592 (2011)
- [熊野 12] 熊野 雅仁, 小関 基徳, 小野 景子, 木村 昌弘: 地理および時間情報をもつ写真データに基づいたホット撮影スポットの抽出, *情報処理学会論文誌*, Vol. 5, No. 3, pp. 41–53 (2012)
- [松原 11] 松原 仁: 特集: 「観光と知能情報」にあたって, *人工知能学会誌*, Vol. 26, No. 3, p. 225 (2011)
- [川村 10] 川村 秀憲, 鈴木 恵二, 山本 雅人, 松原 仁: 観光情報学, *情報処理*, Vol. 51, No. 6, pp. 642–648 (2010)
- [味八木 10] 味八木 崇, 暦本 純一: 集合知センシングによる実世界インタフェース, *情報処理*, Vol. 51, No. 7, pp. 775–781 (2010)

没入型3Dメタバーズ情報空間による観光支援システム

Exploring Immersive 3D Information Metaverse for Planning Sightseeing Tours

岩渕 聡¹ 熊野 雅仁^{2*} 亀井 貴行¹ 小野 景子² 木村 昌弘²
Satoshi Iwabuchi¹ Masahito Kumano² Takayuki Kamei¹
Keiko Ono² Masahiro Kimura²

¹ 龍谷大学大学院理工学研究科電子情報学専攻

¹ Division of Electronics and Informatics, Graduate School of Science and Technology, Ryukoku University

² 龍谷大学 理工学部 電子情報学科

² Department of Electronics and Informatics, Faculty of Science and Technology, Ryukoku University

Abstract:

Recently, considerable attention has been devoted to planning sightseeing-tours by using a large number of geotagged photographs and restaurant reviews in Social Media. We have developed such a sightseeing-tour planning system that 1) extracts hot photo-spots from a set of geotagged photographs with timestamps, 2) classifies and embeds them in a 3D information space on the basis of image similarity of photographs, 3) visualizes restaurant reviews in the 3D information space, and 4) provides an immersive interaction environment in which a user explores the 3D information space in order to find his/her favorite hot photo-spots and plans his/her own sightseeing tour. Using real data from “Flickr” and “Tabelog”, we examined the effectiveness of the sightseeing-tour planning system.

1 はじめに

これまで、観光産業が、地球規模の観光はもちろん国内の地域観光においても、独自の調査と専門的な知識に基づいて、観光案内情報を提供してきたが、Web空間に電子化された観光情報が溢れるに従い、自ら情報を収集し、観光旅行計画を立てるユーザが増えつつある。ただし、Web空間に蓄積された情報は極めて膨大であるため、ユーザとその要求を満たす対象とを引き合わせる方法として、検索システムが普及し、ユーザが自ら能動的に情報を探す習慣を浸透させつつある。

検索システムは、異なるユーザが情報を要求する際、同じクエリを用いると、基本的に同じ検索結果を提示する。このとき、検索システムに求められる資質は、同じクエリを用いる多数のユーザをできるだけ満足に導くことであり、主に多くの人に人気のある情報や、普及している情報がその満足に貢献してきた。そして、検索システムの普及以前と比べ、思いがけない対象との

偶然の出会い (Serendipity)¹ の機会が増えたとされる。そこで、検索システムの発展は、ユーザと検索対象との多様な結び付けを高めることが期待される。しかし、その期待に反して、社会学の研究から、検索システムが人気増幅器として働き、むしろ多様性を損なわせるとする一例が報告された [Evans 08]。観光の場合、人気スポットはさらに人気を増すが、その反面、例えば観光地の新しいスポットと潜在的に嗜好の合うユーザが引き合わされないことで、新しいスポットが日の目を見ぬまま消え去るなど、検索システムが、観光地の多様な発展を阻害する方向に働く可能性が考えられる。

一方、Social Mediaの繁栄が多様な引き合わせの促進に貢献し始めている。例えば、人気スポットを訪れた観光客が、その移動中、偶然に興味を引いた対象を見つけ、TwitterやFacebookなどに情報を投稿し、同じ嗜好を持つ人々に伝搬すれば、やがて人気スポット

¹Serendipity は、目的のものを探す過程で偶然意外なものに遭遇し、当初の目的のものは手に入らなかったが、幸福を得たという、童話を語源とする偶然の幸運を意味する語である。研究領域においても、幸運な発見を意味し、インターネット時代には、検索により、ユーザが予期しない、新たなデータと出会う機会が増えたことから注目された。推薦システムにおいては、目新しさに、思いがけなさ、予見のできなさ、または意外性の要素が加わった概念である。

*連絡先： 龍谷大学
滋賀県大津市瀬田大江町横谷 1-5
E-mail:{kumano,kono,kimura}@rins.ryukoku.ac.jp

として成長する可能性がある。このとき、この観光客は実世界の情報を捉える点で Social Sensor と見なせ、Web 空間に情報を発信する点で Social Media の一部とも見なせる。ただし、もし Social Media を通じて投稿されたスポット情報が、うまく伝搬せず、観光業界が注目するスポットとも一致しないなら、より多様なスポット情報が Web 空間内に埋没していることになる。また、Social Media に基づくスポット情報は、多様な嗜好が内在すると考えられ、多数派だけでなく少数派も含まれると思われるが、多様な嗜好は、いくつ存在するか不明である。そこで、いかに多様な嗜好を抽出し、潜在的に嗜好が合致するユーザといかに引き合わせ、Serendipity を助長するかが問題となる。

ユーザと多様な情報を引き合わせる方法として、推薦システムがある [神島 07, 神島 08a, 神島 08b]。推薦システムは、主に、協調フィルタリングと内容に基づくフィルタリングとに分けられる。協調フィルタリングは、メリットとして、他人の知識をうまく利用し、Serendipity を助長することが知られている。しかし、新規ユーザの加入、新規アイテムの追加に関してまだ多くの情報が得られていない場合、推薦の精度が落ちる、いわゆる cold-start 問題があり、その場合、意外性のある推薦ができて、適切な推薦が行えないデメリットがある。一方、内容に基づくフィルタリングは、あらかじめ僅かな嗜好情報が得られれば、メリットとして、推定に基づき cold-start 状態や少数派の嗜好を持つ場合でも比較的良好な推薦が可能である。しかし、嗜好の学習が進むほど、ユーザの過去の情報に引きずられ、予見できる範囲の推薦になりがちとなり、Serendipity の機会も減りがちとなるデメリットが指摘されている。

ここで、ユーザが、まだ訪れたことのない見知らぬ地域に向けて観光計画を行う場合を考える。ユーザにより、人気スポットを求める人もいれば、人ごみを避けて穴場を探す人や、意外なものとの出会いを求める人もいると考えられる。このとき、ユーザは、現地の何に惹かれるか事前にはわからず、何をキーワードとして検索システムを利用すればよいかわからない場合や、要求をうまく言語化できない場合もある。また、意外なものに出会いたいユーザが、自分の過去の嗜好ではなく、新たな嗜好に沿うものを見つけない場合、過去の情報を与えても、適切な推薦が得られないと思われる。本研究では、多様な嗜好を持つユーザと多様な観光先を結びつける点で、以下の指針に着目している。

1. 情報要求が極めて曖昧な状態から始められる
2. キーワードを初めから入力する必要がない
3. あらかじめユーザの嗜好情報を与えなくてもよい
4. cold-start でもよい
5. 幸運な出会い (Serendipity) を助長する
6. 多数派・少数派の区別なく多様な嗜好を考慮する

7. ユーザの母国語は問わない

検索・推薦システムの技術も取り入れつつ、それらの弱点を補い、以上の指針を実現するために、本研究では、Social Media の情報が可視化された空間に没入し、空間を探索して、ユーザの嗜好と合う対象を探しに行く行動を支援するアプローチに着目する。また、探索だけでなく、可視化されたデータ群を対話的に操作、再可視化したり、情報空間と 2 次元地図空間を往来する、探索・データ操作・空間往来型インタラクションの反復により情報を発掘する手法にも着目する。この方法により、ユーザは、キーワードの入力や嗜好情報などをあらかじめ用意する必要はなく、cold-start でもよく、情報要求が極めて曖昧な状態で可視化された情報空間を探索する中で、幸運な対象、多様な対象と出会う可能性が期待される。また、本研究では、可視化される主な情報に写真を用いるため、言語に依存しない。

20 世紀末から計算機技術を用いて複雑ネットワークを可視化する研究 [三末 09] が進み、近年、3 次元情報可視化空間に没入し、視覚的な方法でデータマイニングを行う 3DVDM [Nagel 08] が注目されている。また、Mirror World としての Google Earth 上で時空間データを解析・探索する研究 [Kisilevich 10a] や、観光に関連して、魅力的な場所を視覚的に探索する研究 [Kisilevich 10b] も行われている。この動向は、近年、Augmented Reality、Lifelogging、Mirror Worlds、Virtual Worlds を包含するインターネットを背景とした Metaverse を実現する、次世代 3DWeb [Smart 07] の構想で統合的に組み込まれるものと思われる。本研究は、Social Media の情報を背景とする複雑ネットワークが可視化された Virtual Worlds の一種と見なせる 3D 情報空間への没入と、Mirror Worlds の一種として見なした Google Map とを往来して観光旅行計画を練る Metaverse のプロトタイプ構築を試み、観光旅行計画を支援することで、観光の活性化を促進する観光支援システムの構築を行う。

そこで、2 章では観光支援システムの着眼点と要素技術を説明し、3 章で提案システムを示し、4 章で Social Media Data を用いた実験を行い、5 章でまとめる。

2 観光支援システム

2.1 Social Media に基づく観光スポット情報

近年、観光業界の情報だけでなく、Social Media の情報を捉えることで、さらなる多様な情報獲得が期待できる新しいアプローチとして、Flickr² など写真共有サイトの Geo-tag 付き写真群を用いる研究 [Crandall 09] が脚光を浴びている。写真は、撮影者の心をつかむ対象に遭遇したとき撮影されることが多いことから、写真

²<http://www.flickr.com/>

が単なる記録ではなく、撮影者の何らかの嗜好に基づく意見を内在化させていると考えることができる。つまり、大量の写真群は、意見の集合と見なすことができ、写真の画像から得られる情報や、付随するメタ情報をうまく集約すれば、集合知 [Surowiecki 04] としての、新しい観光スポット情報が得られる可能性がある。

Crandall らは、大量の Geo-tag 付き写真と、写真の画像特徴を用いて空間的なクラスタリングを行い、多くの人が訪れる人気スポットや、ランドマークのある主要地域を示した [Crandall 09]。この研究は、魅力的な地域を探す研究 [Kisilevich 10c] や、写真に付与されたタグ情報も利用することで、地域ごとの地理的トピックを抽出し、地域間を比較して地域の文化的特徴を発見する研究 [Yin 11]、観光マップを作成する研究 [王 11] にも波及している。また、Geo-tag 情報だけでなく、撮影時間というメタ情報を用いることで、旅行する人々の写真撮影行動から旅行行動をマイニングする研究 [Arase 10] や、旅行の計画を支援する研究 [Yin 10]、旅行計画の経路を生成する研究 [Lu 10]、さらに、旬の期間がある格別な地域と期間とのペアをホット撮影スポットと呼んで時空間を抽出する研究 [熊野 12] にも波及している。

これらの Social Media に基づく訪問先に関する情報は、実際に人々が訪れた場所であり、写真を撮影したり、コメントを記載したり、Social Media に投稿するという労力を惜まず、他の人々に見てもらいたい情報である傾向が強いものと考えられるため、Social Sensor が、推薦に値する対象を抽出する Social Filter としても働くことが期待される。また、それらのデータをうまく集約して得られるスポットは、観光業界が伝えるスポットとは異なる多様な場所も含まれていることも期待される。本研究では、熊野らが抽出したホット撮影スポットを対象として、撮影者たちの嗜好に基づいた可視化を行い、3D 情報空間に没入し、探検し、インタラククションを行うことで、嗜好に合致する観光スポットに引き合わせる、観光旅行計画支援法を提案する。

2.2 ホット撮影スポット抽出

正の整数 T に対して、 T 日の期間 $[1, T]$ 内に撮影された写真データ全体の集合を、

$$\mathcal{D}_T = \{d_n; n = 1, \dots, N\}$$

とする。ここに、各写真データ d_n には、Geo-tag 情報 x_n 、時間情報 t_n が付随しており、そのことを明記するために、

$$d_n = (x_n, t_n), \quad (n = 1, \dots, N)$$

と記述する。ただし、 $x_n = (x_{n,1}, x_{n,2})$ であり、 $x_{n,1}$ と $x_{n,2}$ はそれぞれ写真 d_n が撮影された緯度と経度、 t_n は d_n が撮影された日、 N は写真データの総数である。

緯度と経度の情報を用いれば、地球表面上の点は 2 次元 Euclid 空間 $\mathbf{R}^2$ 内の領域

$$\Omega = [-\pi/2, \pi/2] \times [-\pi, \pi] \quad (\subset \mathbf{R}^2)$$

上の点と同一視される。写真データ集合 $\mathcal{D}_T$ から、多くの写真が撮影される人気撮影スポットが近接して存在する地域 $R_m (\subset \Omega)$, ($m = 1, \dots, M$) を抽出し、その地域において格別の期間 $I_m = [T_{m,0}, T_{m,1}]$, ($m = 1, \dots, M$)、すなわち、他の地域と比較して顕著に人々がその地域で写真を撮影している期間を検出する。各 R_m を主要撮影地域、 I_m を R_m のホット撮影期間と呼ぶ。ここに、 M は抽出した主要撮影地域の総数であり、 R_m は半径 h_0 のある円板近傍に含まれる Ω 内の領域、 $1 \leq T_{m,0} < T_{m,1} \leq T$, ($m = 1, \dots, M$) である。ただし、 $h_0 (> 0)$ は、主要撮影地域のサイズを規定するパラメータである。ここで、 R_m と I_m のペア (R_m, I_m) をホット撮影スポットと呼び、与えられた T 日間の写真データ集合 $\mathcal{D}_T$ からホット撮影スポット群 $\{(R_m, I_m); m = 1, \dots, M\}$ を抽出する。

また、($m = 1, \dots, M$) に対し、地域 R_m 内で期間 I_m に撮影された写真群、つまり、 D_m に属する写真をホット撮影スポット m のホットスポット写真と呼ぶ。

$$D_m = \{d_n = (x_n, t_n) \in \mathcal{D}_T; x_n \in R_m, t_n \in I_m\},$$

ホット撮影スポットの抽出においては、 $\mathcal{D}_T$ に属する各写真の撮影場所 x_n , ($n = 1, \dots, N$) を初期値として mean-shift 法 [Comaniciu 02] を適用し、 $\{R_1, \dots, R_M\}$ を主要撮影地域として出力する。次に、抽出された各主要撮影地域 R_m に対して、そのホット撮影期間 $I_m = [T_{m,0}, T_{m,1}]$ を検出する。任意の $m \in \{1, \dots, M\}$ に対して、 $q_m(t)$ を R_m 内で第 t 日に撮影された写真の数とし、各 $q_m(t)$ が

$$q_m(t) = q_m^*(t) + q_0(t) \quad (1)$$

のように分解されるとモデル化する。ここに、 $q_0(t)$ は m に依存しない正整数で、地域によらず一般的に第 t 日に撮影される写真数を表す確率変数である。また、 $q_m^*(t)$ は、地域 R_m に特徴的な撮影動向を表すもので、通常の日には m によって異なる正定数値 $w_{m,0}$ をとり、ホット撮影期間 I_m において $w_{m,0}$ より大きい正定数値をとる階段関数である。任意の主要撮影地域に対して、そのホット撮影期間の候補全体は $\mathcal{J} = \{J = [T_0, T_1]; T_0, T_1 \in \mathbf{Z}, 1 \leq T_0 < T_1 \leq T\}$ であり、それらを、

$$\mathcal{J} = \{J_i; i = 1, \dots, T(T-1)/2\}$$

と番号づけする。ここで、各 R_m におけるホット撮影期間 (すなわち、他の地域と比較して顕著に多数の写真が撮影された期間) を効率的に検出するために、撮影された写真の数に関して、地域 R_m , ($m = 1, \dots, M$) と期間 J_i , ($i = 1, \dots, T(T-1)/2$) の独立性を検定する。具

体的には、まず Fisher 直接確率検定に従って、 R_m と独立性が低い (すなわち、Fisher 直接確率の値が小さい) 期間を候補 $\mathcal{J}$ から探索する [熊野 12]。

本研究では、ホット撮影スポットのホットスポット写真群 D_m を Social Media から得られるデータとして、ホット撮影スポット (観光スポット) に潜在する集合的嗜好を抽出し、クラスタリングする手法を与える。

2.3 潜在的な多重トピックモデルに基づくホット撮影スポットのクラスタリング

近年、画像認識分野において、画像から Visual Word を抽出し、画像を Bag-of-Visual Words による特徴ベクトルで表現する手法が注目されている。つまり、ホットスポット写真群 D_m の写真は、Visual Words で表現された文書と見なすことができる。ただし、撮影者個人が複数の嗜好を持つ場合があり、一つの写真に異なる嗜好が同時、つまり多重に混在して反映される場合が考えられるため、写真に潜在する多重嗜好を表現することが望まれる。また、ある撮影者の多重嗜好は、写真ごとに異なる場合があり、多重嗜好が類似する写真群を大分類できれば、多様な嗜好をカテゴリー化することができる。しかし、多重嗜好の場合、嗜好の成分数は未知数であり、写真群を嗜好に基づいて大分類する場合でも、その分類数は一般的に未知である。本節では、文書 (写真) に潜在する多重トピック (多重嗜好) を推定する潜在的な多重トピックモデルとして、多重トピックの成分数を自動決定できる HDP-LDA に着目し、写真群を大分類する際、クラスタ数を自動決定できる Newman クラスタリングについて説明する。

2.3.1 Visual Words

画像から Visual Words(Keypoints) を抽出する手法では、典型的に SIFT (Scale-Invariant Feature Transform) 特徴が用いられている [Csurka 04]。しかし、本研究では、大量の写真を扱うことから、SIFT 特徴より計算速度が速くなる SURF (Speeded-Up Robust Features)[Bay 08] 特徴を用いる。個々の Visual Word vw は、写真群 $\mathcal{D}_T$ の各写真 d_n から抽出されたすべての SURF 特徴を用いて、あらかじめ与えるクラスタ分割数 K に基づいて K-means[Hartigan 79] クラスタリングを行った結果、得られた各クラスタのセントロイドに対応する。これにより、各写真 d_n 内にある個々の SURF 特徴は K 個存在する Visual Word のいずれかに割り当てられるため、各写真は、 $d'_n = \{vw_{n,1}, \dots, vw_{n,W_n}\}$ のように、 d'_n ごとに異なる Visual Word 数 W_n を持つ Bag-of-Visual Words で表現された文書 d'_n と見なせる。ここで、 d'_n の集合を以下のように定義する。

$$\mathcal{D}'_T = \{d'_n; n = 1, \dots, N\}$$

2.3.2 HDP-LDA

潜在的な多重トピック (多重嗜好) の成分 (トピック) が未知数の教師なしデータ群 $\mathcal{D}'_T$ から、パラメータ学習により未知数を自動的に決定できる潜在多重トピックモデルとして HDP-LDA (Hierarchical Dirichlet Process - Latent Dirichlet Allocation) を用いた文書 (写真) 生成モデルについて考える。ここで、棒折り (Stick-breaking) 過程 [Teh 06] により HDP-LDA の説明を行う。

最初に、 $\beta = (\beta_f)_{f=1}^\infty$ は、ディリクレ過程に基づいて "Griffiths, Engen and McCloskey" (GEM) 分布より生成される。

$$\beta | \gamma \sim \text{GEM}(\gamma),$$

これは、言い換えれば、 $f = 1, 2, 3, \dots$ について、

$$\beta'_f | \gamma \sim \text{Beta}(1, \gamma), \quad \beta_f = \beta'_f \prod_{\ell=1}^{f-1} (1 - \beta'_\ell)$$

と言える。ここで、 $\text{Beta}(1, \gamma)$ は、形状パラメータ 1 と γ によるベータ分布に従う。次に、ランダム測度 G_0 は、次のように定義される。

$$G_0 = \sum_{f=1}^{\infty} \beta_f \delta_{\phi_f},$$

ここで、 $\Phi = (\phi_f)_{f=1}^\infty$ は、 $(J-1)$ -次元の単体であり、ディリクレ分布 H から生成されるが、これは多項分布の共役事前分布

$$\phi_k | H \sim H \quad \text{for } f = 1, 2, 3, \dots$$

である。ここに、 δ_ϕ は、位置 ϕ におけるアトムであり、 β_f は、 ϕ に集中する確率測度である。ここで、各 G_0 は、基底測度 H と集中度パラメータ γ に基づいてディリクレ過程 $\text{DP}(\gamma, H)$ に従って生成されることに注意しておく。

次に、すべての $n \in N$ 、に対して、ランダム測度 G_n は、次のように定義される。

$$G_n = \sum_{f=1}^{\infty} \pi_{n,f} \delta_{\phi_f},$$

ここで、 $\pi_v = (\pi_{n,f})_{f=1}^\infty$ は、 $f = 1, 2, 3, \dots$ に対して、以下のように生成される。

$$\pi'_{n,f} | \alpha_0, \beta \sim \text{Beta} \left(\alpha_0 \beta_f, \alpha_0 \left(1 - \sum_{\ell=1}^f \beta_\ell \right) \right),$$

$$\pi_{n,f} = \pi'_f \prod_{\ell=1}^{f-1} (1 - \pi'_\ell)$$

ここで、各 G_n は、基底測度 G_0 と集中度パラメータ α_0 を伴うディリクレ過程 $DP(\alpha_0, G_0)$ に従って生成されることに注意しておく。これにより、HDP-LDA は、個々の文書 d'_n に対応した多重トピックを表現する離散的確率分布 G_n , ($n = 1, \dots, N$) を出力する。

2.3.3 JS Divergence に基づくネットワーク構築

G_n に基づいて、種類が未知数のホット撮影スポット大分類手法を考える。本研究では、インタラクティブに可視化および再可視化を行う上で、処理の高速性も求められる。そこで、大規模なデータに対して高速であり、クラスタ数を自動決定できる Newman クラスタリング [Clauset 04] を用いる。このクラスタリング手法は、大規模な複雑ネットワークのグラフに内在するコミュニティ構造を高速に抽出する手法であるため、グラフ表現されたネットワークを構築する必要がある。そこで、ホット撮影スポットをノードとし、ホット撮影スポット内の写真 D_m に基づいて異なるノード間の類似度を求め、その値をリンクの重みとすることで、重み付完全グラフを構築する方法が考えられる。

ここで、多重トピック (嗜好) を反映した各 $d_n \in D_m$ に付随する離散的確率分布 G_n に基づいて類似度を算出する場合、 G_n は単体であるため、一般的に、Kullback-Leibler divergence $D_{KL}(G_p||G_q)$ などの相違度が利用される。しかし、 D_{KL} は対称性を持たないため不便な一面がある。一方、 $D_{KL}(G_p||G_q)$ を用いて、対称性を満たすよう変換できる Jensen-Shannon divergence $D_{JS}(G_p||G_q)$ がある [Nielsen 10] ($\lambda = 1/2$ で対称)。

$$D_{JS}(G_p||G_q) = \lambda D_{KL}(G_q||\lambda G_q + (1-\lambda)G_p) + (1-\lambda) D_{KL}(G_p||\lambda G_q + (1-\lambda)G_p)$$

この D_{JS} を用いることで、本研究では、離散確率分布どうしの類似度 S を以下のように定める。

$$S(G_p||G_q) = 1 - D_{JS}(G_p||G_q)$$

ホット撮影スポットをノードとしたネットワーク $G = (V, E)$ を構築する上で、二つの異なるホット撮影スポット a と b 間の類似度を求めるため、本研究では、ホットスポット写真 D_a の $\{d'_{a,1}, \dots, d'_{a,N_a}\}$ 対 D_b の $\{d'_{b,1}, \dots, d'_{b,N_b}\}$ とのすべてのペアで類似度 S を計算し、その平均値をリンクの重みとした重み付き完全グラフを構築した。ここで、 V はノード (ホット撮影スポット) 集合、 E はノード間のリンク集合を表す。

2.3.4 Newman クラスタリングに基づくスポット分類

Newman クラスタリングでは、グラフ内に潜在的に存在するコミュニティ i 内のリンク密度が高く、コミュ



図 1: 地図上に可視化した Geo-tag 付写真群 $\mathcal{D}_T$.

ニティ間のリンク密度が低い状態を良いクラスタ構造と見なす。クラスタリングの精度を測るモジュール性指標 (modularity) $Q > 0$ を導入し、 Q が 0 になるランダムネットワークに対し、コミュニティ構造の存在を反映した Q を最大にするクラスタ数を自動決定できる。

$$Q = \sum_i (e_{ii} - a_i^2) \quad (2)$$

ここで、 e_{ii} は「コミュニティ i 内のノードと i 内の別のノード間のリンク総和」の「総リンク数」に対する割合を表し、 a_i は「コミュニティ i 内のノードから出ているリンク総和」の「総リンク数」に対する割合である。また、 $\Delta Q_{ij} = 2(e_{ij} - a_i a_j)$ を計算することで、高速化が行われている [Clauset 04]。

本研究では、多重嗜好を背景として、複数のホット撮影スポットを同類としてまとめる、未知数である大分類のカテゴリーを抽出するため、重み付ネットワーク G に Newman クラスタリング [Zhou 10] を適用する。

2.4 没入型 3D メタバース情報空間

2.4.1 情報空間没入型・探検アプローチのコンセプト

本研究の手法により、クラスタリングの結果、ホット撮影スポットに関する多重嗜好を背景とした大分類のカテゴリーが得られる。ここで、外国人が日本への渡航に向けて旅行を計画する事例を考える。Social Media の情報を用いて旅行計画の支援を行うため、本研究では、観光スポットに関連する写真を提示するアプローチを採用するため、日本語のわからない外国人でも、どのような位置にどのような対象があるかを写真を通じて視覚的に探索が可能となる。そこで、大分類を行ったあるカテゴリー内の写真群をユーザに提示する場合、例えば、Google Map 上に Geo-tag 情報を用いて可視化すれば、図 1 のような印象になると考えられる。しかし、この可視化法では支援をしているとは言い難い。

このような問題に対し、これまでは、キーワードを与えることで検索システムを利用したり、嗜好情報を提供することで、推薦システムを利用するなど、自動

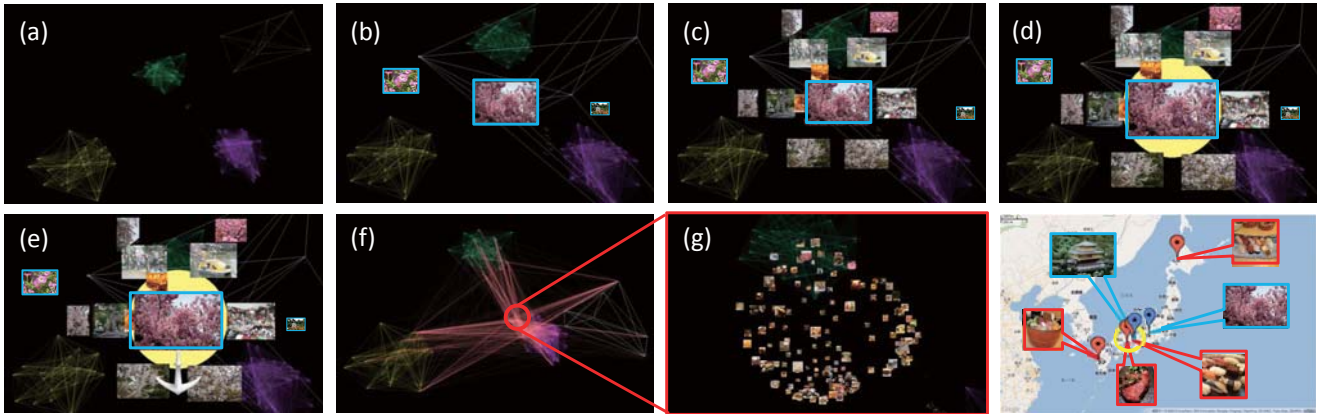


図 2: 3D Metaverse による旅行支援.

フィルタリングを行ってユーザができるだけ楽をする情報提示が行われてきた。しかし、近年、インターネット・ユーザの脳力を衰退させる様々な問題が指摘されており [Carr 10]、他にも、最近の大学生について、例えば「調べる」の意味が「Wikipedia を検索すること」になってしまっており、旧来のように、断片的な情報を方々から集め、情報を整理して自ら考えて理解する行為が失われつつあることを懸念する声さえ聞かれる。

本研究では、Social Media の情報をあらかじめ高度にフィルタリングして提示するのではなく、古来から人が持ち合わせた狩猟・採集能力や視覚的パターン解析能力を刺激しつつ、情報群が可視化された空間にあえて没入し、空間内を探索し、ユーザの嗜好と合う対象を探しに行く行動を支援するアプローチに着目する。

この没入型の探索アプローチは、ユーザがキーワードの入力や嗜好情報などをあらかじめ用意する必要はなく、cold-start でもよく、情報要求が極めて曖昧な状態で可視化された情報空間を探索する中で、かつて冒険や探索により得た幸運な対象、多様な対象と出会う機会を増やし、ユーザの情報探索能力や知識発見能力を刺激するシステムの構築が目指す方向性である。

2.4.2 Social Media に基づく飲食店情報の可視化

ところで、観光スポットを巡る上で、食事も重要な要素である。本研究では訪問先の計画を支援する上で、観光スポットだけでなく、飲食店を加えた異種情報を複合的に扱う問題にも取り組む。飲食店の位置や客が食べた料理の写真情報が投稿される飲食店共有サイトに、食べログ³という Social Media がある。このユーザは、日常のログとして食事内容を写真に撮るだけではなく、それぞれの嗜好に基づいて特別だと思えた料理を撮影する場合も多いと考えられる。また、食べログには、店の評価や食べた料理の評価を登録する機能が

あり、公開されているため、撮影された料理に関する多様な情報をもとに嗜好を分析することができる。

本研究では、食べログの店舗の Geo-tag 情報を付与した写真群（主に料理）を、ホット撮影スポットとの実空間上の距離に応じて近傍に可視化することで、ホット撮影スポットおよび、飲食店情報という複数の異なる情報を同じ空間に可視化する方法についても述べる。

3 提案システム

3D Metaverse 情報空間に没入、探索し、ホット撮影スポットや食べログの写真に出会い、インタラクティブに操作し、情報空間と 2 次元地図空間を往来することで、旅行計画を支援するシステムについて提案する。

3.1 観光旅行計画の支援

図 2 の (a) ~ (h) に示す、実例を用いて説明する。図 2(a) に見える星雲のようなものは、ホット撮影スポットをノードとする、ホット撮影スポットノードの大分類クラスタであり、同クラスタのリンクを同色にしたノード間リンクを 3D 情報空間に可視化した例である。

- (a) 情報空間に没入し、俯瞰より眺め、いずれかの星雲（クラスタ）に向かい、適当に探索する（図 2(a)）。
- (b) 探索中、ホット撮影スポットを象徴する代表写真（図 2(b) 中、青枠）がノードとして見える。
- (c) ホット撮影スポットノードとの距離が近くなると、対応するホット撮影スポット内の写真群が周囲に現れ（図 2(c)）、より詳細な情報が確認できる。
- (d) 好みの一致したノードにさらに近づくともノードの代表写真の背後に黄色の円が現れ、システムが注視状態を認識したことをユーザに伝える（図 2(d)）。
- (e) ユーザが注視対象を訪問先として選択するため、ボタンを押すと、アンカーが現れ、ホット撮影スポットノードが選択状態になったことを示す（図 2(e)）。

このような探索と操作を繰り返すことで、いくつかの訪問先候補にアンカーが着いた状態になる。

³<http://tabelog.com/>

一方、食べログの写真ノードも同時に可視化できるが、食べログノードは、各ホット撮影スポットの実空間上で、1500m 内に存在する飲食店とだけリンクが張られる。ただし、食べログノードは、実空間で複数のホット撮影スポット周辺が重なる位置にある場合、両方のホット撮影スポットノードとリンクが張られる。

図 2(f) は、その食べログノードとホット撮影スポットノードをつなぐリンクを可視化した例である。図 2(f) 中、赤丸の位置は、四つの星雲とリンクでつながっている食べログノード群の存在が暗示される。その部分に近づいた例が図 2(g) である。四つの星雲とリンクでつながった食べログノード集団は、異なる四つの嗜好大分類クラスタに含まれるホット撮影スポットノードとそれぞれ実空間上で近いことが暗示されていると見なせるため、多様で近い複数のホット撮影スポットを観光しつつ、好みの食事にありつける可能性がある。

また、食べログノードも図 2(b)~(e) と同様、アンカーを置くことができる。複数のホット撮影スポットや飲食店にアンカーを置いた状態で、地図空間に遷移した状態が図 2(h) である。例えば、地図空間上では、黄色の円の地域にアンカーを置いたものうち多くが集まっているとユーザが見なせば、訪問先は、この黄色の円内の地域に行くという計画を立てることができる。また、地図空間上の位置関係が悪く、気に入らなければ、また 3D 情報空間に戻り、別の対象を探しに行くことになる。このような複数の空間を往来し、対話的操作を繰り返すことで、提案システムは観光計画を支援する。

3.2 インタラクション機能

ユーザは、Bluetooth に基づく、図 3(a)(b) のようなボタン操作付きハンドル型の無線デバイス (Wii リモコン) を用いて、地面のない宇宙的空間内を自由に移動でき、ボタンを用いて前進、後退が可能となっている。ユーザが操作する様子を示したものが図 3(c) である。

また、ボタン操作により、アンカー設定、全リンク表示・非表示、ホット撮影スポット間のみのリンク表示・非表示、食べログとホット撮影スポット間のリンク表示・非表示、クラスタ内リンクの表示・非表示、食べログノードの表示・非表示などを操作可能とした。

また、大分類クラスタ間のリンクを非表示にすると、ばねモデルによるノード配置処理が各クラスタ内で独立に進行させるよう実装し、ボタン操作で、大分類のクラスタ間を大きく引き離す機能も加えた (図 2(a))。

3.3 クラスタリングにおける類似度の比較法

HDP-LDA の効果を評価するため、各写真 d_n ごとの Visual Words ヒストグラムを得たのち、異なる二つの

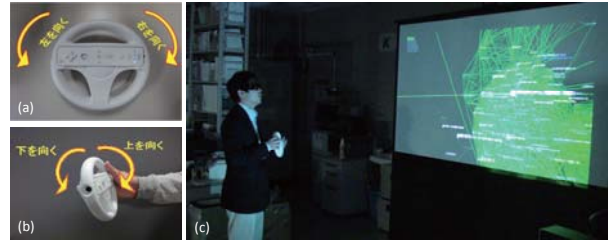


図 3: 無線ハンドル兼ボタン操作インターフェース。

ホット撮影スポットに含まれるホットスポット写真 D_m どちらの写真に総当たりによる $\cos$ 類似度を算出して平均値を算出した。ホット撮影スポットをノードとし、 $\cos$ 類似度の平均値をリンク重みとする完全グラフを構築し、Newman クラスタリングを適用した。

4 Social Media Data を用いた実験

4.1 データセット

本研究では、写真共有サイト Flickr から、2010 年 1 月 1 日から 2010 年 12 月 31 日日本列島が含まれる矩形領域 (緯度:25.8 ~ 45.8, 経度:126.2 ~ 146.8) 内の Geo-tag および撮影時間情報など、EXIF 情報付写真を用いて抽出された、ホット撮影スポット 103 地域に含まれるホットスポット写真 $N=620$ を用いる。また、食べログから、ホット撮影スポット 103 地域の近傍にある 3000 の飲食店を抽出し、4988 枚の Geo-tag 付写真を抽出した。

4.2 実験結果

Visual Words を抽出するため $K=500$ として実験を行った。HDP-LDA を用いて G_n を算出し、類似度 S を用いて完全グラフを構築した。本研究では、この G を対象として、Newman クラスタリングを適用した。

比較法の $\cos$ 類似度に基づく Newman クラスタリングでは、2 つにしかクラスタが分割されなかったが、HDP-LDA を用いた方法では、4 つに分割され、より多くの大分類クラスタが得られた。このネットワーク G を、古典的なばねモデル (KK 法 [Kamada 89]) を用いて情報空間に可視化・レイアウトした。3D 情報空間を探索する実験も行ったが、大量の写真ノードとリンクが表示された環境下で、スムーズに空間内を探索でき、実装した機能はリアルタイムかつ良好に動作した。

5 まとめ

本研究では、Social Media を通じて得た写真というデータ群に潜在する嗜好を反映して、情報空間と地図

空間を往来しながら旅行計画を支援する 3D Metaverse 情報空間の没入型システムを試作した。システム動作は良好であったものの、観光旅行支援の有効性の評価が十分ではないため、今後の課題とする。

ところで、提案法は、検索システムや推薦システムに対峙する手法ではなく、既に指摘した弱点を補い、相互の連携を通じて、統合された Metaverse に基づくシステムを構築することが今後を通じての指針となっている。今後の研究として、ユーザの脳力を衰退させるのではなく、刺激しつつ、提案システムのさらなる発展のため、Social Media のより多くの情報を活用する方法、探検の支援を行う方法、Serendipity の機会をより増やす方法、情報空間への情報の埋め込み法、情報空間の構造化・階層化法の検討や、インタラクションのより効果的な手法を探索する予定である。

参考文献

- [Arase 10] Arase, Y., Xie, X., Hara, T., and Nishio, S.: Mining People's Trips from Large Scale Geo-tagged Photos, in *Proceedings of the 18th International Conference on Multimedia*, pp. 133–142 (2010)
- [Bay 08] Bay, H., Ess, A., Tuytelaars, T., and Van Gool, L.: Speeded-Up Robust Features (SURF), *Comput. Vis. Image Underst.*, Vol. 110, No. 3, pp. 346–359 (2008)
- [Carr 10] Carr, N.: *The Shallows, The Shallows: What the Internet Is Doing to Our Brains*, W. W. Norton & Company (2010)
- [Clauset 04] Clauset, A., Newman, M. E. J., and Moore, C.: Finding community structure in very large networks, *Physical Review E*, pp. 1–6 (2004)
- [Comaniciu 02] Comaniciu, D. and Meer, P.: Mean shift: a robust approach toward feature space analysis, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 24, No. 5, pp. 603–619 (2002)
- [Crandall 09] Crandall, D. J., Backstrom, L., Huttenlocher, D., and Kleinberg, J.: Mapping the world's photos, in *Proceedings of the 18th International Conference on World Wide Web*, pp. 761–770 (2009)
- [Csurka 04] Csurka, G., Dance, C. R., Fan, L., Willamowski, J., and Bray, C.: Visual categorization with bags of keypoints, in *In Workshop on Statistical Learning in Computer Vision, ECCV*, pp. 1–22 (2004)
- [Evans 08] Evans, J. A.: Electronic Publication and the Narrowing of Science and Scholarship, *Science*, Vol. 321, No. 5887, pp. 395–399 (2008)
- [Hartigan 79] Hartigan, J. and Wong, M.: Algorithm AS 136: A K-means clustering algorithm, *Applied Statistics*, pp. 100–108 (1979)
- [Kamada 89] Kamada, T. and Kawai, S.: An algorithm for drawing general undirected graphs, *Inf. Process. Lett.*, Vol. 31, No. 1, pp. 7–15 (1989)
- [Kisilevich 10a] Kisilevich, S., Keim, D. A., and Rokach, L.: GEO-SPADE: A generic Google Earth-based framework for analyzing and exploring spatio-temporal data, in *12th International Conference on Enterprise Information Systems*, Vol. 5, pp. 13–20 (2010)
- [Kisilevich 10b] Kisilevich, S., Mansmann, F., Bak, P., Keim, D. A., and Tchaikin, A.: Where Would You Go on Your Next Vacation? - A Framework for Visual Exploration of Attractive Places, in *GeoProcessing 2010*, pp. 21–26 (2010)
- [Kisilevich 10c] Kisilevich, S., Mansmann, F., and Keim, D. A.: P-DBSCAN: A density based clustering algorithm for exploration and analysis of attractive areas using collections of geo-tagged photos, in *1st International Conference on Computing for Geospatial Research & Application* (2010)
- [Lu 10] Lu, X., Wang, C., Yang, J.-M., Pang, Y., and Zhang, L.: Photo2Trip: generating travel routes from geo-tagged photos for trip planning, in *Proceedings of the 18th International Conference on Multimedia*, pp. 143–152 (2010)
- [Nagel 08] Nagel, H. R., Granum, E., Bovbjerg, S., and Vittrup, M.: Visual Data Mining, in Simoff, S. J., Böhlen, M. H., and Mazeika, A. eds., *Visual Data Mining*, chapter Immersive Visual Data Mining: The 3DVDM Approach, pp. 281–311, Springer-Verlag, Berlin, Heidelberg (2008)
- [Nielsen 10] Nielsen, F.: A family of statistical symmetric divergences based on Jensen's inequality, *CoRR*, Vol. abs/1009.4004, (2010)
- [Smart 07] Smart, J., Cascio, J., and Paffendorf, J.: Metaverse Roadmap Overview, <http://www.metaverse-roadmap.org/overview/> (last accessed 12.10.06) (2007), Accelerated Studies Foundation, Retrieved 2010-09-23
- [Surowiecki 04] Surowiecki, J.: *The Wisdom of Crowds: Why the Many Are Smarter Than the Few and How Collective Wisdom Shapes Business, Economies, Societies and Nations*, Doubleday (2004)
- [Teh 06] Teh, Y. W., Jordan, M. I., Beal, M. J., and Blei, D. M.: Hierarchical Dirichlet Processes, *Journal of the American Statistical Association*, Vol. 101, pp. 1566–1581 (2006)
- [Yin 10] Yin, H., Lu, X., Wang, C., Yu, N., and Zhang, L.: Photo2Trip: an interactive trip planning system based on geo-tagged photos, in *Proceedings of the 18th International Conference on Multimedia*, pp. 1579–1582 (2010)
- [Yin 11] Yin, Z., Cao, L., Han, J., Zhai, C., and Huang, T.: Geographical Topic Discovery and Comparison, in *Proceedings of the 20th International Conference on World Wide Web*, pp. 247–256 (2011)
- [Zhou 10] Zhou, T. C., Ma, H., Lyu, M. R., and King, I.: User-Rec: A User Recommendation Framework in Social Tagging Systems, in Fox, M. and Poole, D. eds., *AAAI*, AAAI Press (2010)
- [王 11] 王 佳な, 野田 雅文, 高橋 友和, 出口 大輔, 井手 一郎, 村瀬 洋: Web 上の大量の写真に対する画像分類による観光マップの作成, 情報処理学会論文誌, Vol. 52, No. 12, pp. 3588–3592 (2011)
- [熊野 12] 熊野 雅仁, 小関 基徳, 小野 景子, 木村 昌弘: 地理および時間情報をもつ写真データに基づいたホット撮影スポットの抽出, 情報処理学会論文誌, Vol. 5, No. 3, pp. 41–53 (2012)
- [三末 09] 三末 和男: ネットワーク可視化技術 - 大規模ネットワークと動的ネットワークへの挑戦, 電子情報通信学会論文誌, Vol. 92, No. 2, pp. 112–117 (2009)
- [神鷹 07] 神鷹 敏弘: 推薦システムのアルゴリズム (1), 人工知能学会誌, Vol. 22, No. 6, pp. 826–837 (2007)
- [神鷹 08a] 神鷹 敏弘: 推薦システムのアルゴリズム (2), 人工知能学会誌, Vol. 23, No. 1, pp. 89–103 (2008)
- [神鷹 08b] 神鷹 敏弘: 推薦システムのアルゴリズム (3), 人工知能学会誌, Vol. 23, No. 2, pp. 248–263 (2008)

テキスト分析における試行錯誤の支援に向けて — TETDM のインタフェースに関する一考察 —

Towards Supporting Trial and Error Process on Text Analysis — Redesign of the Interface for TETDM —

大塚 直也^{1*} 松下 光範¹
Naoya Otsuka¹ Mitsunori Matsushita¹

¹ 関西大学 総合情報学部

¹ Faculty of Informatics, Kansai University

Abstract: 本研究の目的は探索的にテキスト分析を行うユーザの支援である。ユーザの行うテキスト分析は一般に one-shot のプロセスではなく、試行錯誤を繰り返しつつ有益な情報を発見するプロセスである。現在、様々なテキストマイニングモジュールを切り替えながらこのような分析を行う基盤として TETDM が提案されているが、現在の実装では円滑にモジュール切り替えを行えるインタフェースになっておらず、試行錯誤を伴う分析を行う場面で利用するには十分ではない。そこで我々は、TETDM のインタフェースが満たすべき要件について考察を行い、改良インタフェースのプロトタイプを実装したので報告する。

1 はじめに

現在、膨大なテキストデータを日常的に利用できる環境が整ってきている。それに伴い、それらのテキストデータを閲覧だけに利用するのではなく、分析して有益な情報を発見したり新しい知見を見つけたりするための資源として利用することにも期待が高まっている。テキストマイニングは、このような期待の下で研究が進められている技術であり [1]、構造化されていないテキストデータから新たな知識を取り出すことを目的としている。テキストマイニングでは自然言語処理技術やデータマイニング技術、情報可視化技術などの多様な技術が適応的に利用されており、その結果が重要文の抽出や文書集合のクラスタリング、文書要約などに役立てられている。

一般にテキスト集合の中から新しい知識を発見する場合、人は様々な観点からテキストを処理し、その結果を解釈しながら探索を進めていく [2]。すなわちテキストマイニングによる知識発見は、ゴール志向の one-shot の探索プロセスで達成されるタスクではなく、試行錯誤を繰り返しながら有益な情報を発見する探索的な情報アクセスを必要とするタスクと捉えることができる。

しかし現状では、大量のテキストデータの中から有益な知識を得たいと考えている人が持つ様々な要求に、

汎用的かつ十分に対応できる情報アクセス手段は乏しい [3]。また、その分析結果を解釈する際にも利用できる視覚化手法は様々な提案されているが、日常的に使えるレベルのシステムは未だ十分とは言えない [4]。

このような問題を解決するため、汎用的なテキストマイニングの実行を目指した統合環境 (Total Environment for Text Data Mining; 以下 TETDM と略す) が提唱されている [5]。TETDM では、世の中に分散しているテキストマイニングツールを画一的に扱える統合環境を構築することによって、テキストデータを扱うユーザの創造的活動を支援するツールの提供を目指している。

現在 TETDM の統合環境は開発途中であり、そのインタフェースはテキストマイニングのような試行錯誤を必要とする情報探索に適したものになっているとは言い難い。そこで本研究では、TETDM の目的に沿って、ユーザの試行錯誤を伴う創造的活動をより円滑に支援するための環境に改善することを目指している。このような観点の下、本稿では TETDM のインタフェースが満たすべき要件についての考察を行い、改良インタフェースのプロトタイプを実装したので報告する。

2 TETDM の概要

本節では、これまでに進められている TETDM の概要と、そこで利用可能なマイニングツールの例を示す。

*連絡先： 関西大学 総合情報学部
〒 569-1095 高槻市霊仙寺町 2-1-1
E-mail: k589149@kansai-u.ac.jp

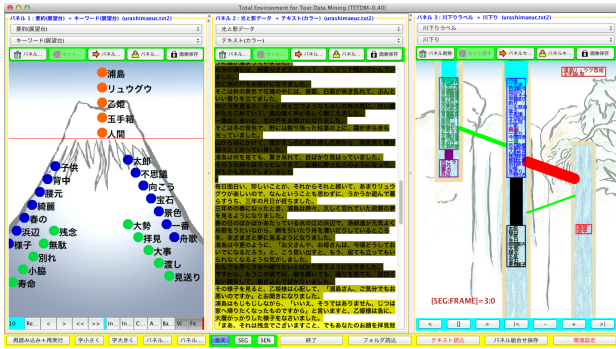


図 1: 複数の可視化結果を表示した時の TETDM の例

2.1 TETDM

これまで、テキストマイニング及びその要素技術について、多くの研究成果が発表されている [3]。しかし、これらの研究により提案されたツールは、論文用の試験的なものが多く実際に世の中で使われている技術はごく一部に限られている。

これらのツールは各研究者によって独自に構築されていることが多いため、個別に導入しなくてはならず、異なるツール間のデータ受け渡しを行うために手作業でデータフォーマットを整えたり、複数のツールの処理結果を比較するために新たに比較用のインタフェースを実装したりするといった手間が必要であった。そのため、テキストデータを用いて情報を多角的に分析したい場合や従来のマイニングツールと新たに開発したマイニングツールの特徴を比較したい場合に、複数のツールを用いることが容易ではなかった。

このような問題を解決するために、TETDM プロジェクトでは世の中に分散しているテキストデータマイニングツールを画一的に扱える環境の構築を目指し、それに向けた統合環境 (図 1 参照) の構築を進めている¹。

前述したように、TETDM の目的は同一環境上で複数のテキストマイニングツールを扱えるようにすることである。そのため、同一のインタフェース上で複数のツールを組み合わせたり、横に並べて表示することができる仕様になっている。利用者は、自らの興味や目的に応じてツールを選択しそれらを自由に組み合わせたり、処理結果を比較することによって、多角的なテキスト分析を行うことができる。また、この統合環境上では複数のマイニングツール間のインタラクションを想定しており、あるツールの出力結果を他のツールの入力として用いたり、出力結果に対しての操作が並列に表示されている他のツールの処理結果にも即座に反映されたりするようになっている。

TETDM に入力されたテキストファイルは、共通に行

¹TETDM の最新版は TETDM プロジェクトのホームページ (<http://tetdm.jp/>) 上で公開されている (平成 24 年 10 月 26 日現在の最新版は 0.40 である)。

われる前処理 (例えば形態素解析など) の後、各モジュールによって処理されるようになっている。モジュールとは、テキストマイニングを行う機能ごとにまとめられた要素で、これを TETDM に追加していくことで、取り扱えるテキストマイニング手法を拡張できるようになっている。テキストマイニングの要諦である二つの技術、すなわちマイニング技術と可視化技術は統合環境内のモジュールとして実装されており、それらをユーザが切り替えることで処理内容と処理結果の表示手法を変更できる。

現状の TETDM のインタフェースは、横に並べられた複数のツールパネルから構成されている。その一例を図 1 に示す。この図では、3 つの異なるツールが同時に表示されている。それぞれのツールパネルの上部には 2 つのドロップダウンリストがあり、それを通じてマイニングモジュールと可視化モジュールを選択することができる。

利用可能なマイニングモジュールと可視化モジュールの一覧を表 1、2 に示す。現在、既にいくつかのテキストマイニングツールが TETDM にモジュールとして実装されており、利用可能になっている。

2.2 利用可能なマイニングツールの例

TETDM で利用可能な分析手法とモジュールについて例をいくつか示す。

テキストの主題関連度を可視化するシステムとして「ひなたシステム」がある [6]。このシステムは文書中のすべての名詞と各文の関連度を計算し、各文の背景色を関連度が低いものほど黒く、高いものほど黄色く彩色する。主題と関連する箇所や全体に対する位置や割合を確認できるようにすることで、ユーザの文章理解の支援を試みている。TETDM では「光と影データ」(表 1-1 参照) というマイニングモジュールと「分布」(表 2-1 参照)、「テキスト (カラー)」(表 2-2 参照) の 2 つの可視化モジュールを組み合わせることで、このひなたシステムを実現できる。

また、テキストの主題に基づく一貫性評価を行えるシステムとして「川下りシステム」が考案されている [7]。川下りシステムでは、テキストとその主題を表す単語の集合を入力とし、テキスト中で使われている各単語に主題との関連を表すラベルを与え、それに基づいた単語分布を可視化することによりテキストの主題に基づく一貫性評価を行える。TETDM では、マイニングモジュール「川下りラベル」(表 1-2 参照) と可視化モジュール「川下り」(表 2-3 参照) によって実装されている。また、入力に用いられる主題を表す単語の集合は、統合環境上に実装された展望台システム [8] の出力結果を用いている。

表 1: 主な利用可能マイニングモジュール

モジュール名	内容
1 光と影データ	文章中の各文の主題関連度を評価
2 川下りラベル	テキストの一貫性評価
3 要約 (展望台)	特徴的な単語の抽出、重要分抽出
4 主語抽出	文書内の主語を抽出
5 単語チェック	文章作成時のチェック支援
6 長文チェック	文書中の長文の提示
7 テキスト分析	他モジュールの出力結果の集約
8 Twitter	Twitter の検索結果を入力テキストとする
9 単語間関連度	各単語の単語間関連度を出力
10 タグデータ	単語の出現頻度によるタグクラウド用データを出力
11 アノテーション	アノテーションを付与すべき文の抽出

表 2: 主な利用可能可視化モジュール

モジュール名	内容
1 分布	テキスト中の各文の評価値を横向き棒グラフで表示
2 テキスト (カラー)	テキストの各文や単語に背景色をつけて表示
3 川下り	川下りラベル用
4 キーワード (展望台)	キーワードを富士山画像の上に配置
5 HTML テキスト	html 形式でテキストを表示
6 キーワードマップ	キーワードを 2 次元平面上に配置
7 タグクラウド	単語の出現頻度によるタグクラウド

このように TETDM では、同一のマイニング処理モジュールからの出力に対し、ユーザの目的に応じて異なった可視化手法を適用できる。既存ないし将来の研究成果によるテキストマイニングツールを同一環境上に集め、画一的に扱える環境を構築することにより、複数ツールを用いた分析を行う際に余計な作業をする必要がなくなり、分析作業に集中できることが期待される。また、多くの研究成果の実用化の際や複数のマイニング手法の比較の際に利用が見込まれる。

3 既存のインタフェースの問題点

前述したように、テキスト分析は試行錯誤を伴う探索的な情報アクセスであるため、ユーザは様々な分析手法の中から適当なひとつの手法を選択して分析を行い、得られた結果を解釈・考慮して次の分析手法を試みる、といった探索プロセスを繰り返す。そのため、TETDM のインタフェースは、こうした探索プロセスを円滑に行えるものでなければならない。しかし、現在の TETDM のインタフェースは、そのようなプロセ

スを円滑に行うには十分ではないと考えている。以下にその問題を 3 点に分けて説明する。

1. モジュール選択の問題

現在の TETDM では、分析に用いるモジュールを選択するためのインタフェースとして、ドロップダウンリストが用いられている (図2 参照)。これを用いてモジュールの選択や切り替えを行う場合、ユーザが利用可能なモジュールを把握するためには、格納されているリストを展開する必要がある。また、現在の実装では単純なリストが採用されているため、利用可能なモジュール全体を概観することができない。これは、どのような分析が可能かを熟知していないユーザにとっては、作業の円滑な遂行の妨げになる。

2. モジュールの入出力データの問題

TETDM では、複数のテキストマイニング技術をシステム内のモジュールとして扱うことにより、それらを組み合わせて多角的なテキスト分析を行える。しかし、現在利用可能なツールには、

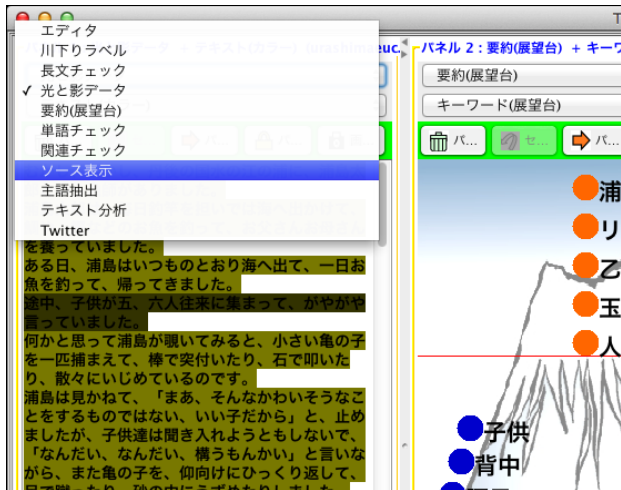


図 2: ドロップダウンリストによるモジュール切り替え

マイニングモジュールと可視化モジュールが一对一の対応となっているものや、一つのモジュールで複数のマイニング処理を行なっているものがある(表 1 参照)。それらのモジュールは、特定のモジュールと組み合わせることのみを前提としており、汎用性が低く、他のモジュールと柔軟に組み合わせることができない。

3. データフローの問題

TETDM に入力されたテキストデータは、形態素解析などの前処理の後、マイニングモジュール及びそれと組み合わせられている可視化モジュールによって処理される。それに加えて、他のモジュールの出力を入力したり、あるモジュールに対する操作が他のモジュールに反映される仕組みが用意されているため、TETDM を用いたテキスト分析では、入力されたテキストデータに対する処理の流れが煩雑となる。しかし、現在の TETDM では、組み合わせられて用いられているモジュール以外の処理の流れが明示されておらず、ユーザは自身の分析行為が他の分析結果に及ぼす影響を把握することが困難である。そのため、ユーザの意図しない操作や処理結果が暗黙的に引き起こされる可能性があり、分析の妨げになる。

これらの問題を解決し、ユーザの探索行為がより円滑に支援できるインタフェースのデザイン指針について次章で述べる。

4 デザイン指針

本研究では、3 章で指摘した問題を解決するために、TETDM のインタフェースが満たすべき要件をユーザ

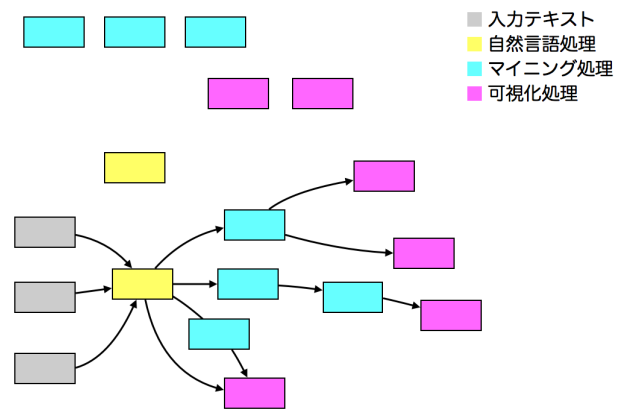


図 3: 提案するインタフェースの概念図

が (1) モジュール切り替えを円滑に行えること、(2) 各モジュールの利用状況を把握できること、(3) モジュール間の処理の流れを把握できること、の 3 つと考えた。

本節では、このようなインタフェースを設計する際に我々が採用したデザイン指針について述べる。

4.1 ユーザの利用に関する指針

TETDM では複数のツールを用いて多面的に分析を行うことを想定しているため、分析内容が多岐にわたり、そのための作業が煩雑になる可能性がある。そのため、ユーザは自身の分析過程を常に把握し、現在行っている分析に引き続いてどのような分析ができるのかを把握できなくてはならない。

このような要求に応えるため、本研究では各モジュールをノード、それらを処理の順に繋ぐ線をリンク、とするグラフ表現をインタフェースに採用することとした。図 3 にその概念図を示す。提案手法では、各モジュールを表すノードをマウスで直接操作することによって、ツールの選択・切り替えを行えるようになっている。このようにオブジェクトを直接操作することにより、直感的なモジュール切り替えを実現でき、3 章で述べたようなモジュール選択の問題を解消できると考えている。

4.2 データの制約

TETDM ではテキストマイニングのプロセスを、自然言語処理、データマイニング、情報可視化の 3 つのプロセスとして捉えている。各プロセスは必ずしも 1 回の処理のみで終わるのではなく、異なる自然言語処理、データマイニングを複数回行うことも想定される。これは、テキストマイニングにおけるデータの流れが複数の処理から構成されることを意味する。提案する

インタフェースではグラフ表現を採用するため、この一連のプロセスをテキストデータの入力から可視化結果までの一本のパスとして捉えることができるようになる。

各モジュールには入力と出力があり、それによっての組み合わせ可能なモジュールの対が決定される。例えば表 1-2 の「川下りラベル」モジュールは plainText を入力とし、ラベルが付与されたテキストを出力する。そして、この出力結果は表 2-3 の「川下り」モジュールでのみ利用される。

現在は、モジュールの粒度がまちまちで、単一のモジュールで複数の処理を行ったり、限定されたモジュールのみとデータのやりとりを行ったりする場合もある。そこで、各モジュールは単一の処理のみを担うものとして整理し、それらの入出力のデータ形式を標準化することで、柔軟にモジュールを選択し、組み合わせ利用できるように粒度を整えることにした。これにより、モジュール間の連携の統一化が図られ、より柔軟にモジュールを組み合わせることができるようになるため、3 章で指摘したモジュール間の入出力データの問題やデータフローの問題が改善されるようになると考えている。

5 実装

本節では、実装を行った改良インタフェースのプロトタイプについて述べる。なお、プロトタイプシステムは TETDM version 0.32 をベースとして実装した。図 4 に、現在実装したプロトタイプのスナップショットを示す。この実装プロトタイプは、モジュール切り替えのための改良インタフェース (図 4、画面左) と、既存システムによるツールパネル (図 4、画面右) から構成されている。

改良インタフェースでは、モジュールの利用状況をグラフとして表現し、各モジュールを表すノードを直接操作することによってモジュール切り替えを行える。

既存の TETDM のインタフェースでは、複数のマイニングツールを横に並べて表示できたが、今回の実装では、同時に複数のマイニングツール並べて表示することはできない。

改良インタフェース上に表示されている矩形はその時点で利用可能なモジュールを表しており、それぞれがモジュールの利用状況を表すグラフのノードとなっている。各ノードには、対応付けられているモジュールの名称を付与した。TETDM で定められているモジュールは、マイニングモジュールと可視化モジュールの 2 種類であるが、改良インタフェースでは、TETDM によって入力テキストに対し共通に行われる形態素解析などの前処理を表す仮想モジュールを追加し、合計 3

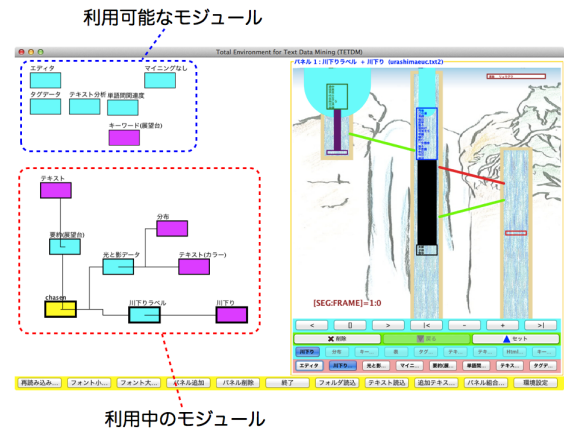


図 4: 実装プロトタイプ

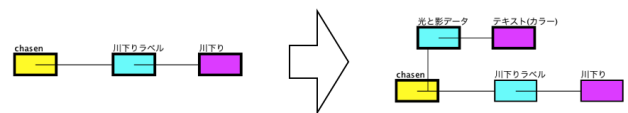


図 5: 「川下り」(左図) に「光と影」を加えた時の様子

種類とした。TETDM によって行われる形態素解析には ChaSen[9] が用いられているため、仮想モジュールには名称として「chasen」を付与している。線で結ばれているノード同士は、それらが示すモジュールが組み合わせて利用されていることを表している。提案インタフェース上に多数のモジュールが表示されている場合、モジュールの種類が判別し難くなる。そのため、モジュールが行う処理の段階によってそれぞれ異なる色に彩色した。

各モジュールを表すノードは、マウスでドラッグ & ドロップすることにより操作可能である。ドラッグしているノードをそれと組み合わせ可能なモジュールを表すノードに近づけることで線で結ばれ選択状態になる。選択状態のノードの輪郭線は強調されるため、現在用いているモジュールと処理の流れを把握することができる。図 5 の左図では、マイニングモジュールとして「川下りラベル」、可視化モジュールとして「川下り」が組み合わせて利用されているため、それぞれのモジュールを表すノードが線で結ばれている。図 5 の右図では、それに加えマイニングモジュールとして「光と影データ」、可視化モジュールとして「テキスト(カラー)」が用いられている。図 5 のように、複数のマイニングモジュールと可視化モジュールのペアが用いられている場合、可視化モジュールを表すノードをクリックすることで、選択ツールを変更することができる。

6 議論

本研究で実装したプロトタイプでは、複数の処理結果を同時に表示させることができない。提案手法によって、複数ツールの同時表示を行うためには、可視化モジュールを表すノードと処理結果を表示する画面の対応付けを行わなければならない。その実現手法について今後検討する必要がある。

探索的なデータ分析を計算機システムで効果的に支援するためには、探索行為だけではなく振り返り行為についても考慮する必要がある[10]。提案手法では、モジュールの利用状況をグラフとして表現している。そのため、提案手法を用いて、グラフを履歴として保存したり、グラフに対して、ユーザが分析過程で得た知見をアノテーションとして付与できる仕組みを用意することで、探索過程におけるユーザの振り返り行為を支援できるのではないかと考えている。

7 おわりに

本稿では、テキスト分析における試行錯誤の支援を目的とし、TETDM のインタフェースについて考察を行い、それに基づいた改良インタフェースのプロトタイプを既存の TETDM 上に実装した。テキスト分析は試行錯誤を伴うため、TETDM のインタフェースは試行錯誤を円滑に行えるものでなければならないと考えている。このような観点の下で実装した改良インタフェースでは、モジュールの利用状況をグラフで表現し、各モジュールを表したノードを直接操作することによってモジュールの切り替えができるようにした。

今後、TETDM のモジュールの入出力の形式を検討し、より汎用的かつ柔軟に利用できる形式に改善するとともに、被験者実験を通じて、その有効性を検証していく。

謝辞

本研究の遂行にあたり、文部科学省科学研究費(課題番号: 22300048)の助成を受けた。記して謝意を表す。

参考文献

- [1] 渡辺勇: ビジュアルテキストマイニング, 人工知能学会誌, Vol. 16, No. 2, pp. 226–232 (2001).
- [2] White, R. W. and Roth, R. A.: *Exploratory Search: Beyond the Query-Response Paradigm*, Morgan & Claypool Publishers (2009).
- [3] 市村由美, 長谷川隆明, 渡部勇, 佐藤光弘: テキストマイニング: 事例紹介, 人工知能学会誌, Vol. 16, No. 2, pp. 192–200 (2001).
- [4] 増井俊之: 情報視覚化の最近の研究動向, 情報の科学と技術, Vol. 54, No. 11, pp. 558–567 (2004).
- [5] 砂山渡, 高間康史, Bollegala, D., 西原陽子, 徳永秀和, 串間宗夫, 松下光範: Total Environment for Text Data Mining: テキストデータマイニングのための統合環境, 人工知能学会論文誌, Vol. 26, No. 4, pp. 483–493 (2011).
- [6] 西原陽子, 佐藤圭太, 砂山渡: 光と影を用いたテキストのテーマ関連度の可視化, 人工知能学会論文誌, Vol. 24, No. 6, pp. 479–487 (2009).
- [7] 砂山渡, 谷川信弘: テキストの主題に基づく一貫性評価と結論抽出への応用, 知能と情報, Vol. 23, No. 5, pp. 727–738 (2011).
- [8] 砂山渡, 谷内田正彦: 文章の特徴を表すキーワードを発見して重要文を抽出する展望台システム, 電子情報通信学会論文誌, D-I, Vol. 84, No. 2, pp. 146–154 (2001).
- [9] 松本裕治: 形態素解析システム「茶筌」, 情報処理, Vol. 41, No. 11, pp. 1208–1214 (2000).
- [10] 松下光範: InTREND: ユーザの探索行為と振り返り行為に着目したデータ分析支援システム, 情報処理学会論文誌, Vol. 49, No. 7, pp. 2456–2467 (2008).

段落間関係の評価に基づく文章構造推敲支援

Text Structure Polish Support based on Paragraph Relationship Evaluation

山手 砂都美* 砂山 渡
Satomi Yamate Wataru Sunayama

広島市立大学大学院情報科学研究科
Graduate School of Information Sciences, Hiroshima City University

Abstract: We have many opportunities to write a text. However, relationship among paragraphs are hard to be grasped. Therefore, we focused on top-down structures and bottom-up structures between paragraphs. Top-down means texts that a paragraph including a conclusion or what the writer wants to say comes at first and bottom-up means texts that such a paragraph comes lastly. In this paper, a system that supports polish of text structure is proposed. Relationships between paragraphs are expressed as a tree structure and writers can confirm whether a text is top-down or bottom-up. Users of the system can polish their texts to be top-down or bottom-up by seeing output of the system.

1 はじめに

文章を書く機会はさまざまなところである。文章を書いた本人が分かりやすく書いたつもりでも、他人が読むと分かりにくかったり、間違った解釈をしてしまったりすることもある。自分の伝えたい意図を正しく伝えるためには、現状の文章の構造を把握するのはもちろんのこと、文章の構造を推敲する必要がある。しかし、自分で書いた文章の構造を正確に理解し、推敲するのは難しく、自分で推敲するには限界がある。他人に文章構造の推敲を依頼せず、自分で簡単に文章構造を推敲していきたい。

そこで本研究では、文章を段落間のつながりに着目した段落間関係の評価に基づく文章構造推敲支援システムの構築を目的とする。

2 関連研究

2.1 段落間関係の評価に関する研究

本節では、段落間関係の評価に関する研究について述べる。

単語の概念関係を用いて段落の一貫性を解析する研究 [1][2] がある。これらの研究は、あらかじめ一貫した形で構成されていると考えられる技術文章を対象とし、

単語間の意味類似度を用いて提案する段落一貫度が有効であるかどうかを検証したものである。本研究では単語間の意味には着目せず、条件付き確率で段落間の関係性を評価していく。

また、文章のセグメント間関係解析に基づく文章構造解析をする研究 [3] がある。この研究では、小さな意味段落内を修辞構造で扱いつつ、意味段落間の関係づけを行っている。本研究では、意味段落間の関係づけを行っていない。また、前者の研究同様に、条件付き確率で段落間の関係性を評価していく。

2.2 文章の構造化に関する研究

本節では、文章の構造化に関する研究について述べる。

文章構造解析に基づく小論文の理論性についての自動採点を行う研究 [4] がある。この研究では、文章の構造化を文に関して行っている。本研究では、文章の構造化を段落間に関して行っている点、理論性について着目している点と着目していない点が異なっている。

また、情報理論による、文章の理論的構造の解析をする研究 [5] がある。この研究では、文章の理論構造を分析して可視化を行っている。本研究では、理論構造の分析を行わず、条件付き確率の大きい段落間をつなぐことで構造化を行っている。

また、理工系学生を対象とした技術文書作成支援システムを作成した研究 [6] がある。この研究では、あらかじめ、論文のルールに沿って書かれていない文、意

*連絡先：広島市立大学大学院情報科学研究科システム工学専攻
〒731-3194 広島県広島市安佐南区大塚東三丁目4番1号
E-mail: {yamate,sunayama}@sys.info.hiroshima-cu.ac.jp

図が分かりにくい文を検出する機能があり、分かりにくい文をクラス図で可視化している。文章の構造を可視化している点は同じだが、本研究では、段落間を対象とした構造を可視化している。

更に、以上の関連研究では、話の分岐として、話を広げている段落、話をまとめている段落に着目していない点で異なる。

2.3 文章の推敲に関する研究

本節では、文章の推敲に関する研究について述べる。

文の理論構造デザインツールの試作と校正・推敲支援ツールを作成した研究 [7][8] がある。これらの研究では、推敲は文章の理解向上を指しており、句読点の統一チェック、長文チェック等、細かい点に着目している。本研究では、推敲に関して細かい点に着目せず、文章の構造に着目し、文章の枠組みに着目している点異なる。

また、文を分かりにくくする要因の分析と改善支援手法の提案をした研究 [9] がある。この研究では、ユーザである学生が陥りやすい誤りを分析し、その結果より、修正すべき箇所の優先順位と改善方法を指摘している。本研究では、ユーザの誤りから改善方法を指摘するのではなく、現状の文章構造から、より良い文章構造を提案して行く。

3 本研究で扱う文章構造

本章では、本研究で扱う文章構造について述べる。

本研究では、トップダウン構造とボトムアップ構造に着目する。トップダウン構造とは、結論や話の全体の構造に関する説明を先の段落で述べた後、その詳細を後の段落で述べている構造のことを指し(図1左)、ボトムアップ構造は、トップダウン構造とは逆に、先の段落で部分的な話の詳細を述べ、結論や全体のまとめを後の段落で述べている構造のことを指す(図1右)。

図2にトップダウン構造とボトムアップ構造の具体例を示し、図3に図2の各段落に使われている単語を示す。

まず、トップダウン構造に着目する。第1段落とつながっている第2段落では単語Aのみ共通に使われており、それぞれの段落の単語の使用割合は、第1段落では、3単語中に1単語、第2段落では7段落中に1単語である。第1段落よりも第2段落では使用単語が増えていることから、第1段落から第2段落に単語Aについて話が展開していると捉えることができ、矢印の向きを付けることが出来る。第3段落、第4段落にも着目すると、第3段落では単語Bについて、

第4段落では単語Cについて話が展開していると捉えられる。

次に、ボトムアップ構造に着目する。第5段落とつながっている第2段落間をトップダウン構造と同様に考えると、単語Aを共通に使われ第2段落から第5段落へ単語Aについて話がまとまっていると捉えることが出来る。同様に、矢印の向きを付けることが出来る。ただし、文章は小さい段落番号から大きい段落番号へ話が進むため、この場合、第5段落から第2段落へ話が広がっているのではなく、第2段落から第5段落へ話がまとまっていると考える。

結論が先に書かれている場合、第5段落を取り除くとトップダウン構造の文章、逆に結論が最後に書かれている場合、第1段落を取り除くとボトムアップ構造の文章になる。

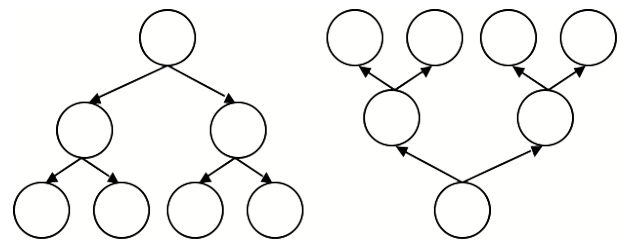


図1: トップダウン構造とボトムアップ構造

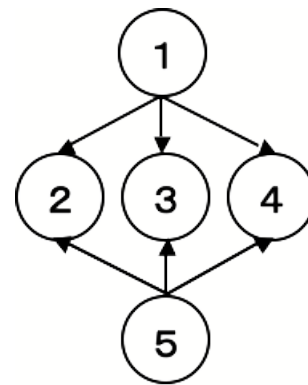


図2: トップダウン構造とボトムアップ構造の例

1段落	2段落	3段落	4段落	5段落
ABC	A DEFGHI	BJKLMNO	CPQRSTU	ABC

図3: 図2の各段落で使われている単語

4 段落間関係の評価に基づく文章構造推敲支援システム (SAT)

4.1 段落間関係の評価に基づく文章構造推敲支援システム (SAT) の構成

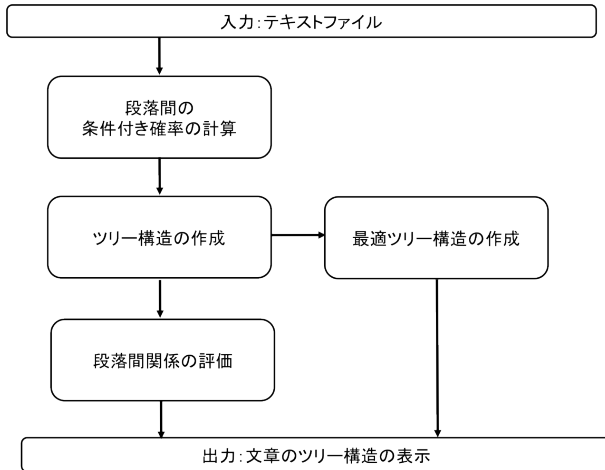


図 4: 段落間関係の評価に基づく文章構造推敲支援システム (SAT) の構成

本節では段落間関係の評価に基づく文章構造推敲支援システム (SAT: Segment Association Tree) の構成について述べる。以下, SAT システムと呼ぶ。

図 4 に SAT システムの全体の構成を示す。SAT システムにテキストファイルを与え、それを元に段落間の条件付き確率を計算を行い、リンクをつなぐ段落間を決定する。次に、段落間の条件付き確率を元にリンクをつなぎ、文章のツリー構造を作成し段落間関係の評価を行う。また、ツリー構造を推敲するために、段落並べ替えによる最適ツリー構造を作成する。最後に、それぞれの結果をインタフェース上へ出力し、ツリー構造と最適ツリー構造の比較を行い、文章構造の推敲をする。

以下、各処理について述べる。

4.2 段落間関係の評価

本節では、トップダウン構造とボトムアップ構造の2つで表される段落間関係の評価方法について述べる。

4.2.1 段落間の条件付き確率の計算

本項では、段落間の条件付き確率の計算方法について述べる。以下の式 (1) で全ての段落間の条件付き確率 $Relation(A, B)$ を計算をする。段落 A で使われてい

る単語集合を W_A 、段落 B で使われている単語集合を W_B とし、それらを数える関数 n へ与える。段落 A で使われている単語集合中、段落 A と段落 B で共通して使われている割合を求めることによって、段落間のリンクに向きをつけることが出来ると考えられる。

$$Relation(A, B) = \frac{n(W_A \cap W_B)}{n(A)} \quad (1)$$

4.2.2 段落間関係の評価

本項では、段落関係の評価方法について述べる。

以下、後述する段落間の関係を表すツリー構造 (段落をノードとし、枝には条件付き確率 $Relation(A_i, A_j)$ が与えられる) が作成されたとき、 n 段落構成の文章のトップダウン構造の評価値を与える際のアルゴリズムを示す。

1. 各段落 A_i について、枝に分かれしている枝数 $Branch_i$ をカウントする。
2. 各段落 A_i の、 $Branch_i$ が 2 以上の段落 A_i の各枝 j について、最も長い葉ノードまでの枝に与えられている、条件付き確率 $Relation(A_i, A_j)$ を加算し、 $BValue_{ij}$ とする。
3. 各段落 A_i について、すべての $BValue_{ij}$ を掛け合わせた $NValue_i$ を求める。
4. $i < n/2$ の段落の $NValue_i$ を加算、 $i \geq n/2$ の段落の $NValue_i$ を減算した合計値を、トップダウン構造の評価値 $TopValue$ とする。

なおステップ 4 で、話を広げるのは前半段落、話をまとめるのは後半段落、の考えから前半の段落にトップダウンがある場合は加算し、後半の段落にある場合は減点する。

またボトムアップ構造の評価値は、トップダウン構造の評価値の計算アルゴリズムと同様で、作成されるツリー構造の天地逆転をしてできるツリーのトップダウン構造の評価値と同じとする。

4.3 文章の強リンク構造の作成

本項では、文章の強リンク構造を作成する方法について述べる。強リンク構造とは、 n 段落構成の文章について、各段落をノード、ノード間のリンクを条件付き確率 $Relation(A_i, A_j)$ として、最少の強いリンクのみで作成したツリー構造を指す。以下に、強リンク構造を作成するアルゴリズムを示す。

1. 全ての段落 $A_i, A_j (0 \leq i, j \leq n, i \neq j)$ 間の条件付き確率 $Relation(A_i, A_j)$ を計算する。

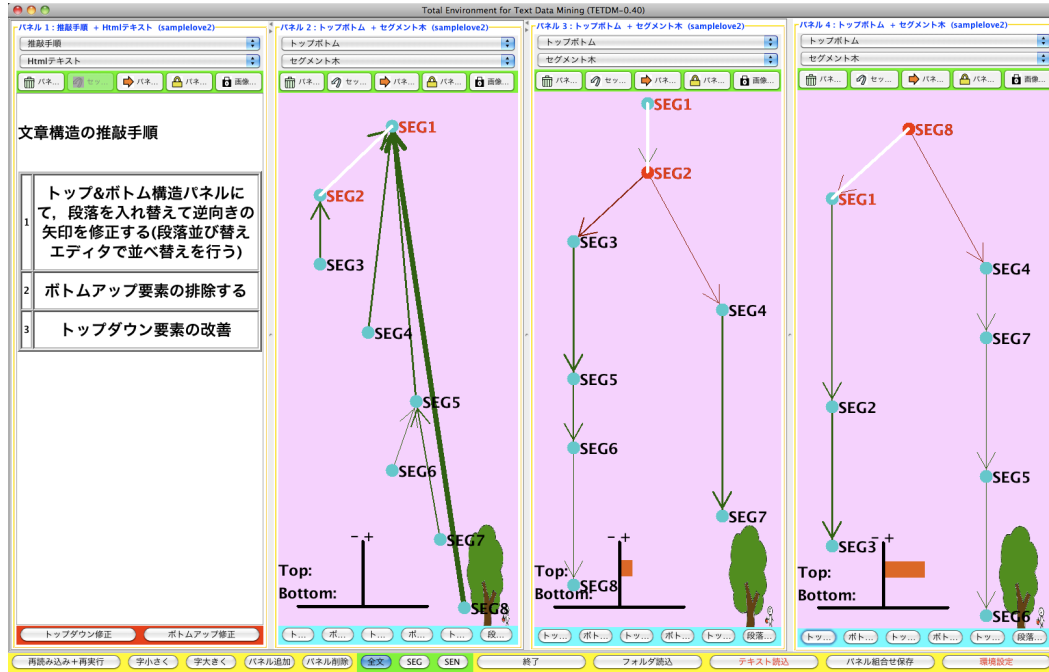


図 5: 文章構造推敲支援システムの使用例

2. 各段落 A_i に関して、段落 A_i との条件付き確率 $Relation(A_i, A_j)$ が最も高い段落 A_j との間にリンクを生成する。(この時点ではリンクの向きは考えない)。
3. n 個の段落のすべてがリンクでつながっていれば (リンクによる段落間のパスが存在すれば) 5.へ。そうでなければ、全段落間の条件付き確率 $Relation(A_i, A_j)$ を大きい順にソートする。
4. ソートされた条件付き確率 $Relation(A_i, A_j)$ の大きい順に、段落 A_i と段落 A_j との間にリンクを生成し、すべての段落がリンクでつながるまでリンクの生成を繰り返す。
5. 第1段落を根ノード (親ノード) とし、つながっているノードを順に子ノードとしたツリー構造を生成する。
6. リンク生成に用いた条件付き確率 $Relation(A_i, A_j)$ をもとに、リンクに $A_i \rightarrow A_j$ の向きを与える。

まず、全ての段落 A_i, A_j 間の条件付き確率 $Relation(A_i, A_j)$ の計算を行い、各段落 A_i に関して、条件付き確率 $Relation(A_i, A_j)$ が最も高い A_j との間にリンクを引く。次に、 n 個の段落の全てがリンクでつながれていなければ、全段落間の条件付き確率 $Relation(A_i, A_j)$ を計算し、まだつながれていない段落間で一番大きい段落間にリンクをつなぐ。一番条件付き確率が高い段落間にリンクを引くことによって、つながりの強いツリー構造が出来ると考えられる。

4.3.1 トップダウン構造の作成

本項は、トップダウン構造の作成について述べる。トップダウン構造のアルゴリズムは、4.3 のアルゴリズムのステップ2で段落 A_j に着目し、下向きの矢印が入ってくる段落 $A_i (i < j)$ の中で、条件付き確率 $Relation(A_i, A_j)$ が一番高い段落間にリンクを生成する。また、他のステップは同様である。下向きの矢印の条件付き確率 $Relation(A_i, A_j)$ に着目することによって、トップダウンの段落からつながっている段落元を探すことができる。

4.3.2 ボトムアップ構造の作成

本項はボトムアップ構造の作成について述べる。ボトムアップ構造のアルゴリズムは、トップダウン構造と同様に4.3のアルゴリズムのステップ2が異なる。段落 A_i に着目し、上向きの矢印が入ってくる段落 $A_j (i < j)$ の中で、条件付き確率 $Relation(A_i, A_j)$ が一番高い段落間にリンクを生成、他のステップは同様である。上向きの矢印の条件付き確率 $Relation(A_i, A_j)$ に着目することによって、ボトムアップ構造からつながっている段落元を探すことができる。

4.4 最適ツリー構造の作成

本節では、最適ツリー構造の作成について述べる。最適ツリー構造とは、文章内の段落の順番を任意に変更

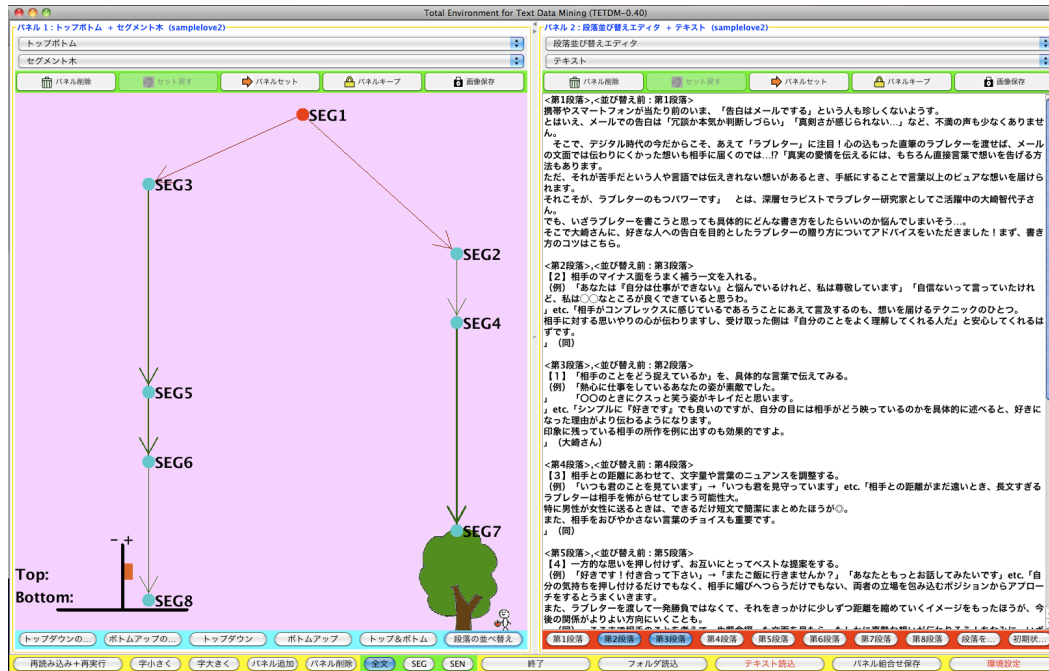


図 6: 段落並び替え可能のインタフェース

する中で、トップダウン構造、またはボトムアップ構造の評価値が最も高くなる構造のことを指す。

1. 文章内の段落 (段落数 n) を並べ替えて出来る $n!$ 通りの全てのパターンを作成する。
2. 作成した各パターンについてツリー構造を作成し、評価値を計算する。
3. 評価値が最も高くなった段落の並びについて、先頭の段落を根ノード (親ノード) つながっているノードを順に子ノードとしたツリー構造を生成する。

これにより現時点のツリー構造と最適ツリー構造との比較を促し、文章構造の推敲を支援する。

4.5 出力：文章のツリー構造の表示

本節では、システムの出力：文章のツリー構造の表示について述べる。4.2 段落間関係の評価、4.4 最適ツリー構造で述べたものを TETDM[11] へ実装している。

SATシステムの出力例を図5へ示す。文章は、mixi[12]のコラムに掲載されていた「ラブレターの書き方のアドバイス」を用いている。段落間をつなぐリンクを線で表し、条件付き確率の大きさによってリンクの太さを変更し pixel で表す。トップダウン構造、ボトムアップ構造の場合は、該当する段落とそのリンク先の色を変える。また、左下にトップダウン、ボトムアップの評価値の値を表示している。

4.5.1 各パネルの表示

図5より、左パネルから順に、文章構造の推敲手順、4.3 節の強リンク構造、4.3.1 項トップダウン構造、4.4 節トップダウンの最適構造を表示しており、これらを用いて文章構造の推敲を行う。強リンク構造では、矢印の向きが全て上を向いており、あとの段落から第1段落をまとめていると解釈が出来る。次に、真ん中のパネルのトップダウン構造では、第2段落は下向きに複数の段落に話を広げていることが分かり、第2段落は話を展開していると読み取ることができる。最後に、右パネルのトップダウンの最適構造では、段落の並び替えによる最適ツリー構造を表示している。

4.5.2 並び替え可能のインタフェース

本項は、並び替え可能のインタフェースについて述べる。

図6に段落並び替え可能のインタフェースを表示する。右側のパネルは文章を示しており、背景が赤いボタンに段落番号を表示している。並び替えを行いたい2つの段落を選んで「段落を並び替える」ボタンを押すことによって段落を入れ替えることが出来る。次に、左のツリー構造パネルで「並び替えの結果」ボタンを押すと、右パネルで並び替えた結果を反映し、最適ツリー構造を表示を行う。このインタフェースを使用することによって、適宜、段落の順番を並び替えてツリー構造を確認することができる。

4.5.3 文章構造推敲支援システム (SAT) の使用の流れ

本項は、文章構造をトップダウン構造に修正したい場合の SAT システムの使用の流れについて述べる。なお、ボトムアップ構造に修正した場合はトップダウン構造の場合と逆になる。

1. 強リンク構造のパネルで、段落を入れ替えて逆向きの矢印を修正する (段落並び替え可能のインタフェースで並べ替えを行う)
2. 同パネルで、つながっている段落間で共通に使っている単語を減らすことでボトムアップ要素の排除を行う。
3. 現状のトップダウン構造より更にトップダウン構造にするため、トップダウンの最適ツリー構造と比較を行い、よりトップダウン構造になる段落間に共通単語を多く使う。また、段落の順序入れ替えを行う。

以下、各手順について述べる。

1 について、段落の順番を入れ替えて文章の修正を行う。トップダウン構造より修正を行うため、矢印は全て下向きにする。段落の入れ替えが出来ない場合、この手順を飛ばす。

2 について、ボトムアップ要素の排除を行う。文章をトップダウン構造にしたい場合、ボトムアップ構造を排除する必要があるため、複数の段落と話がつながっている段落を一緒にする。または、共通に使う単語数を減らして複数段落をつなぐのではなく、1つの段落のみをつなぐ修正する。

3 について、現状の構造よりも更にトップダウン構造になる修正を行う。主に、該当する段落間に共通単語を使用する。

5 結論

本研究では、文章を段落間のつながりに着目した段落感関係の評価に基づく文章構造推敲支援システム (SAT システム) を提案した。文章構造の推敲のために、ツリー構造表示と最適ツリー構造表示を行い、それぞれのツリー構造を比較することにより、文章構造を推敲を行うことを示した。

今後は、実装した SAT システムの精度を計る実験を行うために、実験で使う文章の収集、準備をおこなっていく。

参考文献

- [1] 板倉由知, 白井治彦, 黒岩丈介, 小高知宏, 小倉久和:単語の概念関係を用いた段落一貫性評価指標の有効性, 情報処理学会研究報告, NL-183, pp.107-113, (2008)
- [2] 板倉由知, 白井治彦, 黒岩丈介, 小高知宏, 小倉久和:様々な文書を対象とした段落一貫性の解析:情報処理学会研究報告, NL-192, pp.1-6, (2009)
- [3] 春日隆緒, 田村直良:文章のセグメント間関係解析に基づく文章構造解析:情報処理学会研究報告, NL-155, pp.59-64, (2003)
- [4] 藤田彬, 田村直良:文章構造解析に基づく小論文の理論性についての自動採点, 第9回情報科学技術フォーラム講演論文集 7(2), pp.21-22, (2008)
- [5] 奥出 信一郎:情報理論による, 文章の理論的構造の解析, 電子情報通信学会技術研究報告, pp.1-5, (2008)
- [6] 松本章代, 山田未央佳, 山田翔, 鈴木雅人:理工系学生を対象とした技術文書作成支援システム, 情報処理学会研究報告, CE-98, pp.91-96, (2009)
- [7] 大野博之, 稲積宏誠:文の構造デザインツールの試作と校正・推敲支援ツールとの連携, 電子情報通信学会技術研究報告, ET, 教育工学 108(146), p.73-78, (2008)
- [8] 奥村有希, 大野博之, 稲積宏誠:技術文章作成支援ツールの推敲支援機能の拡張-長い修飾節に起因する悪文の検出手法の提案 (次世代情報教育の構築に向けて/一般), 電子情報通信学会技術研究報告, ET, 教育工学 107(536), (2008)
- [9] 須藤崇志, 丸山広, 中村太一:文を分かりにくくする要因の分析と改善支援手法の提案, 電子情報通信学会技術研究報告, KBSE, 知能ソフトウェア工学 108(65), p41-46 (2008)
- [10] 松本裕治, 山下達雄, 平野義隆, 松田寛, 高岡一馬, 浅原正幸:形態素解析システム『茶筌』, Ver.2.4.0, 使用説明書, (2007)
- [11] 砂山渡, 高間康史, ダヌシカ ボレガラ, 西原陽子, 徳永秀和, 串間宗夫, 松下光範:テキストデータマイニングのための統合環境, 人工知能学会論文誌, Vol.26, No4, p483-493, (2011)
- [12] mixi:<http://mixi.jp/>