

視線運動解析による興味アスペクトの推定

下西 慶^{*1} 石川 恵理奈^{*1} 米谷 竜^{*2} 川嶋 宏彰^{*1} 松山 隆司^{*1}

Learning Aspects of Interest from Gaze

Kei Shimonishi^{*1}, Erina Ishikawa^{*1}, Ryo Yonetani^{*2}, Hiroaki Kawashima^{*1} and Takashi Matsuyama^{*1}

Abstract – This paper introduces a novel model for estimating the aspects of items that we are interested in (aspects of interest) when browsing contents. The proposed model generates items of gaze from the aspects of interest, where the aspects are characterized by associating them with attributes of items such as their specifications and appearances, and the aspects of interest are modeled by the mixtures of the aspects. Then, we can achieve aspects in a data-driven manner and represent aspects of interest flexibly via the estimation of the proposed model. Once we learn the aspects, we can also estimate the aspects of interest from newly observed gaze data. Furthermore, we incorporate influences of content layouts upon regions of gaze such as high conspicuity of center regions. Our experiments demonstrate the appropriateness of proposed model by accurately predicting items that users are interested in and likely to be looked at from estimated aspects of interest.

Keywords : Aspect, Gaze, Content design, Probabilistic generative model, Mental state

本著作物の利用に関する注意：本著作物の著作権は特定非営利活動法人ヒューマンインタフェース学会に帰属します。本著作物は著作権者であるヒューマンインタフェース学会の許可の元に掲載するものです。ご利用にあたっては「著作権法」に従うことをお願いいたします。

1. はじめに

街角に設置されたデジタルサイネージや空港、ショッピングモールなどに設置されたナビゲーションインタフェースといったコンテンツ提示システムの設計において、システムを利用するユーザが提示コンテンツ中のどこに注視を向けているかという視線情報は、その背後に潜むユーザの心的状態を探る重要な手がかりとなる。ユーザの視線情報から心的状態を推定することができれば、ユーザの興味に応じて新たなアイテムを提示する、注視の意図（例：複数対象の見比べや特定対象の吟味）に応じてコンテンツレイアウトを切り替える、コンテンツ閲覧に対する集中状態に応じて情報提示タイミングを制御するというように、システムがユーザに対して適応的かつインタラクティブに情報提示を行うことが可能であり、すでにいくつかのシステムが提案されている [4], [7], [13], [14], [17]。我々は特に、人間の心的状態を推定し適切な情報を推薦することで、ユーザの意思決定を支援するようなシステムの構築を目指しており、本論文ではその基盤技術として心的状態から視線情報の生成過程（注視行動）のモデリングおよび視線情報からの心的状態推定に取り組む。

ユーザがコンテンツを閲覧する状況において、あるアイテムへの注視の背後には、従来扱われてきた「コンテンツ中のどのアイテムに興味を持っているか」[7]に加え「アイテムのどのような側面（アスペクト）に興味を持っているか」という心的状態を考えることができる。これを本研究では興味アスペクトと呼ぶ。たとえば、料理カタログというコンテンツ（4.2節も参照のこと）に対して「健康に良い」「調理が簡単」といった複数アイテムに共通する状態を考えることができ、これらは提示されていないアイテムへの興味を扱う必要がある応用場面、たとえば情報推薦において有効な手がかりとなることが期待される。しかしながら、ユーザの視線情報から興味アスペクトを推定するにあたっては、提示コンテンツのドメイン（たとえば、料理カタログや地図）やユーザの個人差に起因する多様性が課題となる。すなわち、閲覧時に現れうる心的状態をあらかじめトップダウンに想定しておき、注視との相関性を統計的に学習するという従来のアプローチ（詳細は後述）を導入することは難しい。

そこで本研究では、興味アスペクトの表現法および興味アスペクトに基づく注視行動のモデルを新たに提案するとともに、視線情報からのモデルの学習を通して、閲覧時に現れうるアスペクト自体をデータドリブンに獲得するアプローチを導入する（図1(a)）。これにより、獲得されたアスペクトを用いて新たに得ら

*1: 京都大学 大学院情報学研究科

*2: 東京大学 生産技術研究所

*1: Graduate School of Informatics, Kyoto University

*2: Institute of Industrial Science, the University of Tokyo

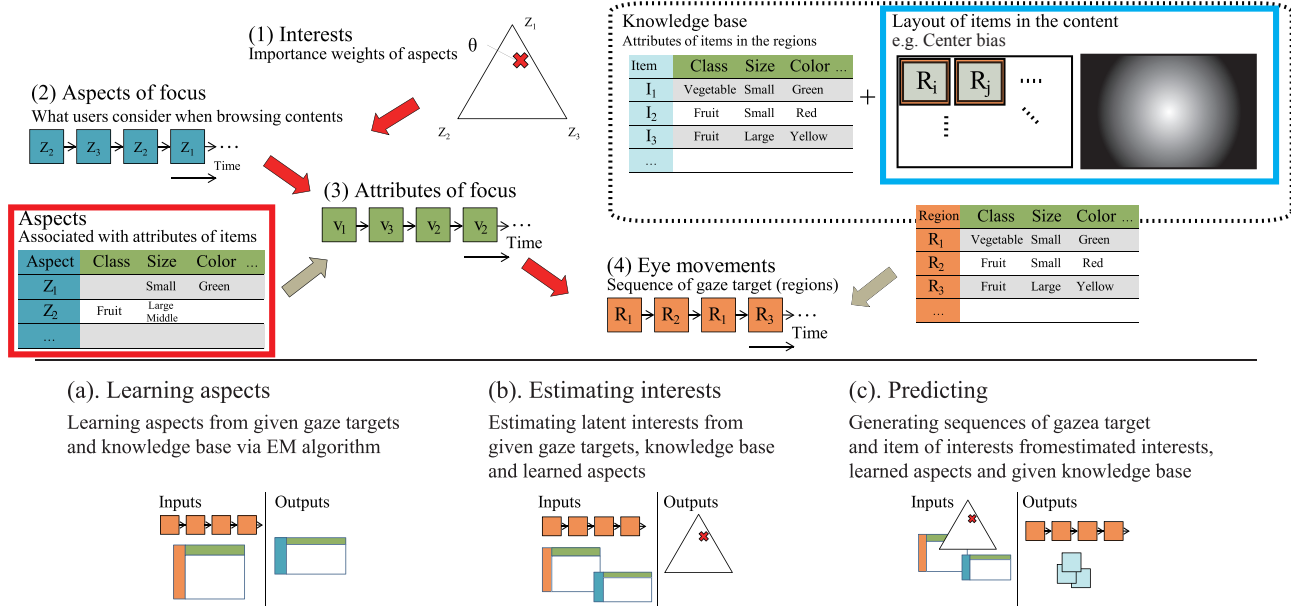


図1 提案モデルの概要
Fig.1 Overview of the proposed model

れた視線情報から興味アスペクトを推定できるようになるとともに (図1 (b)), 推定された興味アスペクトを用いて注視領域や興味アイテムを予測することも可能になる (図1 (c)). 提案モデルは, 確率的生成モデルの1つである確率的潜在意味解析 (probabilistic Latent Semantic Analysis; pLSA) [8] を拡張したものであり, ユーザは各時間において興味アスペクト (図1 (1)) に基づいてアスペクトの1つに着目し (図1 (2)), 着目アスペクトに関連する属性値の1つに着目する (図1 (3)). さらにその着目属性値に基づいてアイテム領域を注視する (図1 (4)) という生成過程を表現している. これはユーザの注視行動を非常に単純化したモデルであるものの, 応用面からは以下のような特長を備えている.

提案モデルの特長

1. 注視アイテム自身ではなく, アイテムのメタ情報 (例: 種別や大きさ) や見えに関する情報 (例: 色や明度) と紐づけてアスペクトを表現する (図1 赤枠部). ここで注視アイテム自体を直接用いないことによって, 先述した複数アイテムに共通する状態を表現可能となり, 情報推薦といった応用場面において提示されていないアイテムへの興味を扱うことも可能となる.
2. 注視行動の生成過程において, コンテンツ自体のデザインが注視アイテムに及ぼす影響を考慮する. 具体的には, コンテンツレイアウト (アイテムの配置) によって特定のアイテム領域に注視が誘導されるという現象を生成過程に組み込む (図1 青枠部). これにより, あるアイテム領域への注視

が興味アスペクトに起因するものかレイアウトに起因するものかを考慮しつつ, アスペクトの学習や興味アスペクトの推定が可能となる.

関連研究

視線運動からの心的状態推定は, 視覚心理, ヒューマンコンピュータインタラクションといった分野を中心に積極的に取り組まれているトピックであり, 興味 [4], [5], 意図 [1], [6], [9], 集中 [20] といった状態の推定手法がこれまでに提案されている. これらの手法は主に, 視線情報 (あるいは視線と閲覧コンテンツの関係性の情報) を特徴ベクトルとして表現するとともに, 実験統制によって与えられたカテゴリー的な心的状態 (たとえば, アイテム A, B, C, ... への興味, 集中度合の高, 低など) に関する識別モデル, 関数を統計的に学習するというアプローチをとる. これにより, 新たに得られた視線情報から, その時々々の心的状態を推定することが可能になる. ただし, 上述のアプローチではユーザが取りうる状態をあらかじめトップダウンに定める必要があり, 多様な興味アスペクトが現れうる本研究に適用することは難しい.

一方で, 視線情報から心的状態をデータドリブンに学習するというアプローチも, 情報検索 (ランキング) [15], [16], [18] や推薦 [19], インタラクションマイニング [11], [21] といった分野において提案されている. ただし, これらのアプローチには提示アイテム [18], [19] (あるいはインタラクションの対象 [11], [21]) への興味のみを推定の対象とするため, 推定時に提示されていないアイテムに対する興味を扱うことはできない, あるアイテム領域の注視が興味に起因するものかコンテンツ



図 2 実験環境．画像は(株)ミツカン (<http://www.mizkan.co.jp/index.html>), および「おいしい」を科学して、レシピにしました。」(サリー(著), サンマーク出版, 撮影: 久保田 健, 2013)より許可を得て掲載．

Fig.2 Experimental environment

デザインに起因するものかが判断できない^{[15], [16]}といった限界がある．これに対して, 提案モデルは先述した特長により, このような問題に対処することが可能である．

本論文の構成

次章ではまず状況設定, 視線情報と提示コンテンツ, および興味の表現方法について整理する．そして, 3. 章において興味アスペクトに基づく注視行動モデルの定式化および視線情報からの興味アスペクト推定法を提案するとともに, 4. 章の実験においてその有効性を評価する．本研究で提案するモデルは, 注視行動に関していくつかの仮定を置いた基礎的なものである．5. 章では, それらの仮定に起因するモデルの限界および拡張について議論し, 6. 章でまとめる．

2. コンテンツ閲覧時の注視行動

商品カタログなどの複数アイテムの情報を含むコンテンツがディスプレイに提示され, ユーザがそれを閲覧している状況を想定する(図2)．実際の状況では, ユーザの入力などに応じて提示アイテムやそのレイアウトが動的に変化することも考えられるが, 本研究では問題の簡単化のため, 提示コンテンツは固定とする．本章では, まず本研究におけるコンテンツの表現方法, つまりアイテムのメタ情報や見えに関する情報をどのように扱うかについて説明する．続いて視線情報の表現方法, および興味アスペクトのコンセプトに関する詳しい説明を述べる．

2.1 提示コンテンツ

2.1.1 アイテムとその属性

N_{all} 個のアイテムからなるアイテム集合 $\mathcal{I} = \{I_1, \dots, I_{N_{all}}\}$ を考える．ここで, アイテムが持つ,

種別や大きさなどといったメタ情報と, 色や明度といったコンテンツ上での見えに関する情報を, それぞれ意味属性, およびアピランス属性と呼ぶ．また, これらの意味属性とアピランス属性とを合わせて, 各アイテムは P 種の共通の属性を持つとする．各属性 $p \in \{1, 2, \dots, P\}$ は Q_p 種類の値(属性値)を取るものとし, 属性 p の属性値集合を $\mathcal{V}_p = \{V_{p1}, \dots, V_{pQ_p}\}$ とする．以下では, この全属性が取りうる値の集合 $\mathcal{V} = \{V_{11}, \dots, V_{1Q_1}, V_{21}, \dots, V_{PQ_P}\}$ を, 添え字を置き換えて $\mathcal{V} = \{V_1, \dots, V_Q\}$ ($Q = |\mathcal{V}|\})$ と表す．

意味属性は, アイテムの性質を直接的に表現できるため, ユーザがアイテムを比較, 吟味する際の判断基準に関連する重要な要素となりうる．しかしながら, これを利用するためには, あらかじめ各アイテムに対する手動のアノテーションがしばしば必要であり, その煩雑さが課題となる．一方で, アピランス属性は, あくまでコンテンツ上でのアイテムの表現(見え方)に関する性質であるため, アイテム自体の性質を陽に表すものではないが, 画像処理技術を適用することによって自動的に獲得できるという利点がある．なお, 4. 章の実験では用いる属性の違いについて比較評価する．

2.1.2 アイテム領域

各アイテムは, コンテンツ上において画像やテキストといったいくつかのメディアによって表現される．このようなメディアの占める各領域を R とする．本論文では簡単のため, 各アイテムがそれぞれ1枚の画像で表現されるとする．これにより, ディスプレイの全画面領域 Ω ($\Omega \subset \mathbb{R}^2$) は, N 個のアイテムに対応した N 個の領域集合 $\mathcal{R} = \{R_1, \dots, R_N\}$ ($R_i \subset \Omega$, $R_i \cap R_j = \phi$ ($i \neq j$), $\bigcup_{i=1}^N R_i = \Omega$) に分割される．

2.2 視線情報

コンテンツ閲覧状況の視線解析にあたっては, どのアイテムが注視されているか, という情報がしばしば有効な手がかりとして用いられる^{[9], [14]}．そのため本研究では, ユーザの視線情報を, 注視されているアイテム領域の系列として表現する．具体的には, まず右目左目それぞれについての視線と画面の交点を計測し, それらの中点を注視点と定義する．さらに, この注視点の系列として獲得された視線運動から, ある時間 t における注視領域を $r_t \in \mathcal{R}$ として, 注視領域系列 $(r_1, \dots, r_T)$ を算出する．この時間 t は, 視線計測時に与えられる固定サンプリング時刻ではなく, 注視の切り替わりに基づいて定められる($r_{t-1} \neq r_t$ とする)．この注視点系列からの注視領域系列の算出において, 留まる時間が閾値(10 サンプリング点(約167ミリ秒))以下の領域は, 計測ノイズや, ある領域が

ら別の領域へと注視領域が移る際に一時的に通過しただけの領域だと考え、注視領域系列から除外する。なお、注視対象を、アイテム自体ではなくコンテンツ上での領域とすることにより、コンテンツのレイアウトに起因する誘目要素の影響をモデルに陽に組み込むことが可能になる。これについては次章で詳述する。

一回の連続した閲覧行動をセッションと呼ぶ。\$S\$ 個のセッションからなるセッション集合 \$S\$ が与えられるとし、このうち \$s\$ 番目のセッションにおいて時刻 \$t\$ に注視された領域を \$R_{n_t}^s\$ と表すこととする。すると、注視領域系列は \$(R_{n_1}^s, \dots, R_{n_{T_s}}^s)\$ となるが、これをまとめて \$\{R_{n_t}^s\}\$ のように表す。この \$\{R_{n_t}^s\}\$ において、領域 \$R_n\$ がどれだけの頻度で注視されたかを、領域注視頻度

$$n(s, R_n) = \sum_{t=1}^{T_s} \delta(R_{n_t}^s = R_n) \quad (1)$$

として表す。また、全セッションでの注視領域系列の集合を \$\mathcal{D}\$ とする。

2.3 興味アスペクト

本研究において導入する興味アスペクトとは「アイテムのどのような側面（アスペクト）に着目しているか」を表現する心的状態である。たとえば、図2のような料理画像が並べられたコンテンツを閲覧する際には「健康に良い」「たくさん食べたい」といったアスペクトに基づいて複数のアイテムを順に注視していくことが想定される。

\$K\$ 種類のアスペクト \$\mathcal{Z} = \{Z_1, \dots, Z_K\}\$ が得られている状況を考える。このとき、ユーザの興味アスペクトは、各アスペクト \$Z_k\$ への重み（興味度合い）によってモデル化することができる。具体的には、\$Z_k\$ への重みを \$\theta_k\$ とし、ユーザの興味アスペクトをパラメタ \$\theta = (\theta_1, \dots, \theta_K)\$ （\$\theta_k \geq 0\$, \$\sum_{k=1}^K \theta_k = 1\$）として定義する。この興味アスペクトは、\$K-1\$ 次元 simplex 上の1点として表すことができる。ここでは簡単のため、同一セッション内においてユーザの興味アスペクトは時間的に変化しないものと仮定する。さらに、セッション内での各時刻における注視領域は、それ以前の時刻における注視領域とは独立であるとする。

3. 提案モデル

本章では、興味アスペクトに基づく注視行動の確率的生成モデルを提案する。提案するモデルはアイテムの持つ知識ベース、およびコンテンツのレイアウトの影響を取り入れたモデルである。また、同モデルを用いた興味アスペクトの推定法や、注視領域の予測手法についても述べる。

3.1 モデルの定式化

1. 章で述べたように本研究では、ユーザがあるアイテムを注視している際、アイテム自身に注目しているのではなく、そのアイテムに関連したアスペクトに注目していると考える。ただし、複数コンテンツ、複数ユーザを扱い、さらにはユーザによる自由な閲覧を想定する場合、現れるアスペクトをあらかじめトップダウンに列挙しておくことは困難である。そこで本研究では、2.1.1 節で述べた属性値とアスペクトを紐づけることによって、アスペクトを複数属性（たとえば「緑色」の「野菜」）との関連度としてデータドリブンに学習、表現するというアプローチをとる。

提案モデルにおける注視行動の生成過程は以下の通りである。まず、セッション \$s \in S\$ における興味アスペクト \$\theta(s)\$ によって各時間 \$t\$ での注目アスペクト \$z_t \in \mathcal{Z}\$ が分布 \$P(z_t = Z_k | s) = \theta_k(s)\$ に従って決定されるとする。次に、ある属性値 \$v_t \in \mathcal{V}\$ が条件付き分布 \$P(v_t = V_q | z_t)\$ に従って選択されるものとする。ここで、\$v_t\$ を時間 \$t\$ における注目属性値と呼ぶ。各アスペクト \$Z_k\$ は、属性値 \$V_q\$ への関連度 \$P(V_q | Z_k)\$ によってモデル化されることになる。最後に属性値 \$v_t\$ を持つアイテムに視線が引き寄せられる形で（後述のようにこれは \$P(r_t | v_t)\$ としてモデル化する）領域 \$r_t \in \mathcal{R}\$ が注視されることになる。

このとき、簡単化のために、アスペクトと属性値の関係はセッションによらないものとし、さらにユーザはある時刻で1つの属性値にのみ注目すると仮定する。さらに、条件付き分布 \$P(r_t = R_n | v_t = V_q)\$, \$P(v_t = V_q | z_t = Z_k)\$ はいずれも時刻によらないものとする。すると、あるセッションにおける注視領域、注目属性値、注目アスペクトの同時確率は \$P(R_n | V_q)\$, \$P(V_q | Z_k)\$, \$P(Z_k | s)\$ を用いて

$$\begin{aligned} P(r_t = R_n, v_t = V_q, z_t = Z_k | s) \\ = P(R_n | V_q) P(V_q | Z_k) P(Z_k | s) \end{aligned} \quad (2)$$

となる（図3（b））。このとき、ある領域 \$R_n \in \mathcal{R}\$ が注視される確率は式（2）を \$V, Z\$ に関して周辺化し、

$$P(r_t = R_n | s) = \sum_{q=1}^Q \sum_{k=1}^K P(R_n, V_q, Z_k | s) \quad (3)$$

として得られる。

3.2 レイアウトの構造を考慮した視線の確率的生成モデル

コンテンツ閲覧中の注視行動をモデル化するにあたっての重要な要素として、コンテンツのレイアウトによる影響が考えられる。たとえば、注視予測に関するいくつかの従来研究では、画面中央に注視が向きやすいといった事前知識が導入されている [3], [12]。提案す

視線運動解析による興味アスペクトの推定

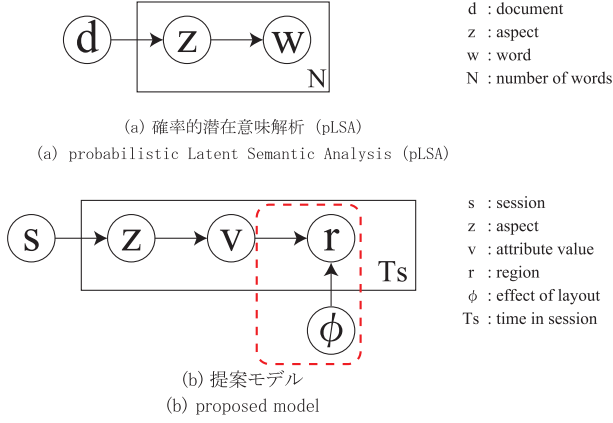


図3 pLSA と提案モデルのグラフィカルモデル

Fig. 3 Graphical model of pLSA and proposed model

るモデルでは、コンテンツレイアウトが注視行動に与える影響を事前にモデル化しておくことで、このような影響に起因する注視行動を考慮することができる。

提案モデルについての前節の式(2)は確率的生成モデルの1つであるpLSAをベースとしたものであり、観測データを説明する少数の隠れ変数(アスペクト、トピック)を獲得することを目的として、自然言語処理やデータマイニングの分野においてもしばしば同様のアプローチがとられる[2]。pLSAでは、基本的に観測データ(単語)間の構造は扱われないが、その一方で視線情報の解析においては、コンテンツ上でのアイテムのレイアウトや注視アイテムの時間的順序など、構造的な情報が重要となる。そこで本研究では、図3(a)に示されるpLSAを、特にレイアウトが視線に与える影響を考慮するため、以下のように拡張する(図3(b)参照)。pLSAでは、注目アスペクト z から単語 w が生成され観測されるが、提案モデルでは、まず注目アスペクト z から外部から観測されない注目属性値 v が生成される。そして、この注目属性値 v から実際に観測されるユーザの注視領域 r が生成される。注視領域 r はコンテンツレイアウトにも影響を受けると考えられるので、図3(b)赤枠点線部のように、レイアウトによる影響(ϕ から r への矢印)を導入する。

このようにレイアウトの影響を扱うにあたって、前節の定式化において、 $P(I_n|V_q)$ を直接考えるのではなく、 $P(R_n|V_q)$ とする点が重要である。すなわち、心的状態から最終的に各時刻の「注視領域」が決定(観測)されるプロセスとは、単に興味に近いアイテムに視線が引き寄せられるのではなく、アイテムの配置(たとえば中央が目を引きやすい)や直前の注視アイテムなどの各種構造によって「変調」されていると考える。これは、図3(b)においてパラメータ ϕ で与えられ

る部分である。レイアウトの影響としては、絶対的な配置、前時刻に注視していた領域からの相対的な配置の2つが考えられるが、本論文では、まず画面上での絶対的な配置による影響を扱うこととする。

具体的には、以下のようにして $P(R_n|V_q)$ を導出する。まず、ある属性値 V_q が与えられた際に、提示されているコンテンツ内のあるアイテム I_n が選択される確率 $P(I_n|V_q)$ が与えられているとする。ここで、コンテンツ内で I_n を提示している領域 R_n がアイテム I_n に関わらず注視される確率が $P(R_n; \phi)$ で与えられるとした時、最終的に V_q が与えられた場合に領域 R_n に視線が向く確率を、混合比 α ($0 \leq \alpha \leq 1$)を用いて、

$$P(R_n|V_q) := (1 - \alpha)P(I_n|V_q) + \alpha P(R_n; \phi) \quad (4)$$

としてモデル化する。また、この $P(R_n|V_q)$ は、パラメタ α, ϕ に依存するが、本論文では $P(R_n|V_q)$ と省略して記す。

これらを用いて、セッション s における視線系列 $\{R_{n_t}^s\} = (R_{n_1}^s, \dots, R_{n_{T_s}}^s)$ の生成確率は、次のようになる。

$$P(\{R_{n_t}^s\}) = \prod_{t=1}^{T_s} P(R_{n_t}^s | s) \\ = \prod_{t=1}^{T_s} \left\{ \sum_{q=1}^Q \sum_{k=1}^K P(R_{n_t}^s | V_q) P(V_q | Z_k) P(Z_k | s) \right\} \quad (5)$$

3.3 モデルパラメタの学習、推定

本節では、モデルパラメタについて図1で示したそれぞれのパラメタの推定手法について述べる。本論文では、レイアウトに起因する注視確率 $P(R_n; \phi)$ として、観測された全セッションのデータから算出された相対領域注視頻度を用いる。また、各属性値からの領域選択確率パラメタ $\mathcal{P}_{RV} := \{P(R_1|V_1), \dots, P(R_N|V_Q)\}$ はコンテンツごとに既知であるとする。

推定すべきパラメタは、各セッションに対する興味アスペクト $\theta(s)$ と、アスペクトと属性値の関係 $\mathcal{P}_{VZ} := \{P(V_1|Z_1), \dots, P(V_Q|Z_K)\}$ である。すなわち、上述パラメタを推定する際には、アスペクト自体も学習データをよく表すものがデータドリブンに獲得される。

いまセッション s における注視領域系列の生成確率が式(5)で表される。このとき、 S 個のセッションにおける注視領域系列集合 $\mathcal{D}$ の生成確率は式(2)(3)より、

$$P(\mathcal{D}) = \prod_{s=1}^S \prod_{t=1}^{T_s} \sum_{q=1}^Q \sum_{k=1}^K P(R_{n_t}^s | V_q) P(V_q | Z_k) P(Z_k | s) \quad (6)$$

となる．本論文では，各パラメタを式(7)に示す対数尤度を最大化するものとして推定する．

$$L = \log P(\mathcal{D}) \quad (7)$$

3.3.1 アスペクトの学習 (図1(a))

得られているセッション集合から，アスペクトを学習する際の入出力は以下になる．

入力 注視領域系列集合 $\mathcal{D}$

出力 アスペクトと属性値の関係 P_{VZ} ，各セッションの興味アスペクト $\theta(s)$

パラメタの学習にはEMアルゴリズムを用いて，対数尤度 L が収束するまで以下の E-Step と M-Step を繰り返すことによりパラメタを推定する．

E-step 現在のパラメタの下で，各時刻に対するアスペクト $Z_k \in \mathcal{Z}$ ，属性値 $V_q \in \mathcal{V}$ の事後確率 $P(Z_k, V_q | s, R_{n_t}^s)$ ($Z_k \in \mathcal{Z}, V_q \in \mathcal{V}, R_{n_t}^s \in \mathcal{R}$) を計算する．

M-step 現在の確率分布の下で，事後確率に関する完全データの対数尤度の期待値を最大化するようなパラメタを選択する．

3.3.2 興味アスペクトの推定 (図1(b))

新規セッション $\hat{s}$ に対する興味アスペクト $\theta(\hat{s})$ を推定する際の入出力は以下になる．

入力 アスペクトと属性値の関係 P_{VZ} ，新規セッションの注視領域系列 $\{R_{n_t}^s\}$

出力 新規セッションの興味アスペクト $\theta(\hat{s})$

パラメタの学習には同様にEMアルゴリズムを用いるが，アスペクトの学習とは異なりここで最大化するパラメタは $\theta(\hat{s})$ のみとなる．

3.3.3 注視領域および興味アイテムの予測 (図1(c))

学習されたアスペクトおよび推定された興味アスペクトより注視領域や興味アイテムの予測は，

入力 アスペクトと属性値の関係 P_{VZ} ，新規セッションの興味アスペクト $\theta(\hat{s})$

出力 新規セッションにおける各領域の注視確率 $\{P(R_n|\hat{s})\}$ および各アイテムへの注目確率 $\{P(I_n|\hat{s})\}$

として，式(2)(3)より

$$P(R_n|\hat{s}) = \sum_{q=1}^Q \sum_{k=1}^K P(R_n|V_q)P(V_q|Z_k)P(Z_k|\hat{s}) \quad (8)$$

および，

$$P(I_n|\hat{s}) = \sum_{q=1}^Q \sum_{k=1}^K P(I_n|V_q)P(V_q|Z_k)P(Z_k|\hat{s}) \quad (9)$$

として与えられる．

4. 実験

提案モデルによって学習されるアスペクトの妥当性および興味アスペクト推定の有効性を評価するため，本研究では実際にカタログコンテンツを利用した実験を行った．実験参加者(計9名)はそれぞれディスプレイ¹の前に着席し，提示されたコンテンツを一定時間閲覧する(図2)．この際の視線運動は，画面下に設置された視線計測装置²によって注視点系列として計測される．

4.1 実験の手順

提案手法によって多様な興味アスペクトが表現できることを確認するため，本研究では実験参加者に複数のタスクを与え，様々な視点からのコンテンツ閲覧を促した．また，タスクを与えることで，各セッション中の実験参加者の興味が一定に保たれると想定した．具体的には，以下の3つのタスクを課した．

- ・ 健康に良さそうな順に3つの料理を選ぶ
- ・ 空腹時に食べたい順に3つの料理を選ぶ
- ・ 自身で調理してみたい順に3つの料理を選ぶ

このように複数アイテムの順序付けをタスクとして課すことにより，より興味のあるアイテム間での比較や吟味が行われることが期待される．

本研究では，提示アイテムへの比較吟味を行っている際の視線運動のモデル構築を目的とする．しかし，ユーザがコンテンツを閲覧する際には，どこにどのようなアイテムが提示されているのか確認する，という情報獲得のフェーズがあると想定される．このような情報獲得フェーズの影響を軽減するために，例えばDodaneら^[5]は，まず提示画像を1枚ずつ順に提示するフェーズを設けるというアプローチをとっている．そこで本実験でもこれにならい，料理画像を1枚ずつ順に提示するフェーズを設けた．また，あらかじめコンテンツ中のアイテムをすべて閲覧した状態で実験を行うことにより，アイテムへのボトムアップな視覚的注意が軽減されることも期待される．各画像を順に提示し終わったのち，全ての画像を一斉に再提示し，この時点からの注視点系列を1分間獲得した．そして，各セッションの終わりに，タスクの結果選択された3アイテムをアンケートにより記録した．セッションごとにアンケートを挟むことで，1つ前のセッションにおける提示コンテンツが次のセッションの視線運動に与える影響を軽減できると考える．

上に述べた3つのタスクそれぞれについて，提示コンテンツを変更して3回ずつ計測を行った．また順序

1: DELL 製ディスプレイ，提示領域は 520x325mm，画素数は 1920x1200pixel

2: Tobii X60 & X120 Eye Tracker，頭部の自由稼働域は 30x22x30cm，サンプリングレートは 60Hz．

表 1 各タスクで用いた画像セット。
Table 1 Image sets displayed in each task.

Image set ID	1	2	3	4	5
Task 1	○	○	○		
Task 2	○	○		○	
Task 3	○	○			○



図 4 提示コンテンツの例, 画像は (株) ミツカン (<http://www.mizkan.co.jp/index.html>), および「おいしい」を科学して、レシピにしました。」(サリー (著), サンマーク出版, 撮影: 久保田 健, 2013) より許可を得て掲載。

Fig. 4 Example of displayed contents

効果を軽減するため、タスクを与える順序や、コンテンツの閲覧順序については、実験参加者ごとにランダムに入れ替えた。本実験を通して獲得された注視点系列の合計セッション数は、9 人 × 9 セッション (3 タスク × 3 コンテンツ) = 81 セッションとなる。

4.2 提示コンテンツと解析手法

ディスプレイに提示するコンテンツとして、料理画像が縦 3 列、横 4 列のタイル状にレイアウトされたものを用いた (図 4)。各料理画像はレシピサイト³ およびレシピ本⁴ から許可を得て転用しており、12 枚の画像セットを計 5 セット用意した。5 つの画像セットのうち、2 セットを 3 つ全てのタスクで利用し、残りの 3 セットをそれぞれ 1 つずつのタスクで利用した (表 1 参照)。また画像をディスプレイに提示する際には、あるアイテムに特定のレイアウトの影響が偏って現れることを防ぐため、提示のたび各画像をランダムに配置するようにした。

提案モデルの妥当性の検証に関して、学習されたアスペクトや推定された興味アスペクトは真値が与えられないため、これらを直接評価することは困難である。そのため本研究では、ユーザが興味アスペクトに基づいてアイテムの注視、選択を行うと想定し、学習されたアスペクトおよび推定された興味アスペクトに基づいて、新たに注視されやすい領域、および興味アスペクト

ム (最終的に実験参加者に選択されるアイテム) の予測を行い、その予測性能をもってモデルの妥当性を評価する。

アスペクトは「コンテンツの見方」であり、コンテンツドメイン (たとえば、料理カタログや地図) ごとに学習されるべきものである。さらには、ユーザのコンテンツに対する見方は多様であり、アスペクトの学習の際には、属性値を幅広く用意しておくことが好ましい。本論文では、料理カテゴリ (米、肉料理、野菜料理、魚料理、麺類)、カロリー、レシピの工程数を意味属性として用いた。具体的には画像の転用元であるレシピに記載された情報を利用し、カロリーについては 4 段階に量子化した。一方、アピランス属性としては画像全体の色情報、具体的には、各画像の前景部分から色相、明度を計算し、16 段階にベクトル量子化したものを用いた。

なお、2. 章で述べたように、意味属性とアピランス属性は、扱う上でそれぞれ利点と欠点がある。両属性が予測性能に与える影響を比較評価するため、以下のように、どの属性を扱うかという点から 3 通りの分析を行った。

SA: 意味属性のみを考慮する

AA: アピランス属性のみを考慮する

SA+AA: 両属性の和集合を考慮する

さらに、ある属性値に興味がある際、その属性値を持つアイテム群から 1 つのアイテムが選ばれる確率は等しいとした。すなわち、提示コンテンツ内で属性値 V_q を持つアイテムの個数 n_q を用いて、

$$P(I_n|V_q) = \begin{cases} \frac{1}{n_q} & (I_n \text{ has } V_q) \\ 0 & (\text{otherwise}) \end{cases} \quad (10)$$

とした。レイアウトの影響を与えるパラメタ ϕ については、タスクやコンテンツ中のアイテムに依存しないことを仮定し、3.3 節で述べたように、実験を通して獲得された全セッションの注視領域系列から学習した。具体的には各領域の相対注視頻度

$$\frac{\sum_{s=1}^S n(s, R_n)}{\sum_{n=1}^N \sum_{s=1}^S n(s, R_n)} \quad (11)$$

を計算し、これを $P(R_n; \phi)$ として用いた。

アスペクト数 K に関しては $K = 3, 6, 9, 12$ の 4 通りについてそれぞれ解析を行い、アスペクト数の変化がどのような結果の違いをもたらすかを検証した。2. 章で述べたとおり、本論文ではセッション中における興味アスペクトの変化は起こらないものと仮定しており、興味アスペクトの推定は、同一セッション内のデータを用いて行った。この際のアスペクトの学習について

3: (株) ミツカン (<http://www.mizkan.co.jp/index.html>)
4: 「おいしい」を科学して、レシピにしました。」(サリー (著), サンマーク出版, 撮影: 久保田 健, 2013)

は、全 81 セッションの中から 1 セッションを評価用とし、残りのセッションを学習に用いた。

予測精度の算出方法としては、まず学習用セッションを用いて学習されたアスペクトに基づき、評価用のセッションの前半部分より興味アスペクトを推定した。ここで、興味アスペクトを推定するために用いるセッションの割合を変化させることで観測データの量による予測精度の変化を分析した。次に、推定された興味アスペクトと学習されたアスペクトから、評価用セッションにおいて注視されやすい領域、興味アイテムの予測を行った。セッション毎に実験参加者にアンケートをとり、実際に選択されたアイテムと、評価用セッション後半によく注視された領域とを真値とし、情報検索などの分野で用いる適合率を評価の指標とした。ここで、興味アイテムの予測において上位 3 つを考慮していることをふまえて、注視されやすい領域の予測についても、注視頻度の高い順に上位 3 つを正答とした。またこれら 3 つの正答は、順序を考慮せず同等に扱った。

4.3 比較手法

比較手法としては、同様に興味アイテムの予測を行う手法として、協調フィルタリングの基本的手法である最近傍推薦を用いる。この評価値としては、提案手法と揃えるため、各属性値にどれだけ注目しているかを利用する。具体的にはまず、セッション毎に獲得された注視点系列から、各領域の注視頻度 $n(s, R_n)$ を算出する。次に注視アイテム頻度 $n(s, I_n)$ を求めるが、本論文では R_n と I_n が一対一に対応するため、 $n(s, R_n) = n(s, I_n)$ となる。ここで各アイテムは、各属性値を持つかを 0 と 1 の 2 値であらわす属性値ベクトル V_{I_n} ($V_{I_n} \in \{0, 1\}^Q$) によって表される。つまり、あるアイテムが属性値 $v_q \in \mathcal{V}$ を持つ場合、そのアイテムの属性値ベクトルの q 番目の要素を 1 とする。この属性値ベクトル V_{I_n} と、注視アイテム頻度 $n(s, I_n)$ を用いて、興味属性値ベクトルを

$$V_{\text{comp}}(s) = \sum_{n=1}^N n(s, I_n) V_{I_n} \quad (12)$$

として求める。

あらかじめ学習用の各セッションにおいて同様に算出された興味属性値ベクトルのうち、興味属性値ベクトル $V_{\text{comp}}(s)$ とのコサイン距離が最も小さい興味属性値ベクトル V_{near} を選び、これを推定興味属性値ベクトルとする。そして、興味アイテム予測として推定興味属性値ベクトル V_{near} と最近傍の属性値ベクトル V_{I_n} を持つ 3 アイテムを選ぶ。さらに、注視領域予測としては、予測された興味アイテムと対応する領域を注視されやすい領域として予測する。

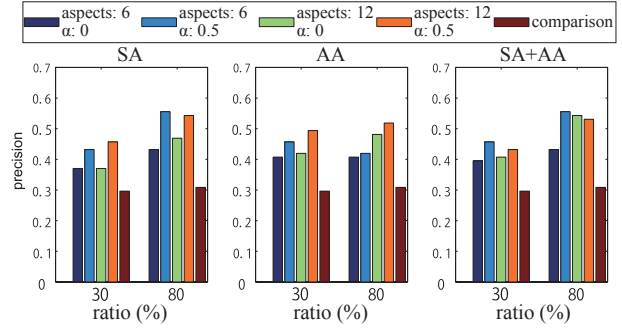


図 5 注視領域の予測精度。

Fig. 5 Prediction accuracies of gaze targets.

この比較手法は、提案手法と同様にアイテムの持つ属性値を考慮しているが、提案手法がモデルベースであるのに対し、比較手法はメモリベースとなっている。両者を比較することで、モデル化を行うことの有効性を評価できると考える。さらには、モデルにレイアウトの影響を取り入れるにあたって、レイアウトを考慮した場合と考慮しない場合との比較を行い、その性能の違いを評価する。つまり、式 (4) で $\alpha > 0$ の場合と $\alpha = 0$ とした場合とを比較する。

4.4 結果

4.4.1 注視領域予測

2 通りのアスペクト数 ($K = 6, 12$)、2 通りの混合比 ($\alpha = 0, 0.5$) に対して、注視されやすい領域を予測した結果を示す (図 5)。このとき、4.2 節で述べたように、考慮する属の種類によって、SA, AA, SA+AA の 3 通りの予測精度を算出した。グラフの横軸の値 x は、興味アスペクトの推定を行うにあたっての評価用セッションの使用率 (セッション開始から何 % を推定に用いたか) を示している。縦軸は、最も注視されやすいと予測された領域に対し、実際の注視確率と比較して得られた予測精度である。全ての条件で、提案手法を用いることで最近傍推薦よりも高精度で注視領域を予測できるという結果が得られた。

図 5 では、アスペクト数が一定の場合、全ての条件において、レイアウトの影響を考慮しない場合 ($\alpha = 0$) よりも考慮した場合 ($\alpha = 0.5$) の方が予測精度が高くなっている。そこで、レイアウトの影響を考慮することがどのように予測の精度に影響するかを調べるため、混合比 α を変えながら予測精度の変化を求めた (図 6)。折れ線の色の違いは、データの使用率 x を表している。閲覧開始すぐにおいては、レイアウトの影響を考慮した方が注視領域の予測精度が向上するが、時間がたつにつれ、レイアウトの影響を考慮しすぎると逆に予測精度が下がってしまうといった結果が得られた。

視線運動解析による興味アスペクトの推定

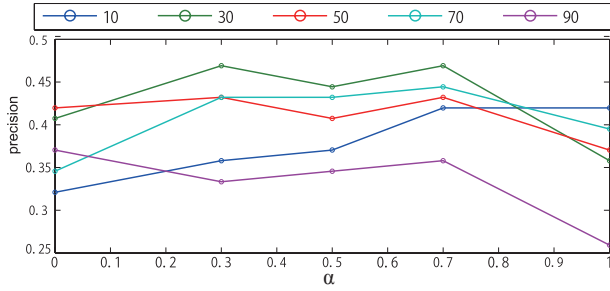


図6 予測精度と混合比 α の関係 (アスペクト数: 9, 属性: SA). 凡例の値は興味アスペクト推定時のセッション使用率 (%) を示す.

Fig. 6 Relationship between α and accuracies of predicting gaze targets.

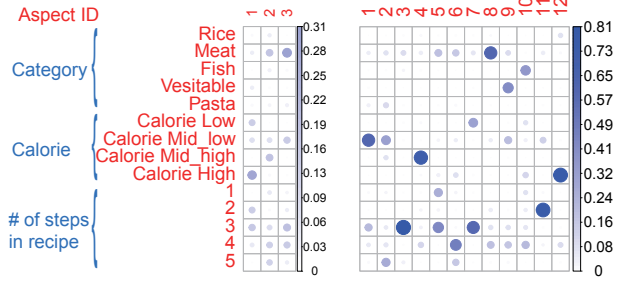


図8 学習されたアスペクトの例 ($\alpha = 0.5$, 属性: SA). 左がアスペクト数を 3 とした時の結果. 右がアスペクト数を 12 とした時の結果.

Fig. 8 Examples of learned aspects.

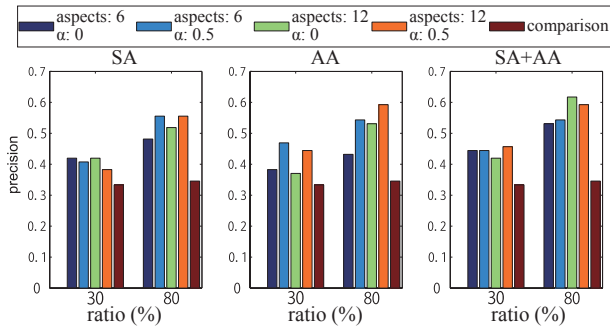


図7 興味アイテムの予測精度.

Fig. 7 Estimation accuracies of items of interest.

4.4.2 興味アイテム予測

同様に2通りのアスペクト数 ($K = 6, 12$), および2通りの混合比 ($\alpha = 0, 0.5$) に対して全セッションから興味アイテムの予測を行った結果を示す (図7). 興味アイテムの予測においても, 全ての条件で, 提案手法を用いることにより最近傍推薦を用いるよりも高精度の予測が行えることがわかった.

4.4.3 アスペクト

全セッションを用いて学習されたアスペクトを, 図8に示す. これは, 各列ごとに1つのアスペクトと属性値の関係, $P(V_q|Z_k)$ を示しており, 値が高い属性値ほど, そのアスペクトの関連が強いということになる. ここでは, 属性値と関連付けられたアスペクトの持つ意味の解釈が行いやすいということから, 意味属性のみを扱っている. また混合比 α は 0.5 とした. また, アスペクト数を 12 とした場合のアスペクトを用いて, コンテンツ1を閲覧していた際の各セッションの興味アスペクト $\theta(s)$ は図9のように推定された.

4.5 考察

本実験のように学習データがあまり多くない場合においても注視領域, および興味アイテムの予測が可能なることから, ユーザの行動をモデル化することの有効性が示された (図5, 7). さらにレイアウトの影響を

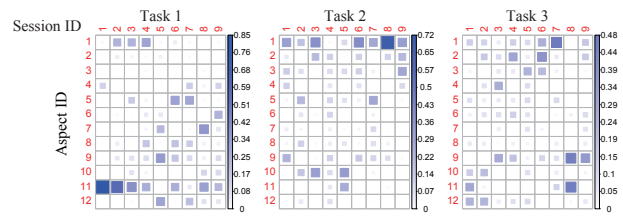


図9 コンテンツ1閲覧時の各セッションに対して推定された興味アスペクト $\theta(s)$ の例 ($\alpha = 0.5$, 属性: SA).

Fig. 9 Examples of estimated aspects of interest

考慮することが, 注視領域の予測精度を向上させることが示された (図5). 一方で, レイアウトの影響を考慮しすぎると逆に予測精度が下がってしまうという結果が得られており (図6), これは, 予測精度を向上させるためには適切な混合比 α の選択が必要となることを示している. また, 閲覧開始すぐ ($x = 10\%$) においてはレイアウトの影響を考慮することで推定精度が向上するのに対し, 閲覧開始からしばらくすると ($x = 90\%$), レイアウトの影響を考慮することで推定精度が低下する傾向にある. これは, セッションの前半部においては事前知識を利用したものの方が予測を行いやすいためだと考えられる. すなわち, この混合比 α は, 時間の経過に伴い, 動的に変化させる必要があると示唆される.

図5, 7より, 扱う属性による予測精度はあまり変化しないということがわかった. つまり, アノテーションが必要となる意味属性を使わずとも, 自動的に属性値を獲得できるアピランス属性のみで十分な予測が行えることを示している. ただし, 学習されたアスペクトの解釈を行いたい場合などには, 意味属性を用いることが有効であると考えられる. しかしながら, 本研究では料理のみをアイテムとして扱っていたため, 異なるジャンルのアイテムが混在している場合についても同様の結果が得られるかについては, 検討の余地が

ある。

学習されたアスペクトの解釈について、アスペクト数を少なくした場合 ($K = 3$) には、各アスペクトは多くの属性値と紐付いてしまい、各アスペクトの意味が解釈しづらくなっている (図 8 左)。これに対して、アスペクト数を多くした場合 ($K = 12$) には、各アスペクトと紐づけられる属性値が比較的疎となり解釈しやすくなる。今回得られた結果において、例えば、6 番目のアスペクトは工程数が多い「手間がかかる」といった料理の側面を表していると考えられ、9 番目のアスペクトは、野菜が含まれカロリーが低いという「健康に良い」という料理の側面を表していると解釈される (図 8 右)。

このように、アスペクト数を多く設定した方が解釈しやすい結果が得られるが、予測精度という観点からは必ずしもアスペクト数を多くすることで精度がよくなるとは限らない。たとえば、図 7 において、意味属性のみを扱い、興味アスペクトの推定に使用したセッション使用率が 30% の結果においてはアスペクト数を 6 とした結果の方が高い予測精度となっている。この理由の 1 つとして、学習データが少ないにも関わらずアスペクト数を多くしたことで過学習が引き起こされたということが考えられる。すなわち、モデルを使用する目的に応じて適切なアスペクト数を選択する必要があると考えられる。

各セッションに対して推定された興味アスペクト $\theta(s)$ は図 9 のようになり、タスクによって現れやすいアスペクトが異なることがわかる。さらに、同一タスクであったとしても、特にタスク 3: 自分で調理してみたいものを選ぶ、などユーザによって着目するアスペクトが多様なものとなる場合も扱えることがわかる。

5. 討議

現時点での提案モデルではいくつかの仮定をおり、これに起因する以下の課題がある。

(i) 注目属性値の扱い方

提案モデルでは、ユーザが各時刻において注目する属性値は 1 つであると仮定している (3.1 節)。しかし実際には、たとえば「赤く小さいもの」といったようにユーザが複数の属性値に注目しているという状況も考えられる。このような状況において、ユーザにアイテムを推薦することを考えると、「赤い」もしくは「小さい」という属性値のうち 1 つを持つアイテムを推薦するのではなく、実際には両者の属性値どちらも持つアイテムを扱える必要がある。このためには、複数の属性値の組をまとめて新たな属性値として扱うことなどが考えられるが、組み合わせによる属性値の数の増加によって学習が難しくなるといった問題がある。

(ii) アスペクトの表現能力の限界

3.1 節で述べたように、学習されたアスペクトは、あらかじめ用意されたアイテムの属性値の組に関連付けられるため、用意されていない属性値に関連するようなアスペクトは考慮することができない。これに対しては、あらかじめ大量の属性値を用意しておき、データをもとに属性値選択を行うといった方法を今後検討する必要がある。

(iii) 時間情報の利用方法

提案モデルはアイテム領域の注視頻度を扱っており、注視の持続時間などを考慮していない (2.2 節)。また、各領域の注視は過去 (直前など) の注視領域とは独立におこると仮定しており、注視された順序などの情報は考慮されていない。しかしながら、従来研究より注視対象の遷移パターン^{[9], [14]}や注視のタイミング構造^[5]がユーザの心的状態を推定する上で有用であることがわかっており、このような時間情報を扱えるようにモデルを拡張することを今後の検討課題としている。これにはたとえば、直前に注視していたアイテムの持つ属性値や、空間的位置の影響 (近い位置に注視が遷移しやすいなど) を利用することが考えられる。

(iv) 情報獲得行動のモデル化

4.1 節で述べたように、本実験ではセッション開始時に各アイテムを順に提示することによって情報獲得行動による影響の軽減を行った。しかしながら実際には、情報獲得に要する時間はユーザやタスクによって変動すると考えられ、上述のような提示法によってアイテムの吟味閲覧フェーズと情報獲得フェーズとが完全に分離できるとは限らない。これに関して、Ishikawa ら^[9]によって、視線運動からユーザの現在の閲覧行動が比較吟味行動であるか情報獲得行動であるかを推定するという手法が提案されている。このような手法を利用して、吟味閲覧や情報獲得のフェーズを推定した上で興味の推定を行うことが考えられる。

(v) モデルの妥当性評価

学習されたアスペクトや推定された興味アスペクトそのものに対する妥当性評価は、真値を獲得することができないため難しい。本論文では、モデルに基づく注視領域と興味アイテムの予測性能を評価することでモデルの評価としたが、学習されたアスペクトの妥当性を直接的に評価することも今後検討する必要がある。これには、たとえば無作為に抽出された属性値の組と学習されたアスペクトに強く関連する属性値の組を、共に外部観察者に提示し、それらとコンテンツドメインとの合致度を評価させる、といったアプローチを今後検討する。

6. 結論

本論文では、コンテンツ提示システムを利用の際の心的状態の1つとして、興味アスペクトを導入し、興味アスペクトに基づいて特定のアイテムが注視されるという注視行動モデルおよびモデルの学習を通じた興味アスペクトの推定法を提案した。実験では、提案するモデルに基づいて注視領域および興味アイテムの予測を行い、その予測性能を評価することでモデルの妥当性を検証した。コンテンツのレイアウトによる影響を考慮することで、予測性能が向上することが確認された。

今後の展開としては、5.章で述べた課題に対応してモデルを拡張し、そのうえで、本研究で得られた知見をもとに、推定されたユーザの興味に基いて適応的かつインタラクティブに情報を提示するシステムへの応用を目指す。

参考文献

- [1] Bednarik, R., Vrzakova, H., Hradis, M.: What do you want to do next : A novel approach for intent prediction in gaze-based interaction; Proc. of ETRA, pp.83–90, (2012).
- [2] Blei, D. M.: Probabilistic topic models; Commun, vol.55, No.4, pp.77–84, (2012).
- [3] Borji, A.: Boosting Bottom-up and Top-down Visual Features for Saliency Estimation; Proc. of CVPR, pp.438–445, (2012).
- [4] Brandherm, B., Prendinger, H., Ishizuka, M.: Interest estimation based on dynamic bayesian networks for visual attentive presentation agents; Proc. of ICMI, pp.346–349, (2007).
- [5] Dodane, J., Hirayama, T., Kawashima, H., Matsuyama, T.: Estimation of User Interest Using Time Delay Features between Proactive Content Presentation and Eye Movements; Proc. of ICACII, pp.201–208, (2009).
- [6] Eivazi, S., Bednarik, R.: Predicting problem-solving behavior and performance levels from visual attention data; Proc. of IUI, pp.9–16, (2011).
- [7] Hirayama, T., Sumi, Y., Kawahara, T., Matsuyama, T.: Info-concierge: Proactive multi-modal interaction through mind probing; Proc. of APSIPA, (2011).
- [8] Hofmann, T.: Probabilistic latent semantic analysis; Proc. of UAI, pp.289–296, (1999).
- [9] Ishikawa, E., Yonetani, R., Kawashima, H., Hirayama, T., Matsuyama, T.: Semantic interpretation of eye movements using designed structures of displayed contents; Proc. of Gaze-In, (2012).
- [10] Jannach, D., Zanker, M., Felfernig, A., Friedrich, G.: Recommender systems: An Introduction; CAMBRIDGE UNIVERSITY PRESS, (2011).
- [11] Jayagopi, D., Sanchez-Cortes, D., Otsuka, K., Yamato, J., Gatica-Perez, D.: Linking speaking and looking behavior patterns with group composition, perception, and performance; Proc. of ICMI, pp.433–440, (2012).
- [12] Judd, T., Ehinger, K., Durand, F., Torralba, A.:

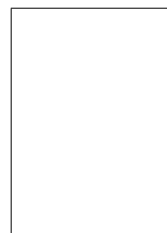
Learning to Predict Where Humans Look; Proc. of ICCV, pp.2106–2113, (2009).

- [13] Li, X., Ji, Q.: Active affective state detection and user assistance with dynamic bayesian networks; IEEE Trans on SMC, pp.93–105, (2005).
- [14] Nakano, Y., Ishii, R.: Estimating user's engagement from eye-gaze behaviors in human-agent conversations; Proc. of IUI, pp.139–148, (2010).
- [15] Pasupa, K., Saunders, C.J., Szedmak, S., Klami, A., Kaski, S., Gunn, S.R.: Learning to rank images from eye movements; Proc. of ICCV Workshops, pp.2009–2016, (2009).
- [16] Puolamäki, K., Ajanki, A., Kaski, S.: Learning to learn implicit queries from gaze patterns; Proc. of ICML, pp.760–767, (2008).
- [17] Qvarfordt, P., Pernilla, Z., Zhai, S.: Conversing with the user based on eye-gaze patterns; Proc. of CHI, pp.221–230, (2005).
- [18] Sugano, Y., Kasai, H., Ogaki, K., Sato, Y.: Image preference estimation from eye movements with a data-driven approach; Proc. of PETMEI, (2013).
- [19] Yoshitaka, A., Wakiyama, K., Hirashima, T.: Recommendation of visual information by gaze-based implicit preference acquisition; Proc. of MMM, pp.126–137, (2007).
- [20] 米谷 竜, 川嶋 宏彰, 加藤 丈和, 松山 隆司: 映像の顕著性変動モデルを用いた視聴者の集中状態推定; 信学論, vol.96, No.8, pp.1675–1687, (2013).
- [21] 中田 篤志, 角 康之, 西田 豊明: 非言語行動の出現パターンによる会話構造抽出 (ヴァーバル・ノンヴァーバル・コミュニケーション, <特集> ヒューマンコミュニケーション～人間中心の情報環境構築のための要素技術～論文); 信学論, Vol.94, No.1, pp.113–123, (2011).

(2002年1月1日受付, 1月1日再受付)

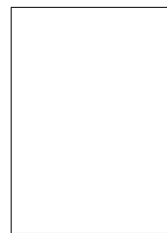
著者紹介

下西 慶



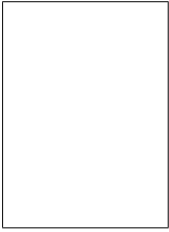
2013年京都大学工学部電気電子工学科卒業。現在、同大学院修士課程在学中。ヒューマンコンピュータインタラクションの分野に興味を持つ。

石川 恵理奈



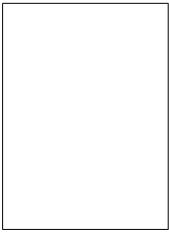
2012年京都大学大学院修士課程修了。現在、同大学院博士課程在学中。2013年より日本学術振興会特別研究員(DC1)。2012年から2013年までカリフォルニア大学サンタバーバラ校客員研究員。ACM, IEEE, 電子情報通信学会学生会員。

米谷 竜



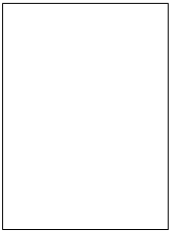
2013 年京都大学大学院博士後期課程修了．博士（情報学）．日本学術振興会特別研究員（DC2）を経て，現在，東京大学生産技術研究所 特別研究員．視覚的注意モデル，視線解析の研究に従事．2010 年 ICPR IBM Best Student Paper Award．2012 年 MIRU 優秀学生論文賞．情報処理学会学生会員．電子情報通信学会学生会員．

川嶋 宏彰（正会員）



2001 年京都大学大学院情報学修士課程了．2007 年より同大学院講師．2010 年から 2012 年までジョージア工科大学客員研究員（日本学術振興会海外特別研究員）．博士（情報学）．時系列パターン認識，ハイブリッドシステム，実世界インタラクションの研究に従事．2004 年 FIT 論文賞．2005 年船井ベストペーパー賞．2007 年 FIT ヤングリサーチャー賞．情報処理学会，IEEE 各会員

松山 隆司（正会員）



1976 年京都大学大学院修士課程修了．京都大学助手，東北大学助教授，岡山大学教授を経て，1995 年より京都大学大学院電子通信工学専攻教授．現在，同大学大学院情報学研究科知能情報学専攻教授．2002 年学術情報メディアセンター長，京都大学評議員，2005 年情報環境機構長．2008 年副理事．工学博士．画像理解，分散協調視覚，3 次元ビデオの研究に従事．最近は「人間と共生する情報システム」，「エネルギーの情報化」の実現に興味を持っている．2009 年文部科学大臣表彰科学技術賞（研究部門）等受賞．国際パターン認識連合，情報処理学会，電子情報通信学会フェロー．日本学術会議連携会員．