

卒業論文

経済情報の波及メカニズムの分析による確率的 グラフィカルモデルを適用した予測

Curriculum Standardization and Learning Analytics for
WebBT in Teaching IR

富山県立大学 工学部 情報システム工学科

2020010 蒲田 涼馬

指導教員 奥原 浩之 教授

提出年月: 令和6年(2024年)2月

目次

図一覧	ii
表一覧	iii
記号一覧	iv
第1章 はじめに	1
§ 1.1 本研究の背景	1
§ 1.2 本研究の目的	2
§ 1.3 本論文の概要	3
第2章 市場間分析を活用した最適なストラテジー構築	4
§ 2.1 市場間の因果性分析	4
§ 2.2 バックテストによる最適なストラテジー	8
§ 2.3 経済情報の波及メカニズムとデータ取得	11
第3章 *****	17
§ 3.1 変数選択とグラフィカル表現	17
§ 3.2 確率的なふるまいのとらえ方, 再現	21
§ 3.3 上がり下がりの判断についての手法	25
第4章 提案手法	30
§ 4.1 経済要因波及と変数選択のシステム開発	30
§ 4.2 各時系列の予測可能性の判定システム開発	35
§ 4.3 提案手法 (予測) のアルゴリズム	38
第5章 数値実験ならびに考察	42
§ 5.1 数値実験の概要	42
§ 5.2 実験結果と考察	45
第6章 おわりに	49
謝辞	50
参考文献	51

図一覧

2.1	e ラーニングとは	5
2.2	e ラーニングの例：progate	5
2.3	ラーニングアナリティクスの概要	10
2.4	レコメンドシステムのイメージ	12
2.5	情報推薦システムの分類	12
2.6	コンテンツベースフィルタリング	13
2.7	協調フィルタリング	13
3.1	レビューと CVR の相関性 [26]	22
3.2	レビュー評価と購買意欲の相関性 [26]	22
3.3	Amazon におけるレビュー例	23
3.4	サクラチェッカー	23
4.1	全 15 回書かれたシラバスの例	31
4.2	授業計画不足シラバス例	31
4.3	シラバス標準化イメージ	34
4.4	教材作成のイメージ	34
4.5	富山県立大学の卒業要件単位	36
4.6	単位区分例	36
4.7	システムのページ遷移	37
4.8	教材更新の流れ	37
4.9	提案システム全体の流れ	39
4.10	蓄積するデータイメージ	40
4.11	新規登録の流れ	40
5.1	新規登録結果	44
5.2	成績入力結果	44
5.3	レビュー時のデータ	45
5.4	もとのシラバス	48
5.5	改善したシラバス	48

表一覧

2.1	eラーニングのメリット・デメリット	6
2.2	教育ビッグデータの例	10
2.3	協調フィルタリングとコンテンツベースフィルタリングの比較 [24]	15
3.1	ある書籍の評価結果	18
3.2	ユーザー×アイテムの評価行列	18
5.1	アンケート項目一覧	43
5.2	アンケート結果	46
5.3	Precision@10	47
5.4	Catalogue Coverage	47
5.5	アンケート結果	49

記号一覧

以下に本論文において用いられる用語と記号の対応表を示す.

用語	記号	用語	記号
特定のユーザー	x	g_c のユーザー数	$size(g_c)$
特定のアイテム	y	g_c の支持度数	$support(g_c)$
利用者数	n	g_c における協調度	$collaborate(g_c)$
アイテム数	m	u_{l_i} が属する頻出投稿者グループ	$G_{u_{l_i}}$
利用者集合 $\{1, \dots, n\}$	$\mathcal{X}$	投稿時間間隔が短いレビュー集合	g_b
アイテム集合 $\{1, \dots, m\}$	$\mathcal{Y}$	レビュー l_i の集中性スコア	$T_score(l_i)$
アイテム y を評価した利用者集合	$\mathcal{X}_y$	g_b のレビュー数	$size(g_b)$
ユーザー x が評価したアイテム集合	$\mathcal{Y}_x$	レビュー l_i と同じジャンルに属するレビュー数	o
対象ユーザー	a	レビュー l_i に出現する名詞集合	K_i
ユーザー x のアイテム y への評価値	r_{xy}	K_i の要素	$term_j$
ユーザー x による評価値の平均	$\bar{r}_x$	l_i と同じジャンルのレビュー集合において $term_j$ を含むレビューの数	$df(term_j)$
アイテム y への評価値の平均	$\bar{r}_y$	レビュー l_i の情報性スコア	$I_score(l_i)$
推薦されたアイテムの数	f	レビュー l_i の類似性スコア	$S_score(l_i)$
レコメンドで推薦されたアイテム集合	$ S_r $	レビュー l_i のサクラ性スコア	$F_score(l_i)$
推薦可能なアイテム集合	$ S_a $	教材 i における信頼性スコア	$K_score(i)$
教材 i につけられたレビュー文章	l_i	教材 i におけるスパムスコアの平均	$\bar{F_score}(i)$
レビュー文章 l_i を <i>bigram</i> によって区切った要素集合	X_{l_i}	教材 i のレビューの評価点の平均	$\bar{l}_i$
頻出投稿者グループ	g_c	レビュー l_i の協調性スコア	$C_score(l_i)$
Precision で考慮する上位ランキングの数	N	Precision で考慮する人数	H

はじめに

§ 1.1 本研究の背景

1996 年の外国為替証拠金取引（Foreign Exchanger: FX）の完全自由化により FX 取引が誕生してから、年々金融市場の規模は拡大している。通信情報技術の発達と金融工学の進歩は、取引単位の小口化と取引手数料の低下により金融市場への参加者を増やし、取引の簡易化と高速化により金融市場全体の流動性を高めた。この 2 点は本来の資産クラスを超えた取引も容易にした。このことは外国為替市場においてもさらなる流動性と市場への参加者をもたらし、元々巨大であった外国為替市場はより一層巨大な市場へと変貌した [1]。通信情報技術の発達がもたらしたのは外国為替市場の規模拡大だけでなくトレーダーにも変化をもたらし、コンピュータが誕生する前や、今ほど性能がない時代はトレーダーの経験や勘といった自身の判断で取引を行う裁量トレードといった取引手法が主であった。コンピュータの性能向上によりそれらを駆使することで自動的にルールに従いトレードを行うといったようなシステムトレードといった物も行われるようになった。また昨今では人口知能を導入することで価格の予測、戦略を獲得するという研究もおこなわれている [2]。市場の予測を身近で行う例として投資があげられる。従来の投資の判断基準として用いられているのが金融市場の要因のみによって得られた分析結果である。そこで用いられる分析は、過去の市場の動きから指標を算出して未来の市場の動向を予測するようなものである。予測に用いられる指標は、直近の市場のデータから一定の値を導出し、その値をもとに今後の市場の値動きを予測することを目的としている。また、指標を単独でもちいて予測を行うことももちろんできるが、それらを複数組み合わせることでより正確な予測を行うことも進められている。また投資を行う上で、このような金融市場のメカニズムを利用して判断することが一般的である。従来の投資の判断基準として用いられているのが金融市場の要因のみによって得られた分析結果である。そこで用いられる分析は、過去の市場の動きから指標を算出して未来の市場の動向を予測するようなものである。市場内の要因を分析するだけでなく他市場からの影響と市場の変動との関連を分析する研究も出てきているように様々な手法を用いて市場の予測を行おうとする研究が行われている。各企業が出している四季報や政治情勢などの市場外的要因も考慮して予測を行う研究もある。市場の変動要因は何であるかについて検証することで市場の変動要因間の相互関係を明らかにするといったような研究も行われている [4]。実際に FX の歴史を振り返ってみると 1985 年のプラザ合意以降に発生した急激な円高進行の局面では、ほぼ並行して原油価格の大幅な低下が進行していた。これらは原油価格低下の影響による好ましくない円高という側面を持つことが知られている [5]。為替リスクとは外国為替相場の変動によって被る損失のこ

とである。一般に、外貨建て資産・負債についてネットベースで資産超ポジションが造成されていた場合、外貨建の取引が予定されていた場合に、為替の価格が当初予定されていた価格と相違することによって損失が発生するリスクと定義される。たとえば貿易業者が輸出をする場合、契約成立から代金回収までには時間がかかるので、この間に為替相場が変動し、為替リスクが生まれる。外為取引を理解する上で大事な部分とは「為替リスク」である。少し具体的に説明する。例えば、日本の企業がもつアメリカの工場がプラント建設をドイツの企業に委託する。このプラント建設は大規模なもので、完成するまで長い時間がかかってしまう。したがって支払いも長期になるケースがある。

§ 1.2 本研究の目的

§ 1.3 本論文の概要

本論文は次のように構成される。

第1章 本研究の背景と目的について説明する。背景では為替リスクについて、目的は確率的な予測および3Dグラフによる可視化を行い、課題に挑むことについて述べる

第2章 金融時系列予測の概要やその利活用について述べる (仮)

第3章 本研究で用いる因果ベイズネットについて述べる (仮)

第4章 本研究の提案手法について述べる。グラフィカル表現や (仮)

第5章 数値実験とその結果、考察についてここでは述べる (仮)

第6章 未定

市場間分析を活用した最適なストラテジー構築

§ 2.1 市場間の因果性分析

eラーニングという言葉が日本国内で聞かれるようになったのは、2000年ごろからである。「e-Japan」という日本型IT社会の実現に向けた構想が打ち出されたことにより、紙などの旧メディアをe化（電子化）していくことに注目が集まり、eラーニングというものが登場した [9], [10].

そして、インターネットのブロードバンド化が大きな転機となり、従来とは比べものにならない高速・大容量通信のインターネット通信ができるようになり、今までのCD-ROM教材を中心とした学習「CBT（Computer Based Training：コンピュータによる教育研修）」からインターネットなどのWebを利用した学習「WBT（Web Based Training：インターネットなどのWeb利用による教育研修）」に移り変わっていった。

インターネットが急速に普及し、それに伴い、企業内のネットワークシステムも整備拡大されるようになった。そして、eラーニングの手法もCBTからWBTに変化していき、eラーニング研修を導入する企業も出始め、国内に浸透していった。

そして、スマートフォンやタブレットといったモバイル端末の登場とともにeラーニングという言葉がより一層の普及を見せるようになった。モバイル端末の普及により、高性能のディスプレイを通して、どこにいても情報にアクセスできるようになった。

また、これらのデバイスを活用したeラーニングは休憩時間や移動時間などの「スキマ時間」を生かした学習にも最適であり、eラーニングの元来的特徴である「いつでも・どこでも学習」をさらに後押しするものとなった。また、直感的な操作性や起動時間の短さ、持ち運びの容易さなどから、学校や塾、企業研修の現場といった幅広い分野での活用が進み、さまざまな利用法や成果が報告されている。

そして、新型コロナウイルスの流行によって、より一層eラーニングが活用されている。外出自粛による勤務形態の変化に伴い、eラーニングを利用する企業が増加している。それに伴い、eラーニング市場も拡大しており、今後もeラーニングの需要・必要性が高まっていくと考えられる [11]。表 2.1 に eラーニングのメリット・デメリットを示す。

eラーニングは同期型と非同期型の2種類に分類することができる [9], [10].

同期型 eラーニング

同期型 eラーニングは同時間帯に eラーニングにアクセスしている学習者の存在がわかるため、臨場感が高く、学習者の孤立を防ぐことや他の学習者の存在によって、学習が動機づけされる可能性が高いと考えられる。また、学習時間も設定されているの

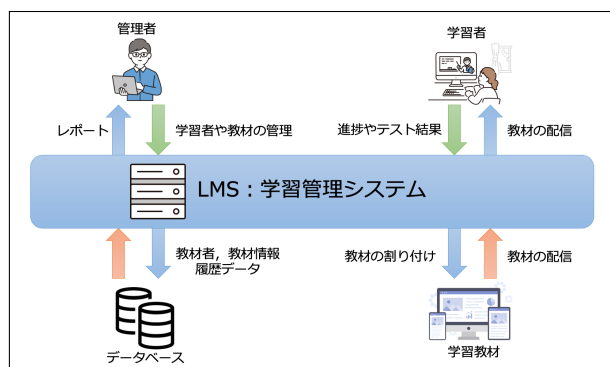


図 2.1: e ラーニングとは

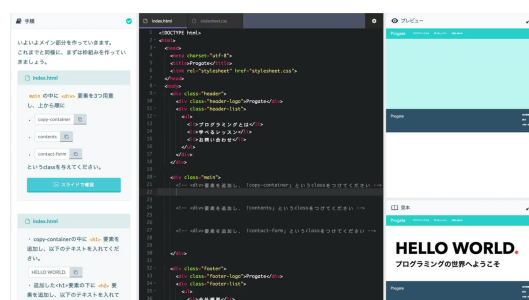


図 2.2: e ラーニングの例：progate

表 2.1: e ラーニングのメリット・デメリット

対象	メリット	デメリット
学習者 受講者	- 自宅や外出先など好きな場所・時間で学習できる - 自分のペースで学習できる - 結果をもとに最適な学習方法が選択され、効果的に習得できる - 学習履歴や進捗が可視化され、わかりやすい - 映像や音声で理解度を深めることができる	- モチベーションの持続が対面と比べ困難 - スポーツなどの実技をともなう学習では、効果的に習得しにくい - インターネット環境とパソコン・スマホなどの端末が必須
提供者 管理者	- 学習者の進捗が一括で管理できる - 教材やプログラムの変更が柔軟に変更することができる - 授業品質の均一化 - 導入以降のコストの削減 - 会場費や移動費などのコスト削減	- 教材を作成する手間やコストがかかる - 教材配信や学習管理のためのシステム（LMS）が必要 - インターネット環境とパソコン・スマホなどの端末が必須 - 管理者または促進者が必要

で、節度が与えられ、学習習慣を確立しやすいと考えられる。

非同期型 e ラーニング

非同期型 e ラーニングはインターネットによって教材を配信し、学習者の都合（時間帯・場所）や能力・スキルに合わせて学習することができる。上記のように、非同期型 e ラーニングは個別の学習者の自由な時間帯に学習が進められるなど、e ラーニング特有の有利な点がある。しかしその反面、学習者が孤立する、学習意欲が湧かない、途中で挫折しやすいなどの欠点が挙げられる。

近年、数多くの大学での ICT 活用教育の普及が進み、学習者は e ラーニング教材を利用する事例が増えている。例として、佐賀大学では 2018 年から e ラーニング科目を導入している。また、広島大学では英語教育に e ラーニングの「ALC NetAcademy NEXT」を導入し、授業の事前課題や授業中の小テスト実施というかたちで利用し、それらを履修した学生を対象に e ラーニングの効果を検証している。その結果として、e ラーニングの利用度合いが高いほど TOEIC のスコアが高くなり、一定の効果をj得ることができている [12]。

しかし、学習者がこのようなシステムを利用するにはコンピュータやネットワーク等を利用する際の基礎的なスキルが十分に身についていないと、学習の序盤から困難に直面し、学習意欲を衰退させてしまう可能性がある。

大学新生に対して、eラーニングを行うにあたっての必要な基礎的スキルや知識をどのくらい持っているか、また、eラーニング導入教育が提示されているかの調査を行った [13]。その結果、大学新生は十分な ICT 活用能力を有していない場合が多く、また、大学が提供する導入教育も不十分な傾向であることがわかったとしている。そして、Web ブラウザ、ネットワークやセキュリティの知識などのスキルを eラーニング利用者は身につけておくことは、学習意欲の面から見ても必要であるとし、導入教育において、その後の継続した情報活用能力である情報フルエンシーを育てていくことを意識した工夫が必要であるとしている。

また、上の例は eラーニング導入についての研究だが、大学における eラーニングは学生の学力をアップし、学力保障につなげるための学び直し（リメディカル教育）や、語学教育、資格取得をメインとした通信教育には欠かせない学習システムである。しかし、eラーニングを続けるには、学生の主体性に依存するところが大きく、学習意欲の向上のための研究も行われている。

上記のように、大学などの高等教育における eラーニングの導入事例、研究やそれらに付随する課題といったものが多く存在する。また、大学教育以外にもさまざまな場面でも eラーニングが導入されている。そこで以下ではそれらの事例について紹介する。

不登校児童生徒に、eラーニングが有効な学習支援システムになるかについての研究を行っている [14]。この研究では「特区事業」において eラーニングを活用しているある市を対象に調査を行っている。その結果、eラーニングには、自分にあった時間帯と場所とレベルから取り組めるメリットがあり、支援員を配し子どもを学習面だけでなく心理面から支えるならば、eラーニングは不登校児童生徒への学習支援システムとして効果的であると判明した。

看護の分野においても eラーニングの導入事例がある。新人看護師の基礎看護技術指導に eラーニングを導入することで、学習者の主体的な学習が促され、看護技術習得率が上昇することを確認できた一方、eラーニング学習で習得した知識や技術を看護技術として確実に習得するための教育体制の構築や、システムやハード面を改善し、eラーニングの機能を充分活用していける働きかけが必要といった課題も残っている [15]。また、看護学生の技術演習で eラーニングを活用した結果、eラーニング教材に視聴平均回数と小テスト結果について、視聴の多い学生に成績がよい傾向があり、質問・アンケートの結果から eラーニングの必要性、関心、意欲が高いことが判明している [16]。

これらは学生らを対象にした eラーニングシステムの研究例だが、大学職員を対象として eラーニングも存在する [17]。佐賀大学では教職員向けの eラーニングとして、適正な研究費管理のためのコンプライアンス教育や、講演会やセミナーの配信などを行っている。学生だけでなく教職員に対しても eラーニングを活用することで、就労時間の異なる職員においてもコンプライアンスなどについて効率よく学習を進めることが可能である。

企業においても eラーニングは数多く導入されており、社会人の新人研修においてしばしば活用されている。企業の新人研修は一般的には3ヶ月間という短い期間で行われる。そのため、eラーニングを導入することで短い研修期間のなかで知識の効率的な定着を図ることが可能になる。企業における eラーニングの活用は、大学等と同様にモチベーションの維持や、受講環境の整備、ロールプレイングなどの実践的な研修の設計といった課題が存在するが、eラーニングを導入することで効率的な学習環境の構築が可能になる。

以上の事例のように、eラーニングが学習効果を高める要因を持っていることがわかる [18].
そこで、以下にそれらの要因をいくつか説明する。

1. 反復学習の最適化

eラーニングによるドリル学習が紙上でのドリル学習と違うところは、ただ決められた順番に問題を提示するのではなく、学習者の解答状況によって最適と思われる問題を選択して提示するアルゴリズムを備えているところである。

一連の問題をどのような順番で提示し、反復学習させていくのが学習者にとって最適なのかという問題は、学習のメカニズムという観点からも、また教育工学的な視点からも重要な研究領域とされている。

2. フィードバックの効果

eラーニングの利点は、個別の学習をしたときに自動的にフィードバックする仕組みを設定できるところである。また、教師が学習者に個人別のフィードバックを与える機能がeラーニングには備えられている。こうしたフィードバックの有無が学習者の意識や成績を向上させる影響があったほか、フィードバックがeラーニングの利用を促進し、そのことが課題や成績に関する意識を高め、それにより自己調整学習方略を高め、さらに自己効力感を高め、成績を高めることができる。

ほかにも、学習内容についてのフィードバック以外にも、講師やメンターからのコミュニケーションとしてのフィードバックも重要である。これらのフィードバックを行うことで、講師やメンターに対して親近感をもち、動機づけを高め、その結果として授業に対する全般的な評価が高くなるという因果関係がある。

3. グループ学習への支援

eラーニングシステムは、グループ学習においても効果的な支援を可能にする。「Re-CoNote」といわれるeラーニングシステムを大学の授業で使用したところ、そのログデータの分析から、相互リンクを利用することで知識構成が促進され、質の高いレポートにつながっている。

以上のように、知識を外化し、それを互いに参照可能にするグループ学習のシステムが学習を効果的に促進するためには、システムの利用だけではなく、質疑や対話についての学習者の姿勢に働きかけることも重要である。

上記の要因や事例が示す通り、eラーニングを導入することで学習効果や学力の向上や表2.1に記してある通りメリットは存在する。その一方でデメリットも存在し、そのデメリットに配慮しながら導入していく必要がある。

§ 2.2 バックテストによる最適なストラテジー

近年、ICTの進歩により、膨大で多様なデータが安価にかつ、容易に取得でき、取り扱うことができる。現在のビッグデータの主な目的は、データを出力する多数の対象の挙動の一般傾向の把握や、予測が中心となっており、ソーシャルメディアの書き込み情報の分析によるトレンド把握などが例として挙げられ、さまざまな試みが実際に行われている。

加えて、子どもから高齢者までのほとんど全ての方が、自分自身のスマートフォンやPCなどの端末を使用するといったICTの活用が広く普及しており、より多くのビッグデータが生み出されている。多くの人間がそれらの端末を使用することで、サイトの閲覧履歴や検索内容といったデータが蓄積されていき、それらのデータを用いたレコメンドシステムや、パーソナライズされた広告の提供といったビッグデータの活用の仕方も行われている。

このようにさまざまな分野においてビッグデータの利活用が進められている。そのような背景のもと「教育ビッグデータ」という言葉の出現が示すように、教育の分野においてビッグデータを用いた問題解決は注目されているトピックの1つとなっている。

教育ビッグデータ

初等・中等教育における電子黒板やタブレットを活用した授業や、大学において1人1台のPCを持参し講義に参加する形式の講義といったe-Learningの広がりにより、教材の閲覧、メモ、成績、学生間の交流など多様かつ大量の学習ログが蓄積されている。また、学習ログのみに限らず、学生のICT機器の操作履歴（どのページを参照したかなど）や学生へのアンケート、学生の履修状況、講義の学習過程の記録、入試の成績、自己評価や他者からの評価など多岐に渡り、これらのデータを教育ビッグデータと呼ぶ。教育ビッグデータの例を表2.2に示す。

教育ビッグデータを通常のビッグデータ（購買行動・公共交通機関などのインフラから生じる大量のログ、ゲノム解析の結果など）と比較すると、以下のような特徴がある [19]。

1. データ量は大きくない

一つの大学に限定すれば、バイト数もそれほど大きくない。通常のビッグデータのデータベースと比較するとデータ量の規模が違う。

2. 対象人数は少ないが、データの種類の急速に増えている

国内最大規模の大学であっても、在籍学生数は10万人以下である。一方、個々の学生に関して収集されるデータは、成績だけではなく、講義の出欠・図書館の利用・eラーニングコンテンツへのアクセスなども含めて、日々急速に増加している。

3. 匿名性が低い

教育ビッグデータとして扱われる学籍番号や履修履歴、学業成績などのデータは個人情報として捉えることが妥当である。そのため、扱う際には匿名化や秘匿性を高める必要がある。また、教育ビッグデータは閲覧する人によっては、簡単に特定できる要素を多く含んでいるため、取り扱いに慎重さが要求される。

4. データの意味解釈が容易

クレジットカードである商品を買ったとしても、本人が使用するのか、いつ使用するのか、なぜそのカードを使用したのかは不明であるが、ある講義に出席したり、オンラインで特定のコンテンツにアクセスしたりする活動の意味は明白である。そのため、学生がどのような行動を、どのような考えを持って起こしているかといった情報を推測しやすい。

5. 因果関係は複雑

個々のデータの意味がわかっていてもデータ同士の関係は不明であることが多い。そのため、それらのデータを分析しても、相関程度の情報しか取得できないケースがある。

6. 多様化・細分化が進行している

教育のICT化が発展するに伴って、個人レベルのマイクロデータが集積している。そのため、データを取り扱う際には取捨選択が重要となってくる。

以上の特徴から、教育ビッグデータの分析においてデータの膨大な量によるデータ処理に要する時間や分散や並列分散処理が問題となるわけではなく、教育ビッグデータを分析したときに、分析結果からどのような関係性、どのような意味を見出す困難さや、さまざまな種類のある教育ビッグデータから意味のあるデータを使用しなければ無駄に多いだけのデータになってしまうといったデータの取捨選択が困難という点が課題となる。

また、3や4に示してある通り教育ビッグデータは扱う人によっては、個人を特定することが可能なデータである。そのためデータへの加工、すなわち「匿名化」処理が必要であるが、匿名化を施したとしても、依然としてプライバシー保護の問題は残るため、これらの個人情報の適切な取り扱いの方策をどうするかといった課題もある。

また、教育におけるビッグデータ・アナリティクスはとくに2000年前後から盛んに行われており、近年、大規模なデータを分析・活用して教育の改善を目指すラーニングアナリティクス（Learning Analytics：学習分析）や教育データマイニング（Educational Data Mining）と呼ばれる研究分野が発展してきている。

ラーニングアナリティクス

ラーニングアナリティクスとは「情報技術を用いて、教員や学習者からどのような情報を獲得して、どのように分析・フィードバックすれば、どのように学習・教育が促進されるかを研究する分野」と定義される [20]。つまり、分析で得られた知見を教授者や学習者にフィードバックし、何らかの意思決定支援を人間が行うことを主たる目的とする。ラーニングアナリティクスとは、さらに、意思決定は人間が行うものであるから、分析結果を分かりやすく可視化することが重要である。図2.3にラーニングアナリティクスの概要を示す。

ラーニングアナリティクスの目的は、個々の授業や科目の改善であることが多く、2011年1月から2016年2月までの、国際会議LAKを含む135件の研究論文の要旨をテキスト分析したところ、予測を目的とするものが多く、特にdropout（落第など）やintervention（介入）を目的とする研究論文が多いことが判明した [21]。

大学におけるラーニングアナリティクスの1つの事例として九州大学の「M2B（みつば）」と呼ばれるシステムがある [22]。このM2Bは「Moodle」といわれる授業やコースにかかわる情報管理やインタラクションの支援を行うことのできるeラーニングプラットフォーム。「Mahara」といわれる長期にわたる学習の記録を保管し共有することのできるeポートフォリオシステム。「BookRoll」といわれるデジタル教科書閲覧システムの3つのシステムを連携させたシステムである。M2Bを使用することで、学生だけでなく教員側に対してもフィードバックを行うことができるため、大学の全体の教育改善・学習改善に繋がっているといえる。

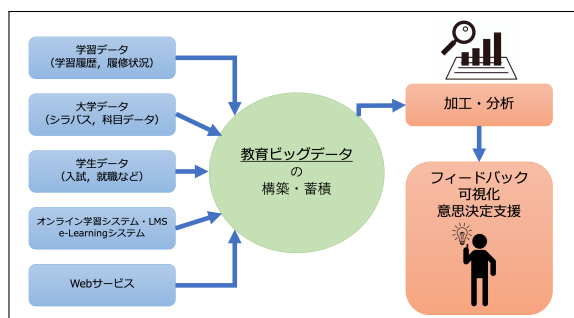


表 2.2: 教育ビッグデータの例

取得時期	教育データ	内容
入学前	出身高校 入試情報 学習ログ 入学前情報	判定値 入試方式、入試成績、志望順位 取り組み状況、学習パターン 志望動機、学習習慣の有無
入学時	学生情報	通学形態、家族構成、通学時間
各セメスター	履修情報 授業 学生生活 成績情報	必修・選択科目、履修科目 出欠状況、遅刻、提出物 サークル、アルバイト、課外活動 科目成績、GPA、テスト結果
4 年次	就職活動	活動履歴、内定状況、志望分野
卒業後	卒業後	満足度（大学、授業）、アンケート

図 2.3: ラーニングアナリティクスの概要

また、ラーニングアナリティクスとは別に教育における「IR」（教学 IR）が発展してきている。教育における IR とは「Institutional Research」の略語で、大学に関するさまざまな情報を収集・蓄積・分析することで現状を詳細に把握し、生徒募集・広報活動、教育活動等の課題を解決することで、学校経営に関する取り組みの改善の推進や意思決定を支援する一連の活動を意味しており、目的としては多くの場合、教育機関全体あるいは学部・学科の意思決定支援であったり、質保証に必要な分析結果の提供であったりする。

教学 IR

IR はアメリカの大学が発祥といわれている。アメリカでは 1960 年代には大学への導入が行われているように歴史が長い。一方、日本において教学 IR という言葉が使われ始めたのは 2010 年代であり、政策誘導によって普及し始めている通り、日本において教学 IR の歴史は浅い。アメリカの大学と日本の大学における教学 IR の目的は少し異なっている。アメリカでの主な目的は学生の獲得から卒業までを維持すること、つまり休退学防止という課題解決のための手法として導入された。一方、日本での主な活用場面は学修成果の改善である。つまり、内部の質保証が主目的であり、いわゆる「教学マネジメント」の文脈で導入・実践されている [23]。

上記のようなラーニングアナリティクスや教学 IR を行い学生が使用するシステムを実装しようとした際に、それらを阻む要因（制度・ルール、実用性など）・課題はいくつか存在する。

1. IR の目的・役割を明確にする

政策誘導によって、近年多くの大学で IR 組織ができているが、目的があって組織化しなければ、担当者が何の業務をすればいいのか手探り状態に陥ったり、他部署から見ても一体 IR は何をしているのか見えず、協力が得られないことが予想される。大学内での課題・問題を明確にし、そのために IR が必要なのか、何をさせるのか、明確な目的と役割を与えることが重要である。

2. 人事制度の問題

IR 部門がどこに組織化され、どのような責任を果たすのか、そのためにどういった能力を持った人員を配置するのかなど、IR の目的に応じて組織過程を作成する必要がある。また、IR を担当するような専門分野を有する専門家を教育機関として雇用

する余裕がない場合、専門家ではない一般の職員に負担がかかってしまう可能性が考えられる。

3. 情報公開の文化が広がっていない

IR 業務を行うには、目的に応じたデータを収集することが必要であり、学内外のどこに必要なデータが、どのような形で保管されているかを正確に把握することが重要である。しかし、情報公開に前向きな大学が少ないため、情報の入手・分析が困難であったり、ラーニングアナリティクスや教学 IR で得られた知見を公開できないことがある。

4. 大学経営の不自由さ

海外の大学等の教育機関とは異なり、日本は大学に対する手当てや寄付、補助が薄いため、資金運用に関する制限が生じてしまう。そのため、新たな知見や分析結果を得られたとしても、それらを活かすような運用が困難になってしまっている。また、資金面だけでなく、大学内での認知も重要となってくる。そのため、こういったことを IR が行い、その結果どういったことに活用できるのか、どういう点で協力が必要なのかなど、大学にいる教職員に説明をし、周知徹底を図ることも重要である。

上記のようにラーニングアナリティクスや教学 IR における問題が存在する。また、上記以外にも学生に対して分析した結果を見せる際のアプローチ方法や学生に対して無理難題な内容で分析結果を表示していないかなど、教育機関における問題は多数存在する。しかし、このような問題・課題が多く存在していても、ラーニングアナリティクスや教学 IR を広めていくことで大学教育や学生の学習環境の向上に貢献できる可能性は大いにあるといえる。

§ 2.3 経済情報の波及メカニズムとデータ取得

近年、ソーシャルメディアの発展やスマートフォンなどの ICT 機器の普及により、多くの人間がインターネットに触れる機会が増えている。中でも YouTube¹をはじめとする動画配信サイト、Amazonをはじめとする EC（Electronic Commerce：電子商取引）サイトなどを利用する機会が多くなっている。その際に、1つの動画、1つの商品を閲覧している際にこれもおすすめですよといったように別の動画や商品をおすすめされることが多くなっている。それらのレコメンド技術を「情報推薦システム」と呼び、現在この情報推薦システムに注目が集まっている。レコメンドシステムのイメージを図 2.4 に示す。

情報推薦システムが必要になった背景は ICT 技術の発展に伴う情勢の中で大きく分けて 2つ存在する。

- ・大量の情報がインターネットを通して発信されるようになったこと。
- ・大量の情報の蓄積や流通が容易になったことにより、誰もが大量の情報を得ることができるようになったこと。

こういった要因により、情報を受け取る側は大量に存在する情報の中から自分の望む情報を取捨選択せざるを得なくなってしまった。それにもかかわらず、日々インターネット

¹<https://www.youtube.com/>

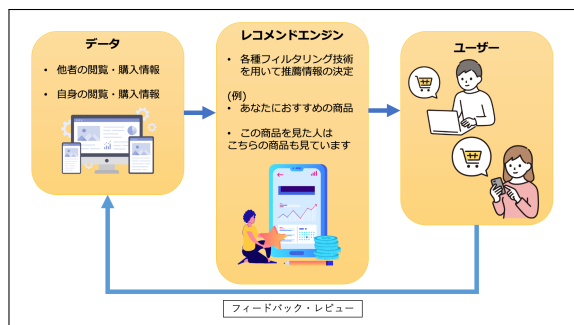


図 2.4: レコメンドシステムのイメージ



図 2.5: 情報推薦システムの分類

にはどこの誰かとも知らない人が発信した玉石混合な情報が更新されている。こうした情報があるにもかかわらず、その情報を利用できない状況を「情報過多」、あるいは「情報爆発」という。

そこで、情報過多といった状況を打破するためや、多種多様な消費者のニーズに応えるための1つの手段として、「情報推薦システム」が多く活用されている。情報推薦システムとは、特定のユーザーに対してアイテムへの嗜好を予測し提示するシステムのことであり、情報推薦システムはAmazonをはじめとするさまざまなECサイト、動画サイトにも導入されており、関連商品の推薦や関連動画の推薦のようなユーザーが閲覧した商品や動画に基づいて、それらと似たものを推薦するようになっている。

推薦システムには永続個人化推薦と一時的個人化、非個人化推薦の3種類に分類することができる [24]。それぞれのイメージを図 2.5 に示す。

永続個人化推薦

たとえ同じ入力や行動をシステムに対して行っている利用者でも、利用者の個人情報や過去の利用履歴、閲覧履歴、購入履歴に応じて異なる推薦をすること。また、ユーザーの趣味嗜好についての情報を蓄積して、それをもとに推薦することである。例えば、過去の購入・利用履歴に基づいて推薦を行ったり、性別や年齢といった個人情報に応じて推薦するアイテムを変えたりすることである。

一時的個人化推薦

システムを利用する1つのセッションで同じ入力や振る舞いをした利用者には、同じ推薦をすること。例えば、利用者がある本を閲覧するという行動をシステムに対して行ったとき、その本に関連する情報を示すといったものがある。(Amazon でいう「この商品を買った人はこんなものも買っています」といったもの)

非個人化推薦

全ての利用者について、全く同じ推薦をすること。例えば、ビュー数の多い記事や、編集者のピックアップによる記事の推薦、売り上げ順位リストなどである。このような、全ユーザを対象(=非個人)として行った推薦(情報の優先度の付け方)を指す。

近年、推薦のパーソナライゼーションという言葉がよく使われている。これは上記にある永続的個人化推薦に該当し、個人の趣味嗜好を考えたうえで個人にあった情報を抜き出し、それらの情報を個人に適した形で提示するといったことである。

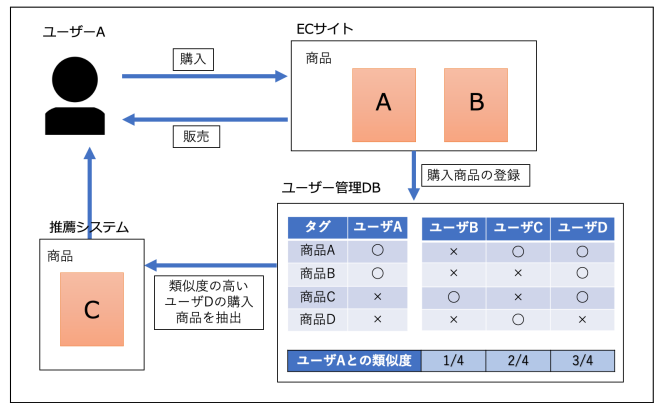
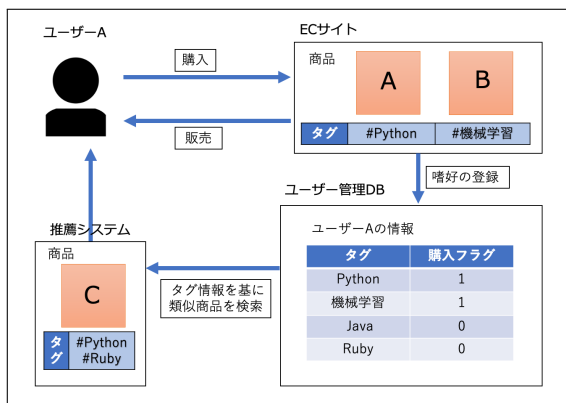


図 2.6: コンテンツベースフィルタリング

図 2.7: 協調フィルタリング

上記に示した通り、情報化社会と呼ばれている現代において必須といっても過言ではないパーソナライゼーションを行うために情報推薦システムでは主に利用者の嗜好を予測し、嗜好に則した情報推薦を行う必要がある。

情報推薦における嗜好の予測方法としてルールベース、コンテンツ（内容）ベースフィルタリング、協調フィルタリング、ハイブリッド法の4つに分けることができる [25]。以下にそれぞれの手法について述べ、本研究で使用する協調フィルタリングについて詳しい説明を述べる。図 2.6、図 2.7 には推薦システムとして広く使われているコンテンツベースフィルタリングと協調フィルタリングの概要を示す。また、表 2.3 はコンテンツベースフィルタリングと協調フィルタリングの比較を示す表である。

ルールベース

一番シンプルなレコメンドアルゴリズム。システムの運営者があらかじめ「このような行動をとった人、このような属性の人にこのような商品や情報を提供する」というルールを設定し、それに従ってレコメンドを行うものである。例えば、クッキーを購入しようとしている人にはコーヒーをレコメンドするといったように、利用者の行動や嗜好を予想して最適だと思われるルールをあらかじめ制定する。

コンテンツ（内容）ベースフィルタリング

コンテンツベースフィルタリングは、アイテム（利用者に推薦する事象や対象）ごとの特徴量をベクトルで表し、利用者の嗜好に近いアイテムの推薦を行うものである。例えば、外食先を推薦するシステムについて考える。この際に利用者は自分の今日食べたいものや自身の住んでいる地域などを検索ボックスに入力し、自分が利用したいと思う飲食店を探す。利用者が「原宿 イタリアン」と検索すると、この「原宿」と「イタリアン」という特徴量についてお店がソートされ、類似度の高い飲食店として利用者にレコメンドする。このように利用者の入力に基づいて情報を推薦するのがコンテンツベースフィルタリングである。

コンテンツベースフィルタリングの中でも上記のような利用者が自分の好むものを直接指定する方法を直接指定コンテンツベースフィルタリングと呼び、利用者の嗜好データから利用者プロフィールを作成し、アイテムデータベースと比較することで利用者の嗜好を測る手法を間接指定コンテンツベースフィルタリングと呼ぶ。

表 2.3: 協調フィルタリングとコンテンツベースフィルタリングの比較 [24]

分類	協調フィルタリング	コンテンツベースフィルタリング
多様性	○	×
ドメイン知識	○	×
スタートアップ問題	×	△
利用者数	×	○
被覆率	×	○
類似アイテム	×	○
少数派の利用者	×	○

協調フィルタリング

協調フィルタリングとは、利用者がシステムを利用する前からあらかじめ利用者の嗜好データをまとめたデータベースを構築、保持し、それらのデータベース内に存在するユーザーの情報をもとに利用者がどのような嗜好パターン（どのようなアイテムを好み、どのようなアイテムを嫌うかといった傾向）にあるのかを分析し、嗜好が似ているユーザー同士は似たようなアイテムを好み、似たようなアイテムを嫌うといった仮定のもと、嗜好パターンが類似しているユーザーを見つけ出し、利用者が好みそうなアイテムを推薦するものである。

協調フィルタリングはメモリベースとモデルベースの2つに大きく分類することができ、そこからメモリベースはユーザーベース協調フィルタリングとアイテムベース協調フィルタリングに分けることができる。以下に協調フィルタリングにおけるメモリベースとモデルベースについて説明し、さらにメモリベースの中のユーザーベース協調フィルタリングとアイテムベース協調フィルタリングについて説明する。

ハイブリッドフィルタリング

コンテンツベースフィルタリングと協調フィルタリングを組み合わせた手法で、データが多い場合により効果的なレコメンドシステムである。コンテンツベースフィルタリングは、推薦がパターン化されやすく、協調フィルタリングは、訪問回数の少ないユーザーに対して機能しにくいといった欠点があるが、ハイブリッドフィルタリングはそれらの欠点が補われ、幅広いユーザーに適切な推薦を行うことができる。

例えば、AさんとBさんの嗜好が似ていて、かつ商品Aと商品Bの属性が似ているとする。この時Aさんが商品Aと商品Bを購入し、Bさんが商品Aを購入していると、Bさんが商品Bを気にいって購入する可能性が高いと考えられるといったような推薦システムである。

メモリベースとモデルベース

メモリベースとは、ユーザーの嗜好についてまとめてあるデータベースとアイテムについてのデータベースをあらかじめ作成しておき、実際にシステムが利用されるときにそれらのデータベースに基づいて、利用者の嗜好パターンを読み取り、嗜好パターンに合わせたアイテムを推薦するといったものである。メモリベースは利用者がシステムを使用するごとにデータベースを参照し、ユーザーの嗜好を定量化しているため、ユーザーの削除やア

アイテムの追加など、データベースの行列の削除・追加された場合でも柔軟に対応できるメリットがある一方、推薦の都度データベースを参照して行っているのでデータのサイズに比例して推薦結果の表示するまでに時間が必要といった計算コストが高くなるといったデメリットもある。

モデルベースとは、利用者がシステムを使用する以前からあらかじめ「Aさんの好むものはBさんも好む」といったような嗜好パターンをモデルとして構築しておく。そしてシステムを利用するときには、データベースではなく、あらかじめ構築しておいたモデルに基づいて情報を推薦するといったものである。モデルベースはシステムが利用される前にモデルを構築するため、モデルの学習が必要なものの登録データの参照は必要なく、メモリベースに比べて推薦処理自体の計算コストを低くすることが期待できるといったメリットがある一方、データベースに変更を加えると1からモデルの構築を行わないといけない必要があるためデータベースの変化に柔軟に対応できないというデメリットもある。

ユーザーベース協調フィルタリングとアイテムベース協調フィルタリング

ユーザーベース協調フィルタリングは、どのユーザーがどのアイテムを購入したかを基に、ユーザー×アイテムの表をデータベースとして持っている。そして新しいユーザーAが商品を購入したとき、どのアイテムを推薦すべきかをAさんと似た購入履歴を持つユーザーを探すことで決める。このようにユーザー間の関係性をもとに推薦するアイテムを決めることがユーザーベース協調フィルタリングである。例として、「あなたに似たユーザーは○○も購入しています」といったものがある。

アイテムベース協調フィルタリングは、どのユーザーがどのアイテムを購入したかを基に、アイテム間の類似性を示すアイテム×アイテムの表をデータベースとして持っている。そして新しいユーザーが商品Xを購入したとき、どのアイテムを推薦すべきかXと関連の高いアイテムを探すことで決定する。このようにアイテム間の関連性をもとに推薦するアイテムを決めることがアイテムベース協調フィルタリングである。例として、「この商品を購入した人は、こんな商品も購入しています」といったものがある。

以上のように、協調フィルタリングには多くの種類のシステムが存在する。そのため、どのような情報推薦システムを作成するかに応じて、どの予測手法を使うかが重要になってくる。

§ 3.1 変数選択とグラフィカル表現

情報推薦の手法として、協調フィルタリングという手法があり、それをアイテムベース協調フィルタリングとユーザーベース協調フィルタリングの2種類が大別することができる。そこで、本項ではその2種類に情報推薦について解説する。まず最初にアイテムベース協調フィルタリングについて解説し、次に本研究で使用するユーザーベース協調フィルタリングについてより詳しく解説する。

アイテムベース協調フィルタリング

アイテムベース協調フィルタリングは、アイテム間の類似性を計算して、特定のアイテムに対して類似したアイテムを推薦するための手法である。この手法は、アイテムに関する評判や評価、アイテム間の類似性などを考慮し、アイテム間の類似性を計算することで実現される。

アイテム間の類似性を計算するためには、さまざまな手法が用いられる。例えば、Jaccard係数は、2つのアイテムが共に評価したユーザーの比率を計算する。コサイン類似度は、2つのアイテムのユーザーの評価ベクトル間のコサイン角を計算する。これらは、アイテム間の類似性を計算するための基本的な手法であり、その他にもさまざまな計算方法が存在する。

アイテムベース協調フィルタリングは、新しいアイテムに対しても推薦が可能であり、特定のユーザーに対して特に適しているアイテムを推薦することも可能である。ただし、アイテム間の類似性を計算するために大量のデータが必要であり、計算に時間がかかってしまうというデメリットが存在する。

例えば、書籍で5段階評価を行っているとする。対象ユーザーは書籍5を読んでいないため評価していないが、どのような評価をするかを予測するといったものを考える。表3.1を用いて説明する。この場合、対象ユーザー以外のユーザA~Dが評価している書籍1~5の評価値から、書籍間の類似度を算出し、類似度行列を作成する。対象ユーザーの書籍1~4の評価値と類似度行列の積から、書籍5の評価値を予測することができ、予測値が大きければ書籍5をレコメンドすることができる。これがアイテムベース協調フィルタリングの考え方である。

また、アイテムベース協調フィルタリングは、Amazonでの推薦アルゴリズムとして使用されており、「あなたと似ている人が買ったアイテム」や「今閲覧しているアイテムと似ているアイテム」、「この商品を買った人はこんな商品も買っています」といったものである。

表 3.1: ある書籍の評価結果

	書籍 1	書籍 2	書籍 3	書籍 4	書籍 5
対象ユーザー	5	3	4	4	??
ユーザーA	3	3	1	5	4
ユーザーB	4	3	4	3	5
ユーザーC	2	5	5	2	1
ユーザーD	3	1	1	3	3

表 3.2: ユーザー×アイテムの評価行列

	国語	数学	化学	物理
ユーザーA	1	3	0	3
ユーザーB	0	1	3	0
ユーザーC	2	1	3	1
ユーザーD	1	3	2	0

ユーザーベース協調フィルタリング

ユーザーベース協調フィルタリング (User Based Collaborative Filtering: UBCF) は、特定のユーザーにおすすめのアイテムを推薦するための手法である。これは、あるユーザーが評価したアイテムをもとに、類似したユーザーの評価履歴を利用して、新しいアイテムを推薦するための方法である。

ユーザーベース協調フィルタリングには、複数の方法があるが、一般的には以下の手順に従って実行される。

Step1：類似度の計算

各ユーザー間の類似度を計算することで、特定のユーザーに類似したユーザーを見つける。この類似度を算出する際にはコサイン類似度、ピアソン相関係数などが用いられる。

Step2：類似ユーザーの選択

Step1 で得られた類似度の高いユーザーを選択することで、特定のユーザーに似たユーザーの評価履歴を参照する。

Step3：推薦アイテムの選択

Step1, Step2 をもとに選択した類似ユーザーの評価履歴をもとに、新しいアイテムを推薦する。

ユーザーベース協調フィルタリングは、類似したユーザーの評価履歴を利用することで、より適切なアイテムを推薦することができ、ユーザーにとって有益な情報を提供できる可能性が高いと考えられる。

そこで、本研究では教学データに対して協調フィルタリング、その中でもユーザーベース協調フィルタリングを適用し、学生がまだ履修していない、または履修できる学年に達していない科目について科目ごとの予測評価値 (予測成績) を算出する。

UBCF ではユーザー×アイテムの評価行列から対象となるユーザーと他のユーザーとの類似度を計算し、対象のユーザーのまだ知らない (評価していない) アイテムについてどのくらいそのアイテムを好むかの嗜好の予測を行う。ユーザー間の類似度は共通して評価しているアイテムについての Pearson 相関で算出し、式 (3.1) で定義する [24]。

$$\rho_{ax} = \frac{\sum_{y \in \mathcal{Y}_{ax}} (r_{ay} - \bar{r}_a')(r_{xy} - \bar{r}_x')}{\sqrt{\sum_{y \in \mathcal{Y}_{ax}} (r_{ay} - \bar{r}_a')^2} \sqrt{\sum_{y \in \mathcal{Y}_{ax}} (r_{xy} - \bar{r}_x')^2}} \quad (3.1)$$

ただし、 $\mathcal{Y}_{ax}$ はユーザー a と x が共通に評価したアイテムの集合、すなわち $\mathcal{Y}_{ax} = \mathcal{Y}_a \cap \mathcal{Y}_x$ であり、 $\bar{r}_x' = \frac{r_{xy}}{|\mathcal{Y}|}$ である。しかしこの時に対象ユーザー a と他のユーザー x が互いに共通して評価したアイテムが1つ以下である場合、Pearson 相関は計算できない。そのためこのような場合は $\rho_{ax} = 0$ として互いの類似度がまったくない状態にして Pearson 相関を算出する。

互いに共通して評価したアイテムが1つ以下ということは、互いにほとんど異なったアイテムについてのみ興味を示しているということなので、ここで互いの類似度を0としてしまってもあまり大きな問題ではない。未評価のアイテムに対する予測評価値は式 (3.1) の類似度で重み付けした対象ではない他のユーザーのアイテム y への評価値の加重平均で予測を行い、式 (3.2) で算出される。

$$\hat{r}_{ay} = \bar{r}_a + \frac{\sum_{x \in \mathcal{X}_y} \rho_{ax}(r_{xy} - \bar{r}_x')}{\sum_{x \in \mathcal{X}_y} |\rho_{ax}|} \quad (3.2)$$

この時、対象ユーザーが既にアイテム y を評価済み ($a \in \mathcal{X}_y$) の状況では、 $\hat{r}_{ay}$ の予測評価値は算出する必要はない。式 (3.2) の第1項は、第2項が中間的な評価で0をとるので、その特徴を補正するためのバイアス項である。また、第2項の分子は上記の加重平均であり、分母は評価しているユーザーが多い、すなわち $|\mathcal{X}_y|$ の集合が大きくなると式 (3.2) の分子が大きくなってしまい加重平均全体が大きくなりやすい問題があるのでそれを補正するための正規化項である。

本研究では、UBCFにおけるユーザーを学生（ユーザー＝学生）、アイテムを学生（アイテム＝学生）が履修することのできる科目、評価値を各科目に対する学生の成績（評価値＝各科目に対する各学生の成績）と置き換え、UBCFを実装する。

UBCFの「同じアイテムに対して高い評価をしているユーザー同士は未評価のアイテムに関してどちらかが良い評価を下していればもう片方は良い評価を下す」という考えを教学データに対して適用して、「同じ科目で高い成績を修めている学生同士は、互いにまだ履修していない科目でも片方が良い成績を修めていれば、もう片方も良い成績を修めることができる」という捉え方をしてUBCFを実装し、学生がまだ履修していない科目について予測成績を算出する。以下では学生の成績データを使った協調フィルタリングの簡単な例を示す。

例えば、表3.2のようなユーザー×アイテムの評価行列があったとする。ユーザーの各科目の評価値は3段階であり3が得意、1が苦手として扱い、0はその科目を未履修のものとする。表3.2ではユーザーAが国語に対しての評価は1である。これは国語の評価が低く、得意ではないことを意味している。

ここで、ユーザーBがまだ履修していない科目: 国語に対する予測評価値 $\hat{r}_{2,1}$ を求める例を示す。まず、式 (3.1) により相関係数を求める。国語を履修済みの学生と対象学生であるユーザーBの相関係数を求める必要がある。ユーザーA、ユーザーC、ユーザーDの3人は国語を履修済みであるので $\mathcal{X}_1 = \{1, 3, 4\}$ の各ユーザーとの相関係数を求める。ユーザーBとユーザーAの相関係数 $\rho_{2,1}$ は共通に履修している科目が数学だけであるので $\rho_{2,1} = 0$ となる。次に、ユーザーBとユーザーCの間の相関係数を計算する。この2人が共通して履

修している科目は数学と化学であるので $\mathcal{Y}_{2,3} = \{2, 3\}$ となり、これらの科目についての $\mathcal{Y}_{2,3}$ の平均評価値はそれぞれ式 (3.3)、式 (3.4) で表すことができ、

$$\bar{r}'_2 = \frac{\sum_{y=2,3} r_{3,y}}{2} = \frac{1+3}{2} = 2 \quad (3.3)$$

$$\bar{r}'_3 = \frac{\sum_{y=2,3} r_{3,y}}{2} = \frac{1+3}{2} = 2 \quad (3.4)$$

そして、相関係数は式 (3.5) で計算される.

$$\begin{aligned} \rho_{2,3} &= \frac{\sum_{y=2,3} (r_{2,y} - \bar{r}'_2)(r_{3,y} - \bar{r}'_3)}{\sqrt{\sum_{y=2,3} (r_{2,y} - \bar{r}'_2)^2} \sqrt{\sum_{y=2,3} (r_{3,y} - \bar{r}'_3)^2}} \\ &= \frac{(1-2)(1-2) + (3-2)(3-2)}{\sqrt{(1-2)^2 + (3-2)^2} \sqrt{(1-2)^2 + (3-2)^2}} \\ &= 1 \end{aligned} \quad (3.5)$$

同様にユーザー B とユーザー D の相関は $\rho_{2,y} = -1$ となる. 次に、予測評価値を計算する. まず、ユーザー B の全ての履修済みの科目の平均評価値を式 (3.6) で計算する.

$$\bar{r}_2 = \left(\sum_{y=2,3} r_{2,y} \right) / 2 = (1+3)/2 = 2 \quad (3.6)$$

最後にこれまでに計算した値を式 (3.2) に代入し、式 (3.7) でユーザー B の国語に対する予測評価値を求める.

$$\begin{aligned} \hat{r}_{2,1} &= \bar{r}_2 + \frac{\sum_{x=1,3,4} \rho_{2,x}(r_{x,1} - \bar{r}'_x)}{\sum_{x=1,3,4} |\rho_{2,x}|} \\ &= 2 + \frac{0(1-3) + 1(2-2) + (-1)(1-5/2)}{|0| + |1| + |-1|} \\ &= 2.75 \end{aligned} \quad (3.7)$$

となり、ユーザー B の国語への予測評価値は 2.75 と計算できる. この値は最大評価値の 3 にかなり近い値となるので、ユーザー B は国語において良い評価値を得られると予測される.

このようにして、学生の各科目における成績データからまだ履修していない科目に対する成績を協調フィルタリングを用いて予測を行う. しかし、協調フィルタリングはその特徴からまだ誰も履修をしたことがない科目についての予測評価値は算出できない. つまり、

同じ学年の学生同士の成績データを使用して協調フィルタリングを行っても、その学年以降に取得できるはずの科目に対して予測成績を算出することはできない。

そこで、本研究では過去の教学データ、つまりは既に大学を卒業した学生の教学データを使用する。現在の学生の教学データを対象とし、過去の学生の教学データを使用して協調フィルタリングを行うことで、いままでの学生が履修してきた科目に対する成績データから、在学中の学生がまだ履修していない科目についても予測評価値を算出することができ、次の学年やその次の学年に取得できる単位については予測成績を算出することが可能になる。

富山県立大学における成績評価方法は成績が良いほうから S, A, B, C , 落単 となっている。これらの成績を GPA の算出時と同じように $S = 4, A = 3, B = 2, C = 1$, 落単 $= 0$ と置き換えて上記した協調フィルタリングを実装する。

§ 3.2 確率的なふるまいのとらえ方, 再現

近年, Amazon や楽天市場¹などの EC サイトにおける市場規模はスマートフォンやタブレット端末の普及などの影響もあり, 年々増加している。また, 新型コロナウイルスの感染拡大の影響により, 在宅需要が高まっている。この現象は EC サイトにとって大きな追い風となり, EC サイトにおける市場規模はこれから先も需要があり, 増加していくと考えられる。

EC サイトはわざわざ店頭に出向かずに購入できる点が非常に便利であり, メリットである。その一方で, EC サイトは実店舗に訪れて実物を見て判断できないというデメリットを抱えている。そこで, ユーザーからの商品へのレビューや評価は, ユーザーの商品に対するイメージをクリアにすることができるので安心して購買に結びつけることができると考えられる。

つまり, レビューや評価はアイテムの信頼性を高め, 購買に結びつけやすいと考えられる。アメリカイリノイ州にある Northwestern University の研究によるとレビューを 5 件以上獲得した場合, 0 件の時と比べてコンバージョンレート (顧客転換率: Web サイト訪問者のうち, 購入や問い合わせなどその Web サイトの最終成果に至った件数の割合) が 270% 増加することがわかっている [26]。図 3.1 はレビュー数と CVR の相関性について示している。

このような理由から, EC サイトにおけるユーザーからの商品へのレビューや評価は他のユーザーへの購買意欲に対して重要な役割を担っている。これらの EC サイトにおける購入者の約 60 % は, 評価値やレビューを含んだ商品サイトから商品を購入する傾向があり, さらに購入者の約 70 % 以上は商品購入前にレビューや評価値を参考にしていると報告されている [27]。

2019 年には食べログといわれるグルメレビューサイトにおいて, 焼肉店が不当に店の評価を下げられたとして, 食べログ運営会社のカクコムに損害賠償を求め, 訴訟を起こしている。訴状としては, 2019 年以降に焼肉店が運営する 21 店舗の評価点が最大で 0.45 点, 平均で 0.2 点下げられ, その結果, 食べログからの来店客数が, 月に 5000 人以上減ったと主張し, 賠償金を求めた [28]。そして, 2022 年において食べログ側は評価点のアルゴリズムを変更を認め, 3840 万円の損害賠償の支払いを命じられた。この訴訟からわかるとおり,

¹<https://www.rakuten.co.jp/>

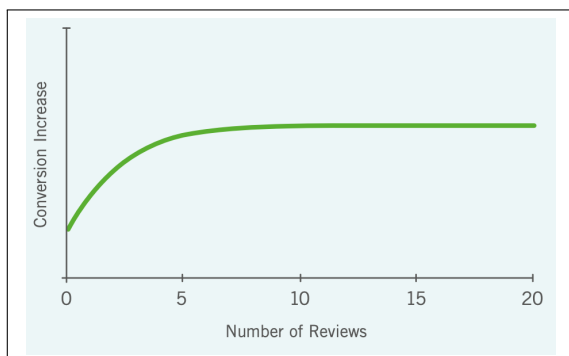


図 3.1: レビューと CVR の相関性 [26]

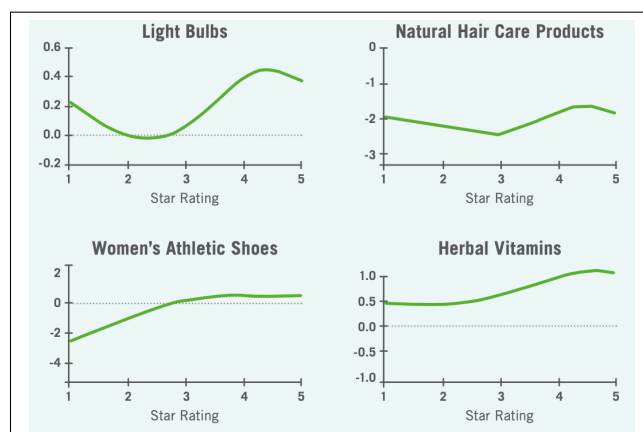


図 3.2: レビュー評価と購買意欲の相関性 [26]

飲食店選びにおける、グルメサイトやレビューの持つ影響力がわかる。また、このような EC サイトにおいて、店や商品に対する評価点が 0.1 下がった場合、客足の低下や、数百万から数億円の利益の低下に繋ってしまうように、レビューには店側に対して非常に大きな影響を及ぼす可能性があるほど重要視される指標となっている。

ユーザーはレビューの内容がポジティブなものであればその商品に対しての購入意欲を高め、レビューの内容がネガティブなものであればその商品に対して購入意欲を下げることになる。しかし、レビューの点数が過剰に高かったりすると、逆にその商品に対して不信感を高めてしまう場合もある [26]。図 3.2 はレビュー評価と購買意欲の相関について表した図である。

このように、商品について投稿されたレビューの内容がその商品の売り上げを直接左右するといっても過言ではない。そのため商品レビューは価値のある情報とされており、マーケティングの観点から非常に注目されており、レビューの評価点や文章に対するデータマイニングなどを用いた研究も盛んに行われている。

しかし、これらの EC サイトにおいてレビューの価値を不当に利用し、利益を目的としたサクラレビューにするステルスマーケティングなどがしばしば行われている。サクラとは、人工的に作られたアカウントやプロフィールを指しており、これらのアカウントは偽のユーザーアイデアや不正な活動を行うときに使用される。また、サクラレビューは上記のサクラアカウントを使用して書かれた商品やサービスに対するレビューを指す。

これらは、商品やサービスの評判を操作するために使用されることがあり、商家や販売者が商品やサービスを宣伝するために書かれたり、競合他社を悪くするためにレビュースパムという偽のレビューの投稿のために書かれることもある。商品購入や店を選択する際のレビュー情報はユーザーにとって大きな情報となることから、これらのサクラによるサクラレビュー問題はしばしば問題として挙げられる。サクラチェッカーといわれるサクラレビューを検出するためのシステムがある²。これは、価格や商品名、レビューについてのそれぞれの項目に応じてサクラかどうか判断するシステムである (図 3.4 参照)。

サクラレビューには、スパムレビューとフェイクレビューの 2 つに分類することができ

²<https://sakura-checker.jp/>

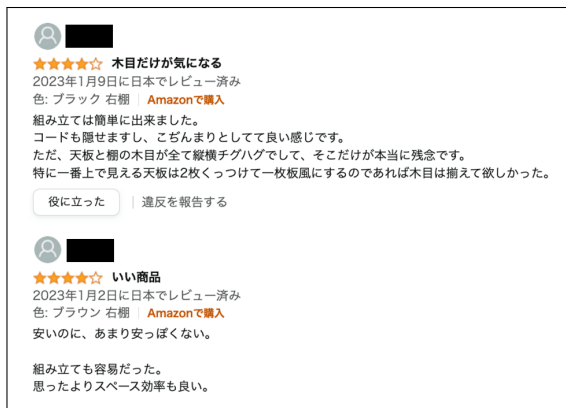


図 3.4: サクラチェッカー

図 3.3: Amazon におけるレビュー例

る [29].

スパムレビュー

スパムレビューとは、対象商品を広告するために意図的に肯定的な評価を与えるレビュー、もしくは対象商品を批判するために不公平または悪意のある否定的な評価を与えるレビューのことである。

フェイクレビュー

フェイクレビューとは、商品を販売している企業が、物品や金銭などを見返りとして、肯定的なレビューを掲載してもらうキャンペーンを持ったレビューである。そのため、フェイクレビューはステルスマーケティングの1種であり、ステルスマーケティングをおこなうためにフェイクレビューグループが存在する。

両者ともそのレビューを読んだユーザーを騙し、ユーザーにとって正確な情報を取得させることを困難にさせたり、誤解を招かせる可能性が考えられる。また、ユーザーに対してだけでなく、商業的にも不正な行為である。そのため、サクラレビューは最悪の場合、ユーザーや店に実害を与える可能性がある。このように、悪質なサクラレビューは利益だけでなく、不当な損害を生じさせる可能性があるものであり、サクラに対する対応は非常に重要な問題といえる。

そのため、これらの問題を解決するためにレビューの信頼性に関する研究はいくつか存在する。しかし、どの研究を見ても、レビューを正確にサクラかどうかを判別するのは困難とされている。そのため、多くの研究ではサクラレビューが持つ特徴を抽出し、複数の観点から評価することで、そのレビューが有用であるか、サクラらしさ（可能性）の算出を行っている。

サクラレビューには、一般的に3種類にタイプに分類することができる [30].

1. Untruthful Opinions

Untruthful Opinions とは、ある商品に対して不当に高い評価を与えて、その商品を宣伝したり、ある商品に対して不当に低いあるいは悪意ある評価を与えることで、その商品の評判を傷つけるといったことを行い、他のユーザーに対して誤解を与えてしまう可能性のあるレビューのことである。

2.Review on Brands Only

Review on Brands Only とは、商品に関するレビューにおいて、商品について具体的なコメントを行うのではなく、商品のブランドやメーカー、販売者についてのコメントをレビューとして投稿することでその製造者や販売元の利益や損失を不当に操作しようとするレビューのことである。

3.Non-Review

Non-Review とは、商品やサービスに対する評価やコメントを含まない、あるいはレビューと称されないレビューのことである。ただの質問やランダムなテキスト、広告がこのタイプに分類することができる。

既存研究として、楽天市場における「みんなのレビュー・口コミ」を対象としたサクラレビューの特徴の分析を行った研究がある [31]。この研究では、サクラレビューに分類されるレビューには「店」「対応」「速い」「メール」「電話」などの商品自体のレビューではなく、店そのものを評価しているレビューが多い傾向があるとしている。逆に、サクラではない一般的なレビューについては「香り」「コンパクト」「軽い」といった五感を通じて感じられる、商品の具体的な評価を行う傾向が強いことがあるとしている。

また、フェイクニュース分類器を作成し、その分類器を通して、フェイクニュースかどうかを判別するといった研究もある [32]。この研究では口コミサイト Yelp に投稿されたレビューを対象にし、分類器を通してそれらがフェイクかどうかを分類している。フェイクレビューとフェイクではないレビューを比較した結果、フェイクレビューの方がレビューを構成する平均の単語の数が少なく、レビューのテキストが短い傾向があるとされている。また、フェイクと判断されたレビューには出現確率が偏っている単語が存在し、「recommend」「friend」といった単語の出現頻度が低いことがわかっている。

上記の3つタイプにおいて、タイプ2とタイプ3についてはその特徴からレビューに付属している文章を解析することで、それらのレビューはサクラであると分類することが可能である。しかし、タイプ1の場合、サクラであるかどうかの判断は非常に困難である。タイプ1に関しては、他のタイプと違い商品に対するレビューを行っているため、普通のレビューと比較しても大きな違いがないため、レビュー文章を解析してもサクラかどうかの判断することは困難である。そのため、サクラレビュー判断に関する研究の多くはタイプ1のようなサクラレビューを検出することが課題である。

しかし、そのレビューが本当にサクラであるのかどうかという問いに対する答えは存在しないため、レビュー文章に何かしらの分析を行い、レビューのサクラらしさを算出してもそのレビュー自体がサクラであるかの答え合わせができない。人間の判断である程度のサクラらしいレビューの答えを作成できたとしても、レビューがサクラかどうかというのは個人の主観であるため100%そのレビューがサクラレビューであるとは言い切れない。

そこで、複製されたレビューに着目してサクラの検出をした研究がある [30]。その研究では、複製、または複製に近いレビューをECサイトから収集し、それらのレビューを手動でサクラとして妥当かどうかの判断を行っている。その結果、サクラとして妥当とされるレビューにはタイプ2,3の特徴が多く含まれていることがわかった。また、それらのレビューの中でタイプ2,3に属さないものは以下の3つの特徴を多く含んでいることがわかった。

- 投稿されたユーザー ID は異なるが、同じ商品に投稿されたレビューの内容が複製もしくは複製に近い本文のレビュー
- 投稿されたユーザー ID は等しいが、異なる商品に投稿されたレビューの内容が複製もしくは複製に近い本文のレビュー
- 異なるユーザー ID で異なる商品に投稿されているが同じあるいは非常に似た本文のレビュー

これら3つの特徴を含み、タイプ2,3に属さないレビューはタイプ1に属する可能性が高い。そのため、複製もしくは複製に近いレビューはその商品を購入したうえで投稿された本物のレビューではなく、サクラグループによって商品の評判を不当に操作するために投稿されたスパムレビューであると判断することが妥当であると結論付いている。

また、サクラはグループでサクラレビューを投稿することで、商品の評価を完全に操作できる可能性があり、そのようなサクラグループを検出するための研究も行われており、サクラかどうかに関しては個々のレビューをサクラかどうか判断するよりも、サクラグループであるかどうかを判断した方がより正確にサクラレビューであるかどうかを判断できるとされている。

他にも、レビューに対してサクラかどうかを判断するのではなく、サクラレビューを投稿しているユーザーに注目し、そのユーザーが他に投稿している商品を見つけることで、各々のレビューにおいて影響を与えている、もしくは受けているユーザーを検出するといった研究も存在する。

§ 3.3 上がり下がりの判断についての手法

科目の推薦を行う際には、ユーザーベース協調フィルタリングを用いた算出された予測評価値をもとに推薦する科目を決定し、それをシステムを使用するユーザーに対して推薦する。しかし、ただ科目のみを推薦するだけでは、学習を支援することはできない。

GPA と授業への取り組み姿勢についての調査を行った研究がある [33]。そこでは、「a. 授業内容について教員に質問する」、「b. 授業中のディスカッションに参加する」、「c. 授業の予習や復習をする」、「d. 授業の発表のために時間をかけて準備する」、「e. 卒業論文の作成を一生懸命頑張った」、「f. 期末テストやレポートの準備もきちんとする」、「g. 板書されていない内容もノートに書き写す」、「h. 授業中に私語をする」、「i. 授業に遅刻や欠席をする」の9つの質問を行っている。その結果、「b. 授業中のディスカッションに参加する」、「c. 授業の予習や復習をする」、「g. 板書されていない内容もノートに書き写す」の3つの要因がGPAにプラスの影響をもたらしていることがわかった。以上の結果から、授業の予習や復習はGPA向上に繋がることがわかる。

そこで、科目を推薦する際に、科目と同時にその科目に応じた教材も同時に提供し、学習の支援を行う。その際に富山県立大学のシラバスに書いてある情報をもとに教材を作成する。シラバスには、開校学期や単位数、授業概要、授業計画、キーワードなどの講義に関する情報が含まれている。そのため、シラバスは学生がその科目を履修する際に参考にするための重要な資料である。

そこで、シラバスの中に書いてある、「授業計画」と「キーワード」を用いて教材を作成する。以上のように、科目を推薦すると同時に、その科目に応じた教材を提供することで、使用するユーザーがそれらの教材を学習することで、学力向上を図っていく。

レビューにおける信頼性については、レビューの特徴からレビュー自体の信頼性を判断することが可能である。これまでに、レビューの信頼性に関するさまざまな研究が行われてきた。これらの研究は、それぞれがサクラの一部についての一面的な特徴のみをとらえたものであり、さらなる改善の必要性が述べられている。また、サクラとこれらの研究との間にはいたちごっこの側面もあり、いくら手法を組み合わせたとしても十分だとはいえない。その結果、最終的に人の主観に頼る必要があると考えられる。

しかし、「レビューの情報源、証拠情報といった信頼性を判断するために役立つ情報が乏しい」、「信頼性判断のために必要な情報を個人が十分に収集するにはコストが大きい」といった理由により、それぞれのレビューが信頼できるか否かを人の手によって判断することは難しい。

そこで、これらのレビューの信頼性の指標として「類似性」、「協調性」、「集中性」、「情報性」の4つの指標を定義し、各指標ごとにスコアを算出し、そのスコアを可視化することで、システム利用者にとってより有効な判断支援を行う研究を行っている [34]。その結果、4つの指標を提示することで、信頼性を判断することが容易になり、有効な判断支援を行うことが可能になったという結果を得られている。つまり、この4つの指標を使用することはユーザーにとって有効な判断支援を行うための材料になり得ることがわかる。これらの指標はそれぞれ0以上5以下の値をとり、この値が5に近いほどサクラらしさが高く、サクラと疑われるものとしている。

そこで本研究でも、レビューの信頼性に関してこれら4つの指標を用いて、サクラ性を算出したあとに、それらの値からサクラ性を考慮した上での教材に対する評価値を算出する。

類似性

複製またはそれに近いレビューには多くのスパムが含まれていることを示している [30]。そこで、他のレビューの文章とどの程度類似しているかを測る指標として「類似性スコア」を定義する。

まずレビュー l_i の文章を bigram で区切る。これは連結する2単語を1つの単位要素として区切る方法であり、bigram によって区切られた単位要素の集合をレビュー l_i を表す要素集合 X_{l_i} とする。次に Jaccard 係数を用いてレビュー l_i と l_j の類似度を式 (3.8) で求める。

$$sim(l_i, l_j) = \frac{|X_{l_i} \cap X_{l_j}|}{|X_{l_i} \cup X_{l_j}|} \quad (3.8)$$

このとき、 $|X_{l_i} \cap X_{l_j}|$ は X_{l_i} と X_{l_j} のどちらのレビューにも登場する bigram の単位であり、 $|X_{l_i} \cup X_{l_j}|$ は X_{l_i} または X_{l_j} の片方に登場する bigram の単位である。そして、レビュー l_i の類似性のスコアを式 (3.9) のように求める。

$$S_score(l_i) = \max_{l_j} (sim(l_i, l_j) | j \neq i, j = 1, 2, \dots, n) \quad (3.9)$$

このとき、 n は l_i と同じジャンルに属するレビューの数である。そして、式 (3.9) を下記の式 (3.10) で正規化を行い類似性スコアを算出する。

$$S_score_{norm}(l_i) = 5 \cdot S_score(l_i) \quad (3.10)$$

協調性

サクラレビューを集団で投稿することで、商品の評判を上げる（または下げる）といった評判を不当に操作するサクラグループは実際に存在する。サクラグループは同じグループのメンバーが同じ商品に対して投稿を行い、商品の評価を協力して不当に操作しようとするものである。そこで、レビューがこのようなサクラグループによって投稿されたものである可能性を測る指標として「協調性スコア」を定義する。

まずサクラグループを発見するために、頻出アイテムセット抽出の方法を用いる。まず、 t_{p_i} をある商品 p_i にレビューを投稿したユーザ ID の集合とし、これらをトランザクションと呼ぶ。また、あるグループと判断されたユーザーが投稿を行った場合、トランザクションの数を増やす。このようにグループに属するユーザーがレビュー投稿を行ったさいに求められるトランザクションの数をそのグループの支持度数と呼ぶ。そして、支持度数が 4 以上でユーザ ID の数が 3 以上となる頻出投稿者グループ g_c を求める。次に、各 g_c の支持度数 ($=support(g_c)$) とユーザ ID 数 ($=size(g_c)$) を用いて、 g_c の協調性を以下のように計算して、以下を式 (3.11) で定義する。

$$collaborate(g_c) = support(g_c) \cdot size(g_c) \quad (3.11)$$

そして、レビュー r_i の協調性スコアを式 (3.12) で求める。

$$C_score(l_i) = \begin{cases} \ln(\max_{g_c \in G_{u_{l_i}}} (collaborate(g_c))) & |G_{u_{l_i}}| \neq \emptyset \\ 0 & |G_{u_{l_i}}| = \emptyset \end{cases} \quad (3.12)$$

このとき u_{l_i} は l_i を投稿した投稿者であり、 $G_{u_{l_i}}$ は u_{l_i} が属する頻出投稿者グループの集合である。さらに、式 (3.12) を下記の式 (3.13) で協調性スコアを正規化する。

$$C_score_{norm}(l_i) = \frac{5 \cdot C_score(l_i)}{\max(C_score(l_i) | j = 1, 2, \dots, N)} \quad (3.13)$$

このとき、 N はすべてのレビューの数である。ただし、投稿履歴が公開されていない投稿者に関してはスコアを求めることができない。

集中性

サクラレビューは時間的に集中して投稿されることを示した [35], [36]。そこで各アイテムのレビューに対して高い（または低い）評価値のレビューがどの程度集中して投稿されているかを図る指標として「集中性スコア」を定義する。評価値とは、各レビュー投稿時にレビューの文章と一緒に投稿されるアイテムに対する 5 段階評価の値である。

どの程度のレビューが集中して投稿されているかを求める方法として、バースト検知手法がある。バースト検知手法は、時系列データに対してある現象の集中的な発生を検出することができる手法である。たとえば、ある時間において Twitter 上で特定の「単語」を含んだツイートの投稿が急激に増えることがある。このようにある現象が急激に増加する現象をバーストと呼び、バースト検知手法はこのような現象を検出する際に用いられる。この手法では、「単語」=「評価値」に置き換えることで、特定の評価値がバーストするタイミングを検知し、そのバーストのタイミングに投稿されているレビューについての集中性スコアと定義する。

バースト検知手法には、「単位時間ごとのイベントの数でふだんより割合が増えているとき」を検知する離散型手法と、「イベント間の時間間隔がふだんより短くなっているとき」を検知する連続型手法がある。この2つの手法を組み合わせることで、ある評価値の投稿数の割合が急激に増加した日を離散型手法で検知した後に、その日の中でのその評価値の投稿数の時間的な変化を連続型手法で計測することが可能になる。以下に、評価値5のレビューの集中性スコアを求める方法を示す。

ある店の m 日目のレビュー集合を B_m とし、時刻の速い順から $\{B_1, B_2, \dots, B_m\}$ と離散時間で送られてくることを考える。このような m 日分のレビュー集合に対して、離散型バースト検知手法を用いて、ふだんよりも評価値5のレビューの割合が増えている日を求める。そのような日を t 日目とし、 t 日目の評価値5のレビュー集合を B_{t_5} とする。

次に B_{t_5} の要素を投稿された時間順に並べた投稿時間列 $l = \{l_1, l_2, \dots, l_{u+1}\}$ を考える。そして、 l_j と l_{j+1} の投稿時間間隔を x_j としたとき、 l の投稿時間間隔列 $x = \{x_1, x_2, \dots, x_u\}$ を求めることで各レビューがどれくらいの時間を空けて投稿されているかを調べる。そして、投稿時間間隔列 x に対して連続型バースト検知手法を用いることで、投稿時間間隔が連続して短いレビュー集合 $g_b \in B_{t_5}$ を求める。各レビュー $l_i \in g_b$ の集中性スコアは、 g_b のレビューの数 $size(g_b)$ を用いて以下の式 (3.14) で求める。

$$T_score(l_i) = In(size(g_b)) \quad (3.14)$$

ただし、このときどのレビュー集合 g_b にも属さないレビューの集中性スコアは0とし、以下の式 (3.15) で正規化を行う。

$$T_score_{norm}(l_i) = \frac{5 \cdot T_score(l_i)}{\max(T_score(l_i) | j = 1, 2, \dots, N)} \quad (3.15)$$

また、他にも MACD で求める（移動平均線収束拡散法）を用いてバーストを検知するで手法も存在する。MACD は短期と中長期の移動平均線の差を求めることで算出でき、MACD と MACD の移動平均線であるシグナルを用いて MACD ヒストグラムを算出できる。MACD、シグナル、ヒストグラムはそれぞれ式 (3.16, 3.17, 3.18) で求めることができる。MACD におけるパラメーター f, s, t は ($f < s$) という条件で設定される。

$$MACD = (\text{時系列値の過去 } f \text{ 期間の移動平均線}) - (\text{時系列値の過去 } s \text{ 期間の移動平均線}) \quad (3.16)$$

$$Signal = MACD \text{ の過去 } t \text{ 期間の移動平均線} \quad (3.17)$$

$$Histogram = MACD - Signal \quad (3.18)$$

情報性

レビューの文章が informative であるほど、そのレビューがサクラではない可能性が高いことを示している [37]。informative であるとは、そのレビューが有益な情報を多く含んでいるということを意味している。また、informative な文章は名詞が多く使用されている傾向があることを示している [38]。つまり、レビュー本文中に名詞が多いとサクラ性を疑われにくいということである。そして、レビュー本文中の名詞が特徴的な名詞、つまり他のレ

ビュー本文で使用されていない名詞であるほどサクラ性が低くなるということでもある。そこで、どの程度 informative なレビューであるかを測る指標として「情報性スコア」を以下の式 (3.16) のように定義する。

$$I_score(l_i) = \ln \left(1 + \sum_{j=1}^{|K_i|} \ln \left(\frac{o}{df(term_j)} \right) \right) \quad (3.19)$$

このとき、 o はレビュー l_i と同じジャンルに属するレビューの数である。また、 l_i に出現する名詞集合を K_i とし、 $term_j \in K_i$ とする。 $df(term_j)$ は l_i と同じジャンルのレビュー集合において $term_j$ を含んだレビューの数とする。

よって、同じジャンルのレビューの中でもあまり他のレビューでは使われていないような特徴的な名詞を多く含んだレビューであればスコアが高くなる。そして、以下の式 (3.17) で正規化を行う。

$$I_score_{norm}(l_i) = 5 \cdot \left(1 - \frac{I_score(l_i)}{\max(I_score(l_i) | j = 1, 2, \dots, o)} \right) \quad (3.20)$$

ここで、 $I_score(l_i)$ 自体はサクラでない可能性を示しており、この値が高いほどサクラである可能性が低いといえる。逆に、 $I_score_{norm}(l_i)$ は他の指標によるスコアと同様にスパムらしさを表しており、この値が高いほどサクラ性が疑われるということである。

提案手法

§ 4.1 経済要因波及と変数選択のシステム開発

本研究では、e ラーニング教材の作成において富山県立大学の Web シラバスを用いて教材を作成する。Web シラバスから情報を抽出し、Web シラバスから得られた情報（授業計画やキーワード）を用いてスクレイピングを行うことで Web 教材を作成する。そして、その Web 教材を学生に提示することで学習支援を行う。

富山県立大学のシラバスには授業科目名、単位数、授業の目標・授業概要、学生の到達目標、授業計画、キーワード等のその授業を履修する際に参考になる情報が多く含まれている（図 4.1 参照）。特に、授業概要や授業計画、キーワードには学生が事前学習する際に必要な情報が含まれており、学生はそれらを参考に学習をする。

しかし、富山県立大学の Web シラバスには、授業計画について情報が不足していたり、全 15 回の授業回数分の量の授業計画が記載されていない科目が多数ある（図 4.2 参照）。また、富山県立大学のシラバスは授業や担当している教員に応じて、さまざまな書き方でシラバスが書かれている。特に授業計画の部分が顕著であり、全 15 回の授業計画を第 1 回、第 2 回と区切るものもあれば、1., 2. という風に区切るものもある。

シラバスとは、「各授業科目の詳細の授業計画。一般に、大学の授業名、担当教員名、講義目的、各回ごとの授業内容、成績評価方法・基準、準備学修等についての具体的な指示、教科書・参考文献、履修条件等が記されており、学生が各授業科目の準備学修等を進めるための基本となるもの。また、学生が講義の履修を決める際の資料になるとともに、教員相互授業内容の調整、学生による授業評価等にも使われる。」として中央教育審議会「学士課程教育構築に向けて」にて定義されている。

つまり、シラバスには授業選択ガイドとしての機能や、学習効果を高める文書、授業の雰囲気伝える文書、授業全体をデザインする文章として役割があると考えられる。そのため、シラバスは学生が授業を理解し、学習するために欠かせないものである。また、シラバスは学生側だけでなく、教員側にとっても重要な役割を持っている。シラバスを作成することで、教員は授業を計画し、授業内容を整理し、学生に適した授業を提供することができるようになる。

そして、シラバスの書き方を統一することも教育機関全体での一貫性を保ち、授業の質を向上させるための重要になってくる。また、書き方を統一することで、学生がこの授業はどういった流れで進んでいくかや、この授業を履修することで何を学べるかといったことが容易に理解しやすくなる。

つまり、シラバスを統一、充実させることは学生に対して、自分の授業を履修したら何

選択した科目に応じたシラバスを記入してもらうようにしてある。その際に、教員名とそれに応じた科目名が書かれている csv ファイルが必要になってくるため、事前に用意しておく。

図 4.3 の中央部分を見てもらうとわかる通り、全 20 回分の項目があるが全ての項目が自由に記入できるわけではない。一部の項目に関しては選択式で選択できるようにしてある。例えば「所属」の部分選択すると教養教育センター、機会システム工学科、知能ロボット工学科、電気電子工学科、情報システム工学科、環境・社会基盤工学科、生物工学科、医薬品工学科の 8 種類が表示され、その中から 1 つ選択してもらうようにしてある。他にも科目区分や単位数、単位区分などは選択式で選択できるようにしておく。

全ての項目を記入してもらい、一番下にある登録ボタンをクリックすると、確認画面に移動し記入した項目が適用されたシラバスが表示される。確認画面には、ダウンロードボタンと登録ボタンがあり、ダウンロードボタンをクリックした場合、前ページで記入してもらった項目が適用されたシラバスを HTML ファイルと css ファイル、javascript ファイル、画像フォルダをまとめた zip ファイルがダウンロードされるようになっている。

また、登録ボタンの場合は記入した項目が適用された HTML ファイルをサーバー上に保存しつつ、最初の名前選択画面に戻るようになっている。また、同時に教材作成時に用いたシラバス情報をまとめてある csv ファイルの該当科目部分を記入してもらった内容で上書きするようにする。このとき、授業計画部分はシラバス表示しやすいように各回の区切りに $\yen$ を追加し改行できるようにし、キーワードの方は 15 回分の記入欄中に間を空けて記入したとしても空白を無くす処理を施しておく。以下ではシラバス標準化のアルゴリズムについて説明する。

Step 1: 教員名の選択

一番始めに表示されるページは教員選択画面である。ここには、情報システム工学科に所属している教員の名前を選択できるようにしてある。この際に教員とその教員の担当科目の科目名が対応している csv ファイルを読み込み、教員名部分のみ抽出し表示する。

Step 2: 科目の選択とシラバス記入

教員は自分の名前を選択すると、次にシラバスを記入してもらうページに遷移する。ここでは、自分の担当している科目を選択式で選択できるようにしてあり、Step 1 でも使用した csv ファイルから科目部分を抽出し選択できるようにしておく。科目を選択してもらい、その他の項目も記入してもらう。その他の項目は全部で 19 種類あり、「授業科目名（英語名）、科目区分、配当学年、職種、所属、開講学期、単位数、単位区分、関連する学習・教育目標、授業の目標/授業概要、学生の到達目標、授業計画（全 15 回分）、キーワード（複数個分）、成績評価基準、教科書・教材参考書等、関連科目・履修条件等、履修上の注意事項や学習上の助言、学生からの質問への対応方法」である。これらの項目は全て富山県立大学の Web シラバスをもとに作成する。

Step 3: シラバス確認と登録

Step 2 でシラバスの内容を全て記入し登録ボタンを選択すると、次に記入した内容が適用されたシラバスが表示される。このページで実際にどのようなシラバスになったかを確認

してもらう。そして、確認出来次第登録ボタンを選択するか、ダウンロードボタンを選択してもらう。

ダウンロードボタンを選択すると、記入した内容を適用した HTML ファイルがダウンロードできる。また、HTML ファイルだけだと css などが適用されていないため、不恰好な状態のままである。そこで、ダウンロードするときにシラバス表示に必要な css ファイルや javascript ファイル、シラバスに埋め込まれている画像フォルダなどと HTML ファイルをまとめて zip ファイルにし、この zip ファイルをダウンロードできるようにする。

登録ボタンを選択すると、記入した内容を適用した HTML ファイルをサーバー上に保存する。それと同時に、もともと Web シラバスをスクレイピングして情報を保存してある csv ファイルの対象科目の部分を記入した内容に書き換える。それらが終了すると、最初の教員名選択の画面に移動する。

Step 4: 記入されたシラバス情報を用いてスクレイピング

Step 3 でシラバスを確認、登録を行うと同時に記入された授業計画を用いて Web ページのスクレイピングを行う。現在のシラバスは、授業全 15 回分の授業計画が書かれていない科目が存在する。そのため、それらの科目は授業計画が 15 回分書かれた科目と比べて教材が少なくなる。シラバス標準化では授業計画を 15 回分記入してもらうため、どの科目も 15 回分の教材が作成でき、この教材は学生用のシステムの教材提供部分で使用する。

ただし、YouTube 動画に関してはスクレイピングを行わない。理由としては YouTube のスクレイピングには API キーが必要となり、この API キーは Google アカウントに対して生成されるものである。本システムは研究室のサーバー上で稼働しているため、自分自身のアカウントに対して生成される API キーを使用した場合、なにかの機会にアカウントを削除した場合 API キーも同時に削除されるため、YouTube のスクレイピングができなくなってしまうからである。

以上の 4 ステップがシラバス標準化の流れである。

次に、教材作成について説明する。上で述べたように、教材を作成する際には Web シラバスから情報を抽出し、そこで得られた情報を用いて、Web 上から教材をスクレイピングする。教材作成のイメージを図 4.4 に示す。

図 4.4 にある通り、まず Web シラバスから各授業の「授業計画」についてスクレイピングを行う。この時、授業計画に 15 回分の計画が書かれていない場合や、15 回分書かれていたとしても同じタイトルで複数回書かれている授業に関しては「キーワード」の方をスクレイピングするようにし、両方とも書かれていない場合は「授業タイトル」をスクレイピングするようにしている。

また、15 回の授業計画が書かれていたとしても、書き方が第 1 回や 1. や①といったような区切り方がバラバラな場合が多いので、その部分は全て取り除き純粋な計画の部分のみ（テキストのみ）取得するようにする。そして、スクレイピングしてきた情報を csv ファイルに出力し、保存する。

次に、保存した csv ファイルを読み込み、全 15 回の授業計画およびキーワードを 1 個ずつ Google Chrome の検索ボックスに入力し、検索結果の上位数件かをスクレイピングする。このときに使用するのが Python のライブラリである BeautifulSoup4 と Selenium である。

BeautifulSoup4 とは、Python 上で実装できるスクレイピング用ライブラリの 1 つである。BeautifulSoup4 に HTML や XML 渡すと、それらのドキュメントをパースし、要素を検索、

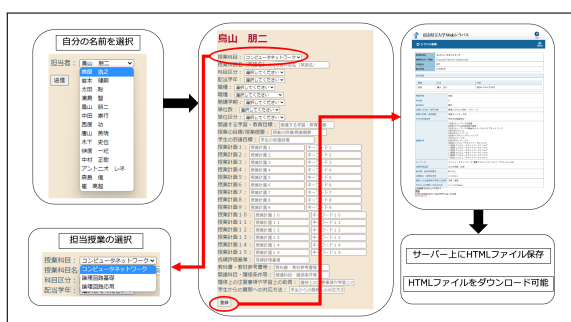


図 4.3: シラバス標準化イメージ

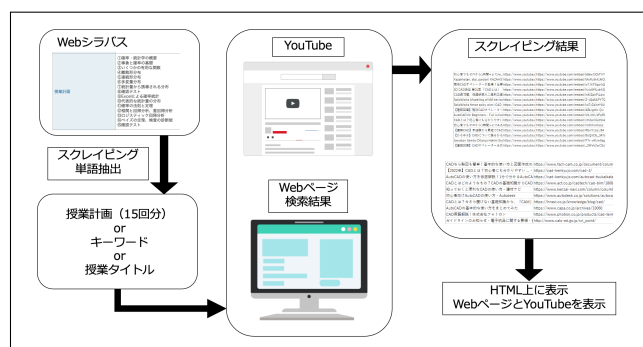


図 4.4: 教材作成のイメージ

操作，抽出することができる．タグ指定やテキスト指定で自分が必要とする情報を取得できる．

Selenium は，Web ブラウザーの自動化を行うためのツールである．複数のプログラミング言語に対応しており，本研究では Python で使用する．Selenium を使用するには Python から WebDriver を使用する必要がある，今回は Google Chrome を使用するため chrome driver を使用する．

BeautifulSopu4 と Selenium は両方とも Google Chrome を用いてスクレイピングを行う．しかし，Google Chrome は bot などによる悪質なアクセスから Web サイトを守る「reCAPTCHA」という機能が存在する．「私はロボットではありません」や「自動車に該当する画像を選択してください」といったものが reCAPTCHA であり，複数回スクレイピングをするとこの reCAPTCHA が表示されてしまう．この reCAPTCHA は Python で自動で突破することは困難であり，有料サービスを使用するといった解決手段があるが今回はそれを利用することができないため，reCAPTCHA に遭遇したら一定時間空けて再度スクレイピングする．

そして，各回のスクレイピングが終了すると，各回ごとの関連した Web ページのタイトルと URL が出力されるので，それを教材とし csv ファイルに保存する．本研究では教材として用いるのは，Web ページだけではなく，動画も教材としてユーザーの提供する．そのため，Web ページだけでなく YouTube から同様な処理を行い，動画情報についてスクレイピングする．ただし，YouTube からのスクレイピングは特定の場合を除いて禁止されており，勝手にスクレイピングすることはグレーゾーンである可能性がある．

そこで本研究では YouTube Data API v3 を使用してスクレイピングする．この YouTube Data API v3 は，YouTube が提供しており，開発者が YouTube のデータや機能にアクセスできる API である．しかし，無料で取得することができるデータ総数は 1 日ごとに上限が設定されているため，動画に関する情報が複数日かけて取得する．

YouTube Data API v3 にはさまざまな機能があり，YouTube の動画再生数やグッドボタンの数，メタタグ，チャンネル登録者数，コメントの取得が可能である．本研究では，そのさまざまな機能の中でも動画タイトル，動画 URL，HTML 上に動画を埋め込む際に必要な HTML コードの 3 つの情報をスクレイピングで取得する．

YouTube でのスクレイピングも Web ページと同様に上位数件を取得するが，YouTube にはスクレイピングによって取得できない動画が存在する．これは，動画の作成者側が他の

HTMLに動画の埋め込みをされたくないといったことが理由であり、これを無理に取得しようとした場合は、著作権の侵害になるので取得をしない。そのため、上位数件をスクレイピングするとあるが、場合によってはこのような動画を除いた上位数件をスクレイピングする。

§ 4.2 各時系列の予測可能性の判定システム開発

本研究ではユーザーベース協調フィルタリングを用いてシステムを利用するユーザーの対して科目推薦を行う。協調フィルタリングの結果を全て使用して推薦すると、予測評価値の低いデータも適用されてしまい、ユーザーにとって興味のない科目まで推薦してしまう可能性がある。そのため、予測評価値が高い科目を優先して学生に推薦するようにする。

また、大学には卒業要件単位や必修単位、選択必修単位が存在する。いくら多くの単位を取得していたとしても、卒業要件単位を満たさなかったり、必修科目を履修し単位を取得できていなかったら卒業はできなくなってしまう。富山県立大学の卒業要件単位は図 4.5 に示すとおりである。図 4.5 は卒業の条件単位数であるが、これ以外にも卒業研究を履修するために必要な単位条件や、3 年に上がる際に必要な単位数も存在する。

選択必修科目だからといって履修しなかったら、アウトな場合が存在する。選択必修科目は必ず履修する必要はないが、決められた科目のグループの中から必ず取得しなければいけない単位であり、これを選択必修単位と呼ぶ。図 4.6 を用いて説明する。図 4.6 では電子・情報工学概論と計測工学は必修単位であり、卒業までにおいて必ず取得しておかないといけない単位である。それ以外の線形代数、工業数学 1、工業数学 2、工業数学 3、工業数学 4、確率システム、情報数学は選択必修単位であり、右に書かれている通りこの 7 科目合計 14 単位の中から 10 単位分以上の科目を履修し、単位を取得しなければ卒業ができなくなってしまう。

このように富山県立大学において、卒業するためには卒業要件単位という大きな制約のなかで、必修単位と選択必修単位という制約を満たす必要がある。本研究では協調フィルタリングで算出された予測評価値の高い科目を優先的に推薦する。しかし、上記で述べているように卒業要件単位や必修単位、選択必修単位を満たさなかった場合、推薦したとしても無意味になってしまう。そこで、これらの要素を考慮しつつ予測評価値が高い科目を優先的に推薦するようにする。

図 4.7 に実際に表示される HTML のページ遷移を示す。まず、使用ユーザー自身の履修状況とそれに応じた推薦科目を表示する。この推薦科目の部分は、上記で述べたように卒業要件単位、必修単位、選択必修単位、予測評価値の高い科目を考慮してソートした状態で表示される。

推薦科目の部分を説明する。ユーザーの履修している科目の成績 (S,A,B,C) と教育ビッグデータ (過去の学生データ, デモデータ) を用いて予測評価値が高い科目を算出する。このとき、単純に予測評価値の高い科目から推薦するのではなく、必修単位を優先して推薦する。このとき、必修単位で予測評価値が低い科目があっても、必修単位を取得しなければ卒業できないため、どんなに予測評価値が低くても上位で推薦するようにする。

その次に選択必修単位を満たすように予測評価値の高い科目順にソートする。それ以降の推薦科目は、すでに必修単位と選択必修単位を考慮している状態のため卒業要件単位を

区分		卒業要件単位	
総合科目	人間	2単位以上	教養小計 44単位
	社会・環境	6単位以上	
	言語・文化	4単位以上	
	精神・身体	3単位以上	
	総合科目計	19単位以上	
基礎科目		13単位	専門小計 79単位
外国語科目	英語	10単位	
	第2外国語	2単位	
キャリア形成科目		7単位	専門小計 79単位
専門基礎科目	卒業研究2以外	71単位	
専門共通科目	卒業研究2	8単位	
専門科目	卒業研究2	8単位	
合 計		130単位	

図 4.5: 富山県立大学の卒業要件単位

専 門 基 礎 科 目	線形代数	◇						半	2	10単位以上(※) 修得すること
	工業数学 1	◇						半	2	
	工業数学 2	◇						半	2	
	工業数学 3		◇					半	2	
	工業数学 4		◇					半	2	
	○確率システム	◇						半	2	
	○情報数学	◇						半	2	
	○電子・情報工学概論	◎						半	※2	
	○計測工学		◎					半	2	

図 4.6: 単位区分例

満たすように予測評価値の高い科目を推薦する。また、その中でも1年、2年、3年、4年といった順番でその学年で履修することができる科目を優先して推薦する。

そして、ユーザーが推薦科目のどれかを選択すると、シラバスと授業計画をキーワードとした教材へのリンクを表示する。このシラバスは富山県立大学のWebシラバスを埋め込んでおり、そのシラバスに表示されている授業計画、キーワードを教材とし、それらへの教材提供ページへのリンクを右側に表示する。

教材へのリンクを選択すると、WebページとYouTubeの教材を提示するページへ移動する。このページには教材であるWebページとYouTubeへのリンクが埋め込まれており、リンクをクリックするとそれらのページへ移動し、ユーザーはそのページをもとに学習するといった流れになっている。また、YouTubeの場合はリンクのみだけでなく、YouTubeの動画も埋め込んであるためこのページ上で動画を見ることができる。

これらの教材が表示される順番は誰もレビューを投稿していなかったり、システムが利用されていない状態だとスクレイピングしてきた順番で昇順に表示される。そしてユーザーがレビューを投稿すると、そのレビューに応じて教材の表示される順番が変化する。

以上が本研究で使用するシステムの科目推薦部分の流れである。次に、教材更新について説明する。

上記で説明した通り、教材の表示の順番はスクレイピングしてきた順番で表示される。そのため、教材として適していないWebページやYouTubeの動画が最初から一番上に表示されていたら、そのままその教材が表示され続けてしまい、学生の教材として学習に使えない状態になってしまう。

そこで、教材へのレビュー機能を追加し、学生がレビューを投稿した際にそのレビューをもとに評価が高い教材を表示頻度を高め、逆に評価が低い教材は表示頻度を下げることによって、学生が教材として利用、学習しやすい教材を優先的に表示するように教材の更新を行う。

レビューにはサクラレビューといわれる故意に評価の上げ下げをするレビューが存在する。そこで教材の更新には、レビューの信頼性の指標を用いて、サクラ性のあるレビューを考慮し教材の更新を行う。

レビューの信頼性の指標には「類似性」「協調性」「集中性」「情報性」の4種類が存在する。その4種類の中でも、協調性に関しては本研究では考慮しないものとする。理由としては、本システムを利用するユーザーは学生を想定しているため、学生がこのシステムに



図 4.7: システムのページ遷移



図 4.8: 教材更新の流れ

において複数人で協調してサクレビューを投稿することは考えにくいといった理由である。そのため、類似性、集中性、情報性の3つの指標を用いることにする。

それぞれの指標の求め方について説明する。類似性は、レビュー文章を2文字単位で区切る。このようにレビューの文章ごとに2文字ずつに区切った文字列集合を作成し、その文字列集合の要素がどれだけ他の文字列集合の中に含まれているかを式(3.8)で各集合の類似度を、式(3.9)で類似性スコアを求め、式(3.10)で正規化し類似度を求める。

集中性はレビューの5段階評価において1または5が極端に増加したタイミング(バースト)を求める。本研究ではMACDを用いてバーストを検知する。MACDにおけるパラメーター $f, s, t(f < s)$ はそれぞれ4, 8, 5でMACDの計算を行う。MACDでは、1日を1つのまとり(日足)として、MACDヒストグラムを求める。そして全期間のヒストグラムの標準偏差を σ としたとき、ヒストグラムの値が平均値 $+3\sigma$ 以上となる日をバーストした日とする。次にバーストした日について15分足でMACDヒストグラムを求め、ヒストグラムの値が平均値 $+3\sigma$ 以上となる期間を求め、その時間にバーストしたと判断する。バーストした期間に投稿されたレビューの数を求めて、式(3.14)で集中性スコアを求め、式(3.15)で正規化し集中度を求める。

情報性は、まず各レビューに対して形態素解析を行う。形態素解析を行うことでレビュー文章を品詞分けを行い、名詞として判断された単語を抽出する。それらの名詞が他のレビューにおいてどれくらい使用されていないかを表す情報性スコアを式(3.19)で求め、式(3.20)で正規化し情報度を求める。

以上の方法でそれぞれの指標スコアを求める。そして、それぞれを足して3で割ったスコアをサクラ性スコア($=F_score(l_i)$)とする。サクラ性スコアは高ければ高いほどサクラレビューとして疑われるものであり、逆にサクラ性スコアが低いとサクラレビューである可能性は低くなる。そしてサクラ性スコアをその教材に対して投稿されたレビュー評価値(5段階評価)の平均で割った数値をその教材の最終的なスコアとし、信頼性スコア($=K_score(i)$)と呼ぶ。

そしてこの信頼性スコアを用いて、このスコアが高い順番で教材の順番を更新し、上位3件の教材を表示する。この教材の順番はユーザーがレビューを投稿するたびに更新されるため、常に最新のスコアが反映された状態でユーザーに表示することができる(図4.8参照)。

以上で説明したのが、科目推薦と教材更新のやり方である。開発したシステムをHTML

上に適用し、システムの作成といったサーバーサイドの部分を「Flask」といわれる Python の Web アプリケーションフレームワークを使用する。

なぜ Flask を使用するかについて説明する。Flask にはあらかじめ実装されている機能が少ないため、自由度が高く動作が軽い。また、柔軟性や拡張性も高く、必要な機能をカスタマイズしたり、別のフレームワークと結合することもできる。他にも、機能が少なく身軽なため、処理速度が高速であり、スムーズな動作を実現できる。

また、最大のメリットとして Python を基盤として作成できる点がある。Python には豊富なライブラリが存在する。(例:Pandas:データ解析ライブラリ, Scipy:数値解析ライブラリ など) 本研究ではスクレイピングを含め、データを処理するさいには Python を用いて処理を行うため、Flask を利用することで HTML から入力されたデータを取得したり、それらのデータを分析し HTML に反映させるといったことが容易に実現できる。

§ 4.3 提案手法 (予測) のアルゴリズム

システム全体の流れを図 4.9 に示す。提案システムのアルゴリズムについて説明する。

Step 1: 教育ビッグデータと教材データの蓄積

ユーザーがシステムを使用する前に、学生 × 科目の行列を複数年分作成しておく。このデータには、富山県立大学を卒業した学生した成績データが入っている。成績データはそれぞれの取得した単位の評価であり、富山県立大学では S, A, B, C, 不可という表し方をしているが、このデータにはデータ処理を行いやすいようにそれぞれの評価を 4, 3, 2, 1, 空白に置き換えたデータが入っている。

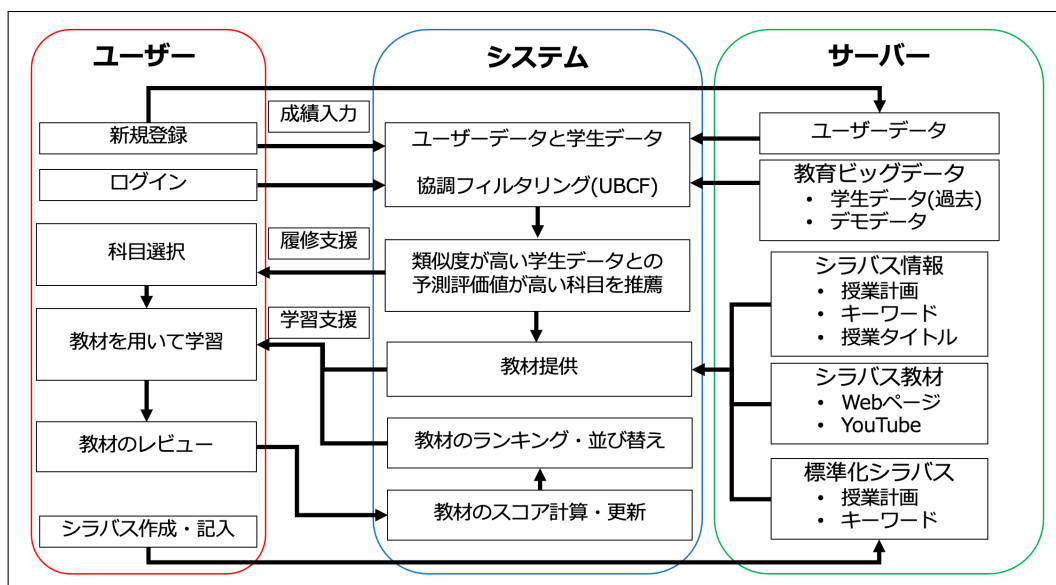
データの中身について図 4.10 を用いて説明する。csv ファイルの 1 行目に科目名を入れてある。今回は富山県立大学の履修の手引きを参考に教養科目、キャリア形成科目、情報システム工学科の専門科目で作成してある。2 行目には必修科目か選択必修科目かどうかを 1 と 0 の 2 値が入っている。3 行目にはどの学期で履修できるのかを表した数が入っている。これは、1 年前期を 10, 1 年後期を 11, 2 年前期を 20, 2 年後期を 21, 3 年前期を 30, 3 年後期を 31, 4 年前期を 40, 4 年後期を 41 といったように表している。4 行目には単位数についての情報が入っている。また、1 列目にはそれぞれの情報名と 1 列 4 行目以降には学籍番号が入っている。そして最後の列に就職先データ (業種) を追加する。

このようなデータを蓄積し、これらを教育ビッグデータとして扱い、サーバー上に保存する。これらのデータは学籍番号や成績といった個人情報が含まれている。本システムの想定しているユーザーは学内の人間であるため、それらの人にとってこのデータは個人を特定できてしまう恐れがある。そこで、本研究では実際の学生データではなくデモデータを作成し、デモデータを用いて協調フィルタリングを行う。

また、学生に提示する教材は科目ごと、授業計画ごとにフォルダを作成し、それぞれそのフォルダに csv ファイルの形で保存しておく。

Step 2: ユーザーの新規登録と成績入力

初めて本システムを利用する場合、新規登録をとして学籍番号とパスワードを登録してもらう。このとき登録されたパスワードはそのまま保存すると、個人情報の流出の恐れがあ



るため、ハッシュ関数を用いてパスワードをハッシュ化しパスワードを不規則な文字列に変換する。そして学籍番号とパスワードをサーバー上の csv ファイルに保存する。この csv ファイルには新規登録されるたびに、学籍番号とハッシュ化したパスワードを新しい行に追加していく。

そして、学籍番号とパスワードを登録し終わったら、次にユーザーの成績を入力してもらう。そこでは情報システム工学科が履修する科目（履修の手引き参照）の成績を入力してもらう。また、成績と同時に希望業種も選択できるようにする。この希望業種の種類は富山県立大学の進路状況のページを参考に製造業、情報通信業、サービス業、建築業、学術研究・専門・技術サービス業、電気・ガス・熱供給・水道業、運輸業・郵便業、金融業・保険業の8種類+希望なしの合計9種類の中から選択できる。

動作としては図 4.11 のような流れで行ってもらう。まずスタート画面として左上の画面が表示される。初めて利用するユーザーは「新規登録はこちら」を選択してもらうと、左下の新規登録用のページに移動する。そこで学籍番号とパスワードを入力し、新規登録ボタンを選択すると、右の成績入力ページに移動する。以上がシステムを初めて利用する場合の一連の流れである。そして、成績登録をすると Step 1 で説明した保存先とは別のユーザー成績専用の csv ファイルに保存する。

Step 3: 協調フィルタリングによる予測評価値の算出

ユーザーにはログイン画面から学籍番号とパスワードを入力してもらいログインするか、新規登録で成績入力を完了したら、履修状況と推薦科目が表示されているページに移動する。ここでは、ユーザーの成績を登録した csv ファイル（Step 2 で作成した）から入力された学籍番号に対応した成績データと希望業種データを抽出し、そのデータと Step 1 で説明したデモデータを用いて協調フィルタリングを行う。

その際、ユーザーの成績データとデモデータには就職に関するデータ（業種データ、希望業種データ）があり、ユーザーの希望業種とマッチしているデモデータの学生を結合す

学籍番号	教養ゼミ I	教養ゼミ II	経済学 I	業種
必修別	1	1	0	
配当開始学年	10	11	10	
単位数	1	1	2	
1715001	3	2	3	2
1715002	3	2	2	2
1715003	2	2	1	4
1715004	4	2		2
1715005	3	2		1

図 4.10: 蓄積するデータイメージ

The screenshot shows a login page titled 'ようこそ' (Welcome). It has a '成績チェック(GET)' section on the right. On the left, there is a '登録済みの方はこちらから' (For those who are already registered) section with a 'Login' button. Below it, there is a '新規登録' (New Registration) button. A red circle highlights the '新規登録はこちら' (New Registration Here) link, and another red circle highlights the '新規登録' button. A red arrow points from the link to the button.

図 4.11: 新規登録の流れ

る．そして，Pearson 関数を用いてユーザーとデモデータの学生の類似度を計算する．類似度が高い順にソートを行い，ユーザーと類似度の高い上位 10 名の過去の学生のデータを抽出する．抽出した 10 名との成績データとユーザーの成績データの加重平均を求めることでそれぞれの科目における予測成績（予測評価値）を導出する．

Step 4: 科目の推薦

Step 3 で求めた予測評価値を用いて使用しているユーザーに対しての科目の推薦を行う．このとき推薦する順番としては，必修単位，選択必修単位，卒業要件単位を満たすように予測評価値の高い科目をユーザーに推薦する．また，その中でも 1 年時科目から優先し 2 年，3 年，4 年時の科目順で推薦する．具体的には 1 年時の必修科目である環境論 1 と選択必修科目である社会学 1 を推薦する場合．必ず上に表示するのは環境論 1 のほうであり，社会学 1 などの選択必修科目に関しては，必修科目の後ろになるようにしてある．

そして，推薦結果とユーザーの履修状況を画面に表示する（図 4.7 の左上参照）．その際に予測評価値の高い科目から上に表示し，予測評価値の低い科目を下に表示する．この科目の推薦は必ず取得する必要はなく，あくまでも履修する際の参考にする程度にしてもらい，ユーザーの履修選択の際の履修支援といったものにする．

Step 5: 教材提供

Step 4 で推薦した科目に関する教材を提供する．教材は Web ページと Youtube の動画の 2 種類である．図 4.7 を用いて説明する．まず表示されている推薦科目のどれかを選択し，次のページに移動するとその科目のシラバスと授業計画，キーワードを教材へのリンクとして表示する．そのリンクを選択することで教材提示ページへ移動する．

そのページでは Step 1 で作成した Web ページ，YouTube 動画の教材が保存された csv ファイルを読み込み，表示する．教材提供ページでは上から評価済み Web ページ，未評価 Web ページ，評価済み YouTube 動画，未評価 YouTube 動画の順で教材が表示される．評価済み Web ページと評価済み YouTube 動画に関してはそれぞれ評価点の上位 3 件ずつ表示する．また，評価済みのものとは別に未評価 Web ページと未評価 YouTube 動画も 3 件ずつ表示する．

Step 6: 教材へのレビュー

ユーザーは提供された教材を用いて学習を行い，その教材に対してレビューを投稿できる．このとき，評価スコアとレビュー本文の 2 種類を投稿できるようにする．評価スコアに関

しては5段階評価をできるようにしてあり、非常に悪いを1、悪いを2、普通を3、良いを4、非常に良いを5としてある。レビュー本文には自由にテキストを入力でき、500文字までの文字数制限を設ける。ユーザーはその教材を用いて学習したあとに2種類のレビューを入力し、評価ボタンを選択することでレビューを投稿できる。

このとき、2種類のレビューの内どちらか1つでもレビューをされていないと、エラーを返すようにする。これはどちらのレビューも信頼性スコアを算出する際に必須であるためである。また、レビュー本文の方は1文字以下で投稿された場合にもエラーを返すようにしてある。

ユーザーがレビューを投稿すると、HTMLから評価スコアとレビュー本文の情報をPythonで受け取り、それをCSVファイルに書き込む。レビューをCSVファイルに書き込む際に、レビュー評価点とレビュー本文に加え、ユーザーの学籍番号、投稿したWebページ、動画名、投稿したWebページのURL、動画URL、投稿された日時がログとして記録される。

Step 7: 教材の更新

ユーザーのレビューが投稿されると同時に投稿されたレビューを反映した教材の評価点の更新を行う。教材の評価点の更新は、投稿されたレビューの類似性スコア、集中性スコア、情報性スコアの3つのスコアを求め、それら3種類のスコアを用いてサクラ性スコアを求める。

そして、過去に投稿されたレビューのサクラ性スコアと新しく投稿されたレビューのサクラ性スコアを用いて、その教材のサクラ性スコアを更新する。このサクラ性スコアは高ければ高いほどサクラとして疑われるレビューであるため、評価スコアの平均が高くてサクラ性スコアが高いと最終的な教材のスコアは低くなる。逆に、評価スコアの平均が低くてもサクラ性スコアが低ければ最終的な教材のスコアは高くなる。

また、評価点の方もサクラ性スコアと同様に過去のレビュー評価点と新しく投稿された評価点を用いて平均の評価点を求める。そして、サクラ性スコアと評価点を用いて最終的な教材の信頼性スコアを求める。この信頼性スコアを用いて教材の更新、並び替えを行う。

教材の更新、並び替えは信頼性スコアの高い順に並び替えられる。信頼性スコアについても評価スコア、レビュー本文と同様にログとして記録される。ログを確認することでどの学籍番号のユーザーがどのタイミングでどのようなレビューを投稿しているかを確認することができる。

教材へのレビューを投稿すると、以上の流れでスコアが計算される。そして、スコアが反映された教材をページの再読み込みの形でユーザーの提示する(図4.8参照)。このときStep 6で説明したようにスコアが高い上位3件の教材を評価済みWebページ、YouTube動画に表示し、まだスコアが算出されていない(レビューが投稿されていない)教材を未評価Webページ、YouTube動画に表示する。この教材の更新の一連の流れはユーザーがレビューを投稿されるたびに行われる。そのため、ユーザーに対して常に最新のスコアが反映された教材を提供することが可能である。

以上の7ステップが学生が使用するシステムのアルゴリズムである。

数値実験ならびに考察

§ 5.1 数値実験の概要

本研究の数値実験としてシステムの有用性，科目推薦の有効性の検証，作成したシラバスの評価を行う。

システムの有用性の検証では実際にシステムを使用してもらい，アンケートに答えてもらう。アンケートの項目は全部で10個あり，その10個には必ず答えてもらう。また，アンケートと同時にコメントを記入できる欄を設けておき，自由にコメントをできるようにする。このアンケートを持って本システムの有用性の検証を行う。アンケート項目は表5.1に示す。以上のアンケートを5段階のリッカート尺度で評価してもらう。

リッカート尺度とは，あるトピックに対して，多段階の選択肢を用いたアンケートを取り，回答者がどの程度同意するか測定する手法である。今回のアンケートでは5段階のうち，1を全く満足していない，2をあまり満足していない，3をどちらでもない，4をやや満足している，5を非常に満足しているといったようにアンケートに答えてもらう。

表5.1を見てわかる通り，アンケート項目全体を通して，基本的にはシステムの使用感に関する質問を多くしてある。本来なら本システムを使用し，成績が向上するのを確認することでシステムの有効性を検証するべきだが，成績向上を確認するには開発したシステムを最低でも半期使用してもらわないと成績が向上したかを確認することができない。そのため，本研究ではアンケート調査を用いてシステムの有効性を示す。

調査の対象は同研究室の修士1年，学部4年，3年生の合計11人に実際に開発したシステムを使用してもらい，アンケートを答えてもらった。実際に使用してもらうにあたり，システムの使用手順について説明を行い，実際に使用してもらう。手順は以下に記してある通り，新規登録から推薦科目の選択，教材のレビュー，ログアウトまでの一連の流れを以下のように説明した。

1. 新規登録のページへ移動し，自分の学籍番号とパスワードを入力し，新規登録してください。
2. 成績登録ページに移動するので，そこで自分の成績を入力してください。
3. 成績を登録し終えたら，自分の成績と推薦科目が表示される画面に移動します。
4. 推薦科目に移動すると，シラバスとそれに応じたキーワードが表示される画面に移動します。
5. 授業計画，キーワードのページに移動すると教材提供画面が表示されるので，レビューを行ってください。
6. レビュー等を行ったら，上にある「キーワード選択画面に戻る」，「科目選択画面に戻る」を押してもらい自分の成績と推薦科目が表示されている画面に移動し，上にある

表 5.1: アンケート項目一覧

システムの操作性はわかりやすいか	システムにレイアウトは親切か
デザインは見やすいか	システムの機能はすぐに理解できたか
システムを使用するのにストレスを感じないか	推薦結果は理解しやすいか
表示されている教材は教材として適しているか	教材は学習に役立ちそうか
このシステムで効率よく学習できそうか	このシステムで学習すると効果が上がると思うか

「ログアウトする」を押して、ログアウトしてください。

7. ログイン画面に移動するので、新規登録で登録した学籍番号とパスワードを入力して、ログインできるかを確認してください

以上の流れの1と2を実行してもらえると、図5.1、図5.2のようなデータが取得できる。図5.1は新規登録後に作成されるcsvファイルであり、1列目に学籍番号で2列目にハッシュ化したパスワードが追加されていく。ログインする際には、登録してもらった学籍番号とパスワードを入力してもらう。このとき、新規登録で作成されるcsvファイルに、入力した学籍番号とハッシュ化したパスワードが一致した場合に成績と推薦科目が表示される画面に移動する。

図5.2は成績入力したあとに作成されるcsvファイルであり、このcsvファイルには1列目に学籍番号で2列目以降には各科目の成績（GPA）が追加されていく。このcsvファイルと過去学生データを用いて推薦科目を選出する。また、成績と推薦科目を表示する画面においてもこのcsvファイルを使用して表示する。

また、教材のレビューを行なってもらう（流れの5）と図5.3に示してある処理を実行する。図5.3に左側は教材提供ページで読み込むcsvファイルであり、このcsvファイルには1列目にWebページのタイトル、2列目にWebページのURL、3列目にWebページを識別するための番号、4列目に信頼性スコアが格納されている。レビューが投稿されると、図5.3の右側のようなcsvファイルにレビュー内容が書き込まれる。このcsvファイルは1列目にWebページのタイトル、2列目にWebページのURL、3列目にWebページを識別するための番号、4列目に投稿したユーザー（学籍番号）、5列目に評価スコア（5段階評価）、6列目にレビューの文章、7列目に投稿された日時、8列目に信頼性スコア、9列目にその教材の評価スコアの平均が格納される。

同じ教材に対してレビューが投稿されると一番下の行に追加されていき、8列目の信頼性スコアと9列目の評価スコアの平均は下の行にいくにつれて更新されていく。そして、更新された信頼性スコアで図5.3の左側のcsvファイルに更新していき、この信頼性スコアによって並び替えを行う。

次に科目推薦の有効性の検証を行う。この検証では科目推薦で推薦された科目がユーザーにとって気になる科目かどうかを検証する指標としてPrecision（適合率）を用いる。また、推薦科目に偏りがなく幅広く推薦できているかを測る指標としてCatalogue Coverage（カタログカバレッジ）を用いる。

Precision

Precision は、レコメンドリストにあるアイテムのうちユーザーが嗜好したアイテム

2155016	255b6e5d33edec2fc
2155017	b47fe49567891cda1
2155018	02b6e8796b892817
2155020	39875d0744a5f839c
2020025	9f54ec2ded21559c1

図 5.1: 新規登録結果

2155016	2	3	3
2155017	3	-1	-1
2155018	2	-1	-1
2155020	3	-1	1
1915026	4	4	-1
1915005	1	2	2

図 5.2: 成績入力結果

の割合のことであり、式 (5.1) で求めることができる。これは 1 に近づくほど推薦したアイテムがユーザーの嗜好にマッチしていることになる。

$$Precision@N = \frac{a \cap p_N}{N} \quad (5.1)$$

このとき N は考慮する上位ランキングの数を表しており、今回は $N = 10$ で設定して行う。つまり本研究では、推薦科目の上位 10 件のうち、ユーザーが気に入ったものが何件あるかを測る。これを人数分行うため式 (5.2) で Precision の全体の平均を求める。このとき $H = 11$ である。

$$Ave_{precision} = \frac{1}{H} \sum_{i=1}^H Precision@10_i \quad (5.2)$$

Catalogue Coverage

Catalogue Coverage は、利用可能なアイテムのうち、1 回のレコメンドでどのくらい多くのアイテムをレコメンドできたかを示す指標である。レコメンドのアイテム方向への広がりやカバー率を表し、この値が大きいほど多種多様なアイテムをレコメンドできたことを示し、式 (5.3) で求めることができる。これを求めることで、推薦科目の偏りなくユーザーに推薦できていることを確認する。

$$CatalogueCoverage = \frac{|S_r|}{|S_a|} \quad (5.3)$$

最後に、作成したシラバスフォーマットで作成したシラバスの評価を行う。シラバスの評価では、15 回の授業計画を含む全ての項目を埋めて作成したシラバスと授業計画に欠損や内容の重複があるシラバスの比較を行う。そのため、まずはシラバスを作成する。今回は授業計画で比較できるように、授業計画が欠損もしくは重複している科目を対象にシラバスを作成する。

そして、どちらのシラバスの方が参考になるかやこういった部分が参考になったかといったアンケートを学生に回答してもらい、アンケート結果をもとに評価・考察を行う。アンケートは、「改善したシラバスの方が学習に役立ちそうか?」、「授業計画が充実している方が、シラバスとして機能するか?」、「作成したシラバスには必要な情報が記載されているか?」の 3 個の項目を「はい」か「いいえ」の 2 値で回答してもらう。

学習ツール指 https://teach	2	10.1158673							
アルゴリズム https://lectui	0	3.43636364							
アルゴリズム https://lectui	1								
【コラム】Py https://www.	3								
ythonブログ https://note.	4								
ythonで学ぶ https://www.	5								
ython超入門 https://scho	6								
ythonで学ぶ https://7net.	7								
新・標準プロ https://www.	8								
【初心者へ上 https://blog.	9								
入門者・初心 https://kato-	10								
[PDF] 機械学 https://www.	11								
アルゴリズム https://zero2	12								
いまさら聞け https://www.	13								
[Python入門 https://atma	14								

アルゴリズム https://lectui	0	2155016	3	普通でした	2022/12/12 19:57	0	3
学習ツール指 https://teach	2	2155016	5	とてもわかり	2022/12/12 19:58	0	5
学習ツール指 https://teach	2	2155016	5	非常に良い	2022/12/12 19:58	0	5
アルゴリズム https://lectui	0	2155016	4	わかりやすい	2022/12/12 19:58	3.43636364	3.5
学習ツール指 https://teach	2	2155016	5	初心者におす	2022/12/12 19:58	10.1158673	5

図 5.3: レビュー時のデータ

§ 5.2 実験結果と考察

表 5.2 はそれぞれの質問項目に対するアンケート結果である。合計すると、全く満足していないが 1 件、あまり満足していないが 10 件、どちらでもないが 21 件、やや満足しているが 56 件、非常に満足しているが 22 件となった。それぞれの質問項目について以下で考察する。

1 個目は「システムの操作方法はわかりやすいか?」といった質問を行った。結果としてやや満足しているが 6 件と非常に満足しているが 3 件となった。この結果からシステムの操作が容易であることが考えられる。このことから本システムを始めて利用する人でもある程度すぐに操作ができるということが考えられる。

2 個目は「システムのレイアウトは親切か?」といった質問を行った。結果としてやや満足していると非常に満足しているがそれぞれ 5 件ずつとなった。この結果から履修状況と推薦結果の表示ページや教材提供ページ、成績入力ページの表示の仕方が適切であると考えられる。

3 個目は「デザインは見やすいか?」といった質問を行った。結果としてやや満足しているが 5 件と非常に満足しているが 2 件ある。この結果から本システムの全体を通したデザインが比較的わかりやすいことが考えられる。しかし、どちらでもないが 3 件、あまり満足していないが 1 件ある。事前に手順を説明するだけではユーザーにイメージさせずらいと考えられる。そのため、システム全体の流れをユーザーに伝えるといったところを改善する必要がある。

4 個目は「システムの機能がすぐに理解できたか?」といった質問を行った。結果としてやや満足しているが 5 件、非常に満足しているが 1 件ある。この結果から本システムには 1 つ目の質問の結果と同様に操作のわかりやすさ・容易さが高いと考えられる。しかし、あまり満足していないが 2 件、どちらでもないが 3 件ある。これは操作画面においてシステム動作についての説明の記載がないのが原因だと考えられる。そのため、システムを利用するにあたり機能説明についてわかりやすく明記する必要がある。

5 個目は「システムを使用するのにストレスを感じないか?」といった質問を行った。結果としてやや満足しているが 4 件、非常に満足しているが 2 件ある。この結果から初めて利用する人でもシステム使用をスムーズに行えることが考えられる。しかし、どちらでもないが 3 件、あまり満足していないが 1 件あることから、4 つ目の質問と同様にどのような

表 5.2: アンケート結果

質問内容	全く満足していない	あまり満足していない	どちらでもない	やや満足している	非常に満足している
システムの操作方法は分かりやすいか?	0	1	1	6	3
システムのレイアウトは親切か?	0	0	1	5	5
デザインは見やすいか?	0	1	3	5	2
システムの機能がすぐに理解できたか?	0	2	3	5	1
システムを使用するのにストレスを感じないか?	0	1	3	5	2
推薦結果は理解しやすいか?	0	2	2	4	3
表示されている教材は教材として適しているか?	0	0	7	4	0
教材は学習に役立ちそうか?	0	2	0	8	1
このシステムで効率よく学習できそうか?	0	1	1	8	1
このシステムで学習すると効果が上がると思うか?	1	0	0	6	4

操作を行えばいいのかなどの記載がないのが原因だと考えられる。

6 個目は「推薦結果は理解しやすいか?」といった質問を行った。結果としてやや満足しているが 4 件、非常に満足しているが 3 件ある。この結果から開発したシステムが推薦する科目はユーザーから納得されやすいことが考えられる。しかし、あまり満足していないとどちらでもないがそれぞれ 2 件ずつあり、一部のユーザーの興味を満たす推薦ができていないことが考えられる。これは、成績と希望業種の 2 種類を考慮しただけでは少ないことが原因である。そのため、実際にユーザーがどのようなジャンルに興味があるのかといった要素を推薦に組み込む必要がある。

7 個目は「表示されている教材は教材として適しているか?」といった質問を行った。結果としてやや満足しているが 4 件、どちらでもないが 7 件となった。本システムは Web ページと YouTube から授業計画をもとに上位複数件を教材として使用している。そのため、スクレイピングしてきた中に授業計画と関係のない Web ページや動画を取得してしまっているのが原因だと考えられる。そのため、授業計画だけではなく授業概要やキーワードを用いたスクレイピングやスクレイピングを行う際に制限を設け、より教材として適しているものを取得する必要があると考えられる。また、教材に対してのレビューデータを蓄積できれば、そのような不適切な教材に対してのスコアが低くなり、表示される頻度が下がるため、より多数の人に使用してもらうことで改善されると考えられる。

8 個目は「教材は学習に役立ちそうか?」といった質問を行った。結果としてやや満足しているが 8 件、非常に満足しているが 1 件となった。この結果から本システムが提示する教材は学習に活用することができる。しかし、あまり満足していないが 2 件ある。これは質問 7 で述べたように、スクレイピングの結果により教材として適していないものがユーザーに提示されたのが原因だと考えられる。これも質問 7 と同様にスクレイピング自体を改善することで解消できると考えられる。

9 個目は「このシステムで効率よく学習できそうか?」といった質問を行った。結果としてやや満足しているが 8 件、非常に満足しているが 1 件となった。この結果から本システムを使用することで効率よく学習ができると捉えることができる。これは、科目推薦のみや教材提供のみではなく、科目推薦と同時に教材を提供することによって得られた結果だ

表 5.3: Precision@10

	該当件数	Precision@10
ユーザー 1	6	0.6
ユーザー 2	7	0.7
ユーザー 3	7	0.7
ユーザー 4	8	0.8
ユーザー 5	5	0.5
ユーザー 6	8	0.8
ユーザー 7	8	0.8
ユーザー 8	4	0.4
ユーザー 9	3	0.3
ユーザー 10	6	0.6
ユーザー 11	8	0.8

表 5.4: Catalogue Coverage

	Catalogue Coverage
ユーザー 1	0.8877551
ユーザー 2	0.85714286
ユーザー 3	0.86170213
ユーザー 4	0.75824176
ユーザー 5	0.75
ユーザー 6	0.69565217
ユーザー 7	0.82105263
ユーザー 8	0.79569892
ユーザー 9	0.7752809
ユーザー 10	0.86021505
ユーザー 11	0.87777778

と考える。しかし、あまり満足していないとどちらでもないがそれぞれ1件ずつあることから、他の要素を組み込む必要があると考える。

10 個目は「このシステムで学習すると効果が上がると思うか?」といった質問を行った。結果としてやや満足しているが6件、非常に満足しているが4件となった。この結果から質問9と同様に本システムを学習に活用することは学習において有用であると考えられる。

また、自由記入であるコメントにおいて「意外な科目が推薦されたが、シラバスやキーワードを確認すると、自身の興味がある内容が含まれていた。」や「しっかりと的確なキーワードが抽出されており、表示される教材も分かりやすいものだった。」、「学校で提供される教科書よりも文章が簡潔かつわかりやすく勉強がはかどりそうだと思います。」などのポジティブな意見が得られた。

以上のアンケート結果の総括として、本システムはユーザーにとって十分有用であることを示せた。しかし、成績と業種以外の要素の組み込みやシステムデザインなどの改善すべき点がアンケートから確認できた。

次に、Precision についての結果を表 5.3 に示す。この実験では開発システムが提示する推薦科目の上位 10 件のうちユーザーが何件に興味を持ったかについて調べる。そこで実際に推薦科目画面を確認してもらい、10 件中何件に興味を持ったかを質問した。表 5.3 から $Ave_{precision}$ が 0.6363 となる。このことから、開発したシステムはユーザーに対して上位 10 件のうち、6 件の割合でユーザーの興味を満たす科目を推薦できていることを示す。アンケートの 6 個目の質問と同様にこれは今回は推薦の際に考慮する要素が成績と希望業種の 2 種類であったため、この要素とは別に他に考慮する要素を増やすことで Precision の値をよりよい値にできると考えている。開発システムが推薦する科目は絶対取得しなければいけないというわけではなく、参考程度に捉えてもらうものである。そのため、この 0.6363 という値から本システムは推薦システムとして機能していると捉えることができる。

Catalogue Coverage についての結果を表 5.4 に示す。この実験では、開発したシステムが推薦する科目に偏りが無いについて調べる。結果として平均は 0.81277448 となった。つまり、80% 以上の割合で多くの科目を推薦できたことを示している。Catalogue Coverage は使用目的によって値はさまざまであるが、一般的には 70% 以上あれば良いと言われているため、80% という値は十分に機能しているといえる数値である。これは本システムが推

履修科目名	心理学Ⅰ 後期	
履修科目名(英語)	Psychology I	
科目区分	教養	
配当学年	工学部 1年	
担当教員		
職階	氏名	所属
講師	井戸 啓介	教養教育センター
開講学期	後期	
単位数	2	
単位区分	選択	
関連する学習・教育目標	情報システム工学科：(C)-1 知能ロボット工学科：(A)-1、(A)-3 情報システム工学科：(A)-1、(A)-2、(A)-3、(A)-4 電気電子工学科：(A)-1、(A)-2、(A)-3、(A)-4 機械・社会基盤工学科：(A)-1、(A)-2、(A)-3 生物工学科：(A)-1、(A)-2 医薬品工学科：(A)-1、(A)-2	
授業の目標・授業概要	心理学とは行動や心のほたらきを科学的に説明する学問である。この講義では、科学的な考え方から始め、人間が情報を受け入れ、認識する過程の基礎を解説する。脳科学からの知見も取り上げる。	
学生の到達目標	①人間の心理や行動に対する科学的な研究方法を理解し、考察できること ②人間が外界を認識し行動する際の基礎的な特性について理解を深めること	
授業計画	①心理学とは何か ②心理学の領域と研究方法 ③科学的な考え方 ④人間の知覚と意識(1) ⑤人間の知覚と意識(2) ⑥脳神経科学と心理学(1) ⑦脳神経科学と心理学(2) ⑧感覚・知覚(1) ⑨感覚・知覚(2) ⑩認知・記憶(1) ⑪認知・記憶(2) ⑫認知・記憶(3) ⑬まとめ	
キーワード	科学的方法論、脳神経科学、感覚、知覚、認知、思考、記憶、意識	
成績評価基準	期末試験の成績により評価する。 なお、出席率が2/3未満の場合は、原則として単位を認定しない。	
教科書・教材参考書等	教科書：使用しない。 参考書：黒川隆『心理学』（有斐閣）ISBN：4641053693	
関連科目・履修条件等	心理学Ⅱを併せて履修すると理解が深まる。	
履修上の注意事項や学習上の助言	学生の皆さんは授業を受けるに当たっては、予習・復習を怠らないように願っています。	
学生からの質問への対応方法	授業の際、あるいは教員室で受け付ける。オフィスアワーは授業開始に通知する。 来室は随時受け付けますが、電子メールなどで予約してもらえば確実です。（メールアドレス：ido@pu-toyama.ac.jp）	

図 5.4: もとのシラバス

授業科目名	心理学Ⅰ	
授業科目名(英語)	Psychology I	
科目区分	教養	
配当学年	工学部1年	
担当教員		
職階	氏名	所属
教授	清水愛士	情報システム工学科
開講学期	後期	
単位数	2	
単位区分	選択	
関連する学習・教育目標	情報システム工学科：(A)-1、(A)-2、(A)-3、(A)-4	
授業の目標・授業概要	心理学とは行動や心のほたらきを科学的に説明する学問である。この講義では、科学的な考え方から始め、人間が情報を受け入れ、認識する過程の基礎を概説する。脳科学からの知見も取り上げる。	
学生の到達目標	①心理学に関する基本的な理論・方法について理解し、わかりやすく説明することが出来る。②心理学の基礎知識を習得したうえで、自ら考え、人間理解や社会での応用実践の方法を探る視点を持つことが出来る。	
授業計画	1回目：心理学とは 2回目：感覚器官、図と地、反転図形 3回目：記憶、実行意知覚 4回目：古典的学習、オペラント学習、社会的学習 5回目：感覚記憶、短期記憶、長期記憶 6回目：感情、帰属、欲求 7回目：ピアジェの理論、認知発達、分離不安 8回目：フロイトとユングの理論、心理療法 9回目：性格・パーソナリティ(1) 10回目：性格・パーソナリティ(2) 11回目：社会問題と人間行動 12回目：社会的認知、対人魅力、対人関係 13回目：コミュニケーション、社会的スキル 14回目：集団の特性、社会的影響過程 15回目：応用心理学と最新の心理学動向	
キーワード	感覚 知覚 認知 思考 記憶 意識 脳神経科学	
成績評価基準	期末試験の成績により評価する。	
教科書・教材参考書等	教科書：使用しない	
関連科目・履修条件等	心理学Ⅱを併せて履修すると理解が深まる	
履修上の注意事項や学習上の助言	予習・復習を怠らないように努めて下さい	
学生からの質問への対応方法	授業の際、あるいは教員室で受け付ける	

図 5.5: 改善したシラバス

薦するシステムはただ単に予測評価値の高い科目のみを推薦している訳ではなく、必修科目、選択必修科目、卒業要件単位をしっかりと考慮できていることを示す。

最後に、作成した Web シラバスのフォーマットを用いて作成したシラバスともとのシラバスの比較を行う。図 5.4 は授業計画の内容が複数重複しているシラバスで、図??は授業計画を 15 回全て書き換えて作成したシラバスである。書き換え部分としては授業計画の部分であり、図 5.4 の方では 15 回分書かれているが、内容が脳神経科学と心理学(1)、脳神経科学と心理学(2)といったように重複している部分がある。その部分を重複させずに全て違う内容で書き換えたのが図 5.5 である。

これらを学生に対して提示し、シラバスの比較を行なってもらいアンケートへの回答およびコメントをしてもらった。アンケート結果を表 5.5 に示す。アンケート結果から改善したシラバスの方が授業計画が重複しているシラバスよりも学生にとって有用であることが考えられる。今回は改善したシラバスの情報を打ち込んだのは自分であるため、「必要な情報が記載されているか？」で「いいえ」が他の質問項目より多くなってしまった。しかし、この部分は教員側が書くことで改善できると考えている。アンケート結果の総括として、授業計画の充実度によって学生のシラバスへの満足度が高くなることがわかった。

図 5.4 のように重複している科目もあれば、教科書の Unit 1 や Unit 2 を授業計画に記入している科目もあり、この場合教科書を所有していないとどのような内容の授業かが不明であり、学生にとって参考にならないシラバスとなってしまう。また、図 4.2 のように授業計画の部分に 15 回書かれていない科目もあるため、このような科目についても作成したフォーマットを使用してもらうことで、学生にとって参考になるシラバスを作成することが可能になる。以上のことから、作成したフォーマットを使用することで学生にとって有益なシラバスを作成することが可能になると考えられる。

表 5.5: アンケート結果

質問項目	はい	いいえ
改善したシラバスの方が学習に役立ちそうか？	9	2
授業計画が充実している方がシラバスとして機能するか？	11	0
改善したシラバスには必要な情報が記載されているか？	8	3

おわりに

本研究では学生の成績データを用いてユーザーベース協調フィルタリングを行い、ユーザーへの科目の推薦による科目選択の支援、各科目における教材をユーザーに提示することで学習の支援を行うシステムの開発。また、ラーニングアナリティクスに向けたシラバス標準化のためのシラバスフォーマットの作成を行った。

ユーザーに履修した科目と希望業種のデータを入力してもらい、入力されたデータと過去学生データを協調フィルタリングを用いて予測評価値を算出し、そこから必修単位等の条件を満たし、履修可能科目の中から科目の推薦を行った。また、教材の作成では富山県立大学の Web シラバスから授業計画、キーワードを用いて Web ページと YouTube 動画をスクレイピングしユーザーに提示する。

教材に関してもただ提示するのではなくレビュー機能を追加し、各レビューに対してスコアを算出しそのスコアをもとに教材間でランキングを行い、学生に最適な教材を提示できるようになっている。また、シラバス標準化のためのフォーマットを作成した。これは富山県立大学の Web シラバスを参考に作成し、授業計画や授業概要といった全 20 項目を記入してもらい、学生にとって有用な情報を提供できるシラバスにする。また、記入してもらった授業計画を学生用システムの教材作成に活かす。

数値実験では実際に学生用システムを複数人に使用してもらいシステムの使用感に関するアンケートに答えてもらった。その結果、アンケート全項目を通して非常に満足している、やや満足しているが多く、このことから開発したシステムの有用性を示した。しかし、教材に不適切な教材が混ざっていることや、推薦結果の理解の部分でネガティブな意見を確認した。これは授業計画のみのスクレイピングや成績データと希望業種の 2 種類のみしか考慮していない点が原因だと考えられる。また、Precision 値と Catalogue Coverage 値を求めた。Precision の値から本システムではユーザーの興味を満たす科目を 6 割を推薦できていることを確認した。また、Catalogue Coverage の値から推薦可能科目の中から 80% 以上の割合で幅広い科目を推薦できていることを確認した。また、作成したフォーマットをもとにシラバスを作成し、もとのシラバスと比較した。その結果、フォーマットを通して作成したシラバスの方が学生にとってより有益な情報を提供できることを確認した。

今回はアンケートを用いてシステムの有用性を確認したが、実際にシステムを半期以上使用してもらうことでどのような学習効果があるのかを確認し、有用性を示す必要がある。それと同時に、過去学生のデモデータ部分を実際のデータで行う必要もある。そしてシステム全体のデザインや教材作成におけるスクレイピングの見直し、成績と希望業種以外に考慮する要素の増加が今後の課題として挙げられる。

謝辞

本研究を遂行するにあたり，多大なご指導と終始懇切丁寧なご鞭撻を賜った富山県立大学電子・情報工学科情報基盤工学講座の奥原浩之教授，António Oliveira Nzinga Renè 講師に深甚な謝意を表します．最後になりましたが，多大な協力をして頂いた，奥原研究室の同輩諸氏に感謝致します．

2023 年 2 月

清水 豪士

参考文献

- [1] 山田 政寛, “ラーニング・アナリティクス研究の現状と今後の方向性”, 日本教育工学会論文誌, 41 巻, 3 号, pp.189-197, 2018.
- [2] “M2B（みつば）学習支援システム”, 九州大学基幹教育院 ラーニングアナリティクスセンター.
- [3] 鶴田 美保子, “大学生の就職活動を成功させる要因：展望論文”, 金城学院大学論集 人文科学編, 第 15 巻, 第 1 号, pp.109-119, 2018.
- [4] 鶴田 美保子, “大学生の就職活動を成功させる要因 コロナ禍における女子大学生の調査”, 金城学院大学論集 人文科学編, 第 18 巻, 第 2 号, pp.111-122, 2022.
- [5] 渡辺 裕子, “本学学生の民間企業への就職活動の実態と成功要因”, 経済研究所所報, 第 19 号, pp.31-50, 2015.
- [6] 由谷 真之, 森 幹彦, 喜多 一, “電子シラバスを用いた大学教養教育における科目選択支援”, 情報処理学会第 68 回全国大会, 4U-3, pp.4-469 - 4-470, 2006.
- [7] 林坂 弘一郎, 大塚 萌佳, “遺伝的アルゴリズムによる大学生の最適履修時間割推薦”, 神戸学院大学経営学論集, 第 18 巻, 第 1 号, pp.45-60, 2021.
- [8] 堀 幸雄, 中山 堯, 今井 慈郎, “科目ネットワーク上の活性伝播を用いた時間割の自動生成システム”, 情報処理学会論文誌, Vol.52, No.7, pp.2332-2342, 2011.
- [9] “e ラーニングとは e-Learning の意味”, <https://www.digital-knowledge.co.jp/el-knowledge/e-learning/>, 閲覧日 2023. 1. 22.
- [10] “e ラーニングとは — SATT - エスエイティーティー”, <https://satt.jp/e-learning/e-learning.html>, 閲覧日 2023. 1. 22.
- [11] “矢野 経済 研究所、国内 e ラーニング市場の調査結果を発表”, <https://edu.watch.impress.co.jp/docs/news/1403045.html>, 閲覧日 2023. 1. 22.
- [12] “国内大学の e ラーニング事例からみる導入の課題”, <https://www.uicommons.co.jp/topics/a62>, 閲覧日 2023. 1. 22.
- [13] 増岡 由貴, 辰己 丈夫, “大学における e ラーニング導入教育についての考察”, 情報教育シンポジウム 2013 論文集, 2013 巻, 2 号, pp. 141-146, 2013.
- [14] 中條 桂子, 南野 奈津子, “不登校児童生徒の学習支援における e ラーニングの活用に関する考察”, ライフデザイン学研究 第 15 号, pp.371-386, 2019
- [15] 今枝 加与, 内藤 圭子, 松田 奈美, 長濱 優子, 後藤 淳子, 杉本 なおみ, 長谷川 しとみ, “新人看護師の基礎看護技術指導に e ラーニングを導入して”, 日農医誌, 64 巻, 5 号, pp.877-881, 2016.

- [16] 西上 あゆみ, 緒方 巧, 湯浅 美香, “e ラーニングを活用した基礎看護技術の学習支援の評価”, 梅花女子大学看護学部紀要, 第 5 号, pp.31-39, 2015.
- [17] 古賀 崇朗, 福崎 優子, 久家 淳子, “e ラーニングを活用した職員研修の実践”, 2018 PC Conference, pp.287-288, 2018.
- [18] 富永 敦子, 向後 千春, “e ラーニングに関する実践的研究の進展と課題”, 教育心理学年報, 53 巻, pp.156-165, 2014.
- [19] 松田 岳士, 渡辺 雄貴, “教学 IR, ラーニング・アナリティクス, 教育工学”, 日本教育工学会論文誌, Vol.41, No.3, pp.199-208, 2017.
- [20] 緒方 広明, “ラーニングアナリティクス：教育ビッグデータの分析による教育改革”, Nextcom, Vol.45, pp.12-21, 2021.
- [21] Philipp Leitner, Mohammad Khalil and Martin Ebner, “Learning Analytics in Higer Education - A Literature Review”, *Learning Analitics: Fundaments, Applications, and Trends*, pp1-23, 2017.
- [22] 島田 敬士, “大学教育における学習分析の活用事例”, 情報処理学会論文誌 教育とコンピュータ, Vol.6, No.2, pp.16-24, 2020.
- [23] “IR (Institutional Research) とは? コアネット教育総合研究所”, <https://core-net.net/keywords/kw003/>, 閲覧日 2023. 1. 22.
- [24] 神嶌 敏弘, “推薦システムのアルゴリズム”, <https://www.kamishima.net/archive/recsy-sdoc.pdf>, 閲覧日 2023. 1. 22.
- [25] 土方 嘉徳, “嗜好抽出と情報推薦技術”, 情報処理学会論文誌, Vol.47, No.4, pp1-10, 2006.
- [26] “How Online Reviews Influence Sales”, <https://spiegel.medill.northwestern.edu/how-online-reviews-influence-sales/>, 閲覧日 2023. 1. 23.
- [27] Shen Huang, Dan Shen, Wei Feng, Yongzheng Zhang and Catherine Baudin, “Discovering clues for review quality from author’s behaviors on e-commerce sites”, *Proceedings of the 11th International Conference on Electronic Commerce*, pp.133-142, 2009.
- [28] “食ベログ判決が社会に与える影響について”, <https://www.valueup-jp.com/2022/09/15/column-vol-l45/>, 閲覧日 2023.1. 23.
- [29] 油布 翔平, “ソーシャルメディアにおけるアカウント集団特定によるキャンペーンの検出”, *Japan Advanced Institute of Science and Technology*, pp.1-41, 2021
- [30] Nitin Jindai, Bing Liu, “Opinion Spam and Analysis”, *Proceedings of the 2008 International Conference on Web Search and Data Mining*, pp. 219-230, 2008.
- [31] 三船 正暁, 金 明哲, “ネットショッピングにおけるスパムレビューの特徴分析”, 日本計算機統計学会大会論文集, 第 30 回, pp. 9-12, 2016.

- [32] 岡山 光平, 石川 博, 廣田 雅春, “フェイクニュース分類器を用いた口コミサイトのレビュー分析”, *DEIM Forum*, pp.3-5, 2018.
- [33] 前田 正子, “甲南大学マネジメント創造学部（CUBE 生）の GPA 及び能力向上感に影響を与える要因についての調査報告”, *Hirao School of Management Review*, Vol.6 ,pp. 1-17, 2016.
- [34] 伊木 惇, 亀井 清華, 藤田 聡, “レビューを対象とした信頼性判断支援システムの提案”, *情報処理学会論文誌*, Vol.55, No.11, pp. 2461- 2475, 2014.
- [35] Sihong Xie, Guan Wang, Shutang Lin and Philip S. Yu, “Review Spam Detection via Temporal Pattern Discovery”, *Proceedings of the 18th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 823-831, 2012.
- [36] Sheng Zhang, Amit Chakrabarti, James Ford and Fillia Makedon, “Attack Detection in Time Series for Recommender Systems”, *Proceedings of the 12th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp.809-814, 2006.
- [37] Johnson, M. K. and Rate, C.L., “Reality monitoring”, *Psychological Review*, Vol. 88, No. 1, pp. 67-85, 1981.
- [38] Paul Rayson, Andrew Wilson, “Grammatical word class variation within the British National Corpus Sampler”, *Language and Computers*, Vol. 36, No. 1 ,pp. 295-306, 2002.