

卒業論文

教学データの学習エビデンスに基づく GPA向上のための 情報推薦・学習支援システムの開発

Development of Information Recommendation
and Learning Support System for GPA Improvement
Based on Learning Evidence from Teaching and Learning Data

富山県立大学 工学部 電子・情報工学科

1815043 滝沢 光介

指導教員 奥原 浩之 教授

提出年月: 令和4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	2
第2章 教学データ分析と情報推薦	4
§ 2.1 教学におけるビッグデータ・アナリティクス	4
§ 2.2 情報推薦と協調フィルタリング	6
§ 2.3 レビューの信頼性の判断支援	9
第3章 取得科目の推薦と教材の最適化	12
§ 3.1 協調フィルタリングからの科目の推薦	12
§ 3.2 シラバスからの教材作成	13
§ 3.3 信頼性を考慮した教材の提示	14
第4章 提案手法	17
§ 4.1 予測評価値からの適切な授業科目の推薦	17
§ 4.2 関連資料のアップデート	17
§ 4.3 提案手法のアルゴリズム	18
第5章 数値実験並びに考察	19
§ 5.1 数値実験の概要	19
§ 5.2 実験結果と考察	19
第6章 おわりに	20
謝辞	21
参考文献	22

図一覧

3.1 シラバスからの授業計画のスクレイピング結果	14
3.2 シラバスからの授業計画のスクレイピング結果	14

表一覽

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
特定の利用者	x
特定のアイテム	y
利用者数	n
アイテム数	m
利用者集合 $\{1, \dots, n\}$	$\mathcal{X}$
アイテム集合 $\{1, \dots, m\}$	$\mathcal{Y}$
アイテム y を評価した利用者の集合	$\mathcal{X}_y$
利用者 x が評価したアイテムの集合	$\mathcal{Y}_x$
対象ユーザー	a
利用者 x のアイテム y への評価値	r_{xy}
利用者 x による評価値の平均	$\bar{r}_x$
アイテム y への評価値の平均	$\bar{r}_y$

はじめに

§ 1.1 本研究の背景

近年、ICT 技術の発展により、大規模なデータを容易かつリアルタイムに蓄積できるようになった。これらのデータは広い範囲で「ビッグデータ」と呼ばれ、これらのデータに対する分析はデータサイエンスやアナリティクスと呼ばれその需要を高めている。これらのことは教育機関においても例外ではなく、多くの学校、特に大学においては、生徒の学習内容ログや修学状況などの教学データが日々蓄積されている。

教学データにおける分析は 2000 年前後から盛んになり、最近では九州大学のラーニングアナリティクスセンターにおいて、蓄積された学生の様々な情報を整理、分析し学生に対してフィードバックを行うシステム、「M2B(みつば)」が実装されている [1]。

上記の M2B を筆頭に多くの大学で、生徒に対するアナリティクスシステムを提案するといった事例が多くの大学で行われている。

このように生徒の情報が蓄積された中で、大学における成績評価方法の一つとして成績平均値 (Grade Point Average: GPA) が存在する。GPA とは大学における学生の各科目における成績を S, A, B, C, D のレターグレードで表し、それらを、S=4, A=3, B=2, C=1, D=0, のグレードポイントに置き換え、それらの平均を算出したものである。GPA はその人の大学における成績の総合評価といえる。

大学での成績、特に GPA が就職活動において重要視されているかということについてはしばしば、議題に挙げられている。その理由は GPA の算出方法とその特徴が要因となっている。GPA は 60 点から 100 点までの 1 点刻みを点から 4 点までで刻み、0 点から 59 点までを 0 点で表現しているのでもとの素点と比べるとその精度は 0.1 になってしまう。また、GPA は授業の難易度、教授の成績のつけ方などの条件で大きく変化し、各大学においてかならずしも同じ尺度で決められた数字ではないといった特徴がある。

しかし、多くの研究で少なからず大学における成績と就職活動には関係性があり、良い成績を納めている生徒ほど就職活動が優位に働いているとされている [2,3]。さらには多くの大学において、成績優秀者に関して大学院進学における何かしらの優遇があることは事実である。このことから、大学時代における最大の目標である就職または進学の双方において、大学での成績が重要となってくる。

§ 1.2 本研究の目的

大学における GPA を上げる方法として、第一に履修した単位において良い成績を修めることである。しかし、大学生は一つの学期で多くの単位を取得する必要があるすべての履修単位について上位の成績を修めるのは困難である。また、自分が良い成績を修めることができるかどうかはその科目を履修してみないとわからない。さらには、大学における卒業要件単位を満たすために自分の不得意とする科目もある程度取得する必要がある。

それにも関わらず、大学における単位選択は学期ごとに生徒個人で行う必要があるその選択肢は膨大なものになる [4]。その結果、多くの学生、特に入学したての 1 年生に対して単位選択という行動がかなりの負担になる可能性がある。

そのため本研究では、過去の卒業生の教学データに対して分析を行い、システムを使用する学生がまだ取得していない科目において成績評価の予測を行い、高い GPA を取得できるように単位選択の推薦を行う。それに加え、すべての科目について良い成績がとれるように予測成績が低い科目については Web 上から関連情報を推薦するシステムの開発を行う。

まず、教学データにおけるデモデータの作成を行い、そのデータを蓄積する。蓄積したデータに対して、システムのユーザーを対象とし協調フィルタリングを行い、対象ユーザーのまだ取得していない科目について、どれくらいの成績をとれるかといった予測成績値を算出する。

協調フィルタリングによって得られた予測評価値をもとに、ユーザーに対して推薦する科目を決定する。科目を選択する際には、予測成績が高い科目から推薦を行う。また、富山県立大学における卒業要件単位を満たすように科目の選択を行う。卒業要件単位を満たすように単位を選択すると予測評価値が低い科目も選択してしまう可能性がある。そのように、予測評価値が低いにもかかわらず選択されてしまった科目については、Web 上からその科目に関する資料をスクレイピングし、スクレイピングした情報をユーザーに提供を行う。

また、Web 上からの情報について学生に評価を行ってもらい、学生からの評価が低い Web ページは授業の内容とあまり関係なく、教材としてふさわしくないと判断し出現頻度を下げ、より学生が学びやすい教材としてアップデートを行う。また、学生からの評価の信頼性の保証を行うために、学生からのレビューについていくつかの指標をもちいてレビューの信頼性を図る。そして、レビューの信頼性が低いにもかかわらず頻出している教材については、不当な評価を受けて頻出しているト判断し出現頻度を下げ、より学生が勉強しやすいシステムを目指す。

提案手法によって構築されたシステムの有効性を検証する必要がある。そのために、科目の推薦を行う際にしっかりとパーソナライゼーションできているかの確認を行う。また、完成したシステムを実際に学生に使用してもらうことでシステムの使用感や教材のアップデートがしっかりと行えているかを評価してもらい、システムの有用性を示す。

§ 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。背景では大学におけるビッグデータアナリ

ティクスについてと大学における成績評価法の信頼性について述べる．目的は教学データからの最適な単位選択とそれにおける情報の提供することを述べる．

第2章 大学におけるビッグデータアナリティクスの概要とその発展についてまとめる．また，教学データに対する分析手法である，協調フィルタリングとそれにおける情報推薦について述べる．

第3章 システムの基礎となる協調フィルタリングについて述べる．またシラバスからの教材作製の手順を述べる．

第4章 学生にとって最適となる科目の選択，それに関する教材のアップデートについて述べる．また，システム全体の流れを述べる．

第5章 本研究の情報推薦の個人化度合いについて述べる．また教材のアップデートによるユーザーの使用感を検証する．

第6章 本研究で述べている提案手法をまとめて説明する．また，今後の課題について述べる．

教学データ分析と情報推薦

§ 2.1 教学におけるビッグデータ・アナリティクス

教学におけるビッグデータの内容は学生の成績や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) や学生へのアンケート、学生の作業物、授業の学習過程の記録、自己評価や他者からの評価など多岐に渡る。教学におけるビッグデータは一般的なビッグデータと比べて以下のような特徴があるとされている。

1. データ量は大きくない

一大学に限れば、テラバイトの処理で十分であり、ペタバイトクラスのデータベースを構築している通信会社の分析とは良の規模が違う。

2. 対象人数は短いが、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、成績だけでなく、授業の出欠・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、日々増加していく。

3. 匿名性が低い

教学データは個人情報といっても過言ではないほど学生個人の人となりを表しておりその扱いは慎重に行う必要がある。

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ったのかは不明であるが、ある授業に出席したり、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り、学生がどんな考えをもって行動しているかといった情報が筒抜けになる。

5. 因果関係は複雑

個々のデータの意味は明白であるが、各データ同士の関係は不明であることが多く、相関程度しかわからないケースがある。

6. 多様化・細分化が進行している

教育の ICT 化に伴って、個人レベルのマイクロデータが集積している。多くの大学で、教学におけるビッグデータはいくつかのデータベースにわけられており、それらを扱う際は、データの集積と情報の取捨選択が必要となる。

これらの特徴からいえることは、教学におけるビッグデータ・アナリティクスにおける課題は大量のデータの実行時間が問題となるわけではなく、さまざまな意味を持つデータを分析したところでデータ同士の関係性をはっきりさせることが難しいことである。それに加え、3、4の特徴からデータにおける個人の把握が容易であり、個人情報として慎重に扱う必要がある。

また6の特徴に伴い、種類が多いにも関わらず分析する対象が小さいデータを分析してもあまりインパクトのない結果になってしまうと指摘している [5]。このことからラーニングアナリティクスに対する全盛期は終わったと述べて、分析結果における EDM・人工知能・学習化学などのコミュニティと協働する必要性があると述べている。

では、上記で述べたような小さいサイズに対する分析における研究はまったく意味がないものなのか。実はそうではなく、小さいサイズに対する分析を行い教育現場の課題解決を行うシステムを開発し、その有用性を示したシステム開発研究は数多く存在する。以下では小さいサイズに対するシステム開発研究の例をいくつか紹介し、システム実装におけるさらなる注意点を示す。

スモールデータに対するシステム開発研究

授業中の生徒同市または生徒と教師の円滑なコミュニケーションを目的とし、レスポンスアナライザシステムの導入を行っている [6]。システムの機能としては質問投稿、質問に対して返信を行える投票機能、表示機能、投稿・投票を確認できるログ分析機能がある。扱っているデータとしては学生の質問・投票データに加え、学生がどの程度アクティビティに参加したかがわかるログデータだけである。データだけを見ると、一つの授業に限ればかなりサイズの小さなデータとなる。しかしながら、これらのシステムを運用することで、授業中の生徒同氏の行動が活発になり授業全体の活性化が見られたと報告している。また、教師側も学生がどこで躓いているかを把握することができ、円滑な授業の組み立てができるようになったと報告されている。

eラーニングは利用者が事前に録画された動画や事前に作成された資料を際限なく視聴して学ぶことができる。そのような際限なく見ることができるのはeメンタの絶え間ない努力によるものである。しかし、利用者が増えたり利用者により高い質の支援を行おうとするとするとeメンタの負担が大きくなる。その問題を解決するためにeメンタの負担を軽減することを目的とし、システム開発を行っている [7]。システムの機能としては学習者タイプの分類とその表示機能、コース情報・メンタリングガイドライン入力とその表示機能、メンタリングアクションプラン自動生成機能がある。実践で扱われたデータは2クラス分のデータであり、学習者タイプの分類、タイプ別メッセージによる学習行動の変容、終了率など定量的なデータを元に考察を行っている。こちらの研究においてもデータだけで見るとかなり少ないデータ数となっているがeメンタの負担を減らすことができたと報告されている。

上記で挙げた特徴における課題を解決したとしても実際に教学データに対してラーニングアナリティクスのようなシステムを実装しようとしたときに、それらの以下のような規制が実装を阻む要因となる可能性がある。

1. 人事制度に柔軟性がない

教育機関がデータアナリストのような専門職を雇用できず、データ分析担当職員がジョブローテーションの対象になる。

2. 情報公開の文化が根付いていない

情報公開に前向きな大学は少なく、教学 IR やラーニング・アナリティクスの知見を公表できないことがある。さらに、先述した大学ポートレート構築の過程でも明らかになったように、他大学のデータを分析することも困難である。これは、ベンチマークの仕組みができない原因の一つでもある。

3. 大学経営の手札が少ない

米国などと異なり寄付収入が少ないことや、資金運用に関する制限があることなどから、分析結果を示されたとしても、有効な対策につながらないケースがある。

また、上記のような規制や制度の問題を解決し、必要なデータが収集でき、精度の高い分析したとしてもその分析結果を実際の教育現場で活用するには以下のようなさらなる問題を解決する必要がある。

1. 学生自身が改善的内要因への対応

例えば、経済的事情でアルバイトをせざるを得ない学生のアルバイト時間が長すぎることで成績低下の主要な要因であることが判明した場合、アルバイトの時間を短縮させることは困難である。また、学業に関する問題であっても、入試成績が大きな要因と特定されて、しかもリメディアル教育が実施されていない教育機関である場合、教育機関側も学生本人も対応は容易でない。

2. 学生のモチベーションに悪影響を与えないアプローチ

システムから留年の可能性が高いという判定が出た場合、そのまま学生に伝えてしまうと、学修の継続をあきらめてしまうことになりかねない。反対に、成績が上位で特に問題ないという結果を伝えた場合にも、モチベーションを下げるおそれがある。

このように教育ビッグデータにおけるラーニング・アナリティクスにおける諸問題は多く存在するが、多様なデータを分析し、意味を読み取るノウハウの提起強において教育ビッグデータにおけるラーニング・アナリティクスが貢献出来る価値は大きいといえる。また教育工学はの研究にはラーニングアナリティクスという概念が提唱される以前から、ラーニングアナリティクスの分析を行っている研究は少なくない。つまり、教育工学はラーニング・アナリティクス研究と重複する分野をカバーしているだけでなく、教育の諸問題を工学的アプローチで解決しようとするが故に、教育 IR とラーニング・アナリティクスの両者と協同する意義がある分野である。

§ 2.2 情報推薦と協調フィルタリング

近年、パーソナライゼーションといった言葉がよく使われている。パーソナライゼーションとは無数にある情報の中から、個人にあう情報を抜き出し、それらの情報を個人に適した形で表示することである。そういったパーソナライゼーションの核となる技術が情報推

薦である。情報推薦の目的は、利用者にとって重要と思われる、事象や対象、情報を利用者の目的に沿う形で提示することである。有名なものでは、Amazonで行われている「この商品を購入した人はこちらの商品も購入しています。」といったものである。

そもそもなぜこのようなパーソナライゼーションといった考えが必要になったのか。その背景はICT技術の発展に伴う情勢の中で大きく分けて二つ存在する。

- ・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
- ・大量の情報が誰でも、容易に手に入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態を情報過多、あるいは情報爆発という。こうした現象に対処するために生み出されたのが、情報推薦技術を利用したパーソナライゼーションである。

このような情報化社会と呼ばれる現代では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の情報推薦技術では主に利用者の嗜好を予測し、嗜好則した情報推薦を行う。情報推薦における嗜好の予測方法として内容ベースフィルタリングと協調フィルタリングの大きく二つに分類される。以下ではそれぞれの方法について述べたのちに、本研究で使用する協調フィルタリングについてさらに詳しく説明を加えていく。

内容ベースフィルタリング

内容ベースフィルタリングは、アイテム(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する。利用者が「原宿 イタリアン」と入力すると、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、類似度の高い店をおすすめの店として紹介する。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内容ベースフィルタリングである。

協調フィルタリング

協調フィルタリングは、利用者がシステムを利用する以前から利用者のアイテムに対する嗜好データをデータベースに蓄積した利用者データベースを保持している。協調フィルタリングではそれらの情報をもとに利用者がどの嗜好パターン(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、嗜好が似ている利用者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下、嗜好パターンが類似している利用者を見つけ出し、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法とモデルベース法の二つに分類される。メモリベース法はその中でもさらに利用者間型メモリベース法とアイテム間型メモリベース法の二つに分類される。情報推薦の全体像は図nのようになっている。

以下では協調フィルタリングにおけるメモリベース法とモデルベース法の二つの分類について説明し、さらにメモリベース法の中の利用者間型メモリベース法とアイテム間型メモリベース法について説明を記述する。

メモリベース法とモデルベース法

メモリベースは利用者がシステムを使用する以前は特に何もせず、ただ利用者データベースを保持しているだけである。利用者がシステムを使用する際に、利用者データベースを参照し嗜好パターンを読み取り、それに併せた情報を推薦する。メモリベースはシステムが利用されるごとにデータベースを参照するのでデータベースの削除や追加といった変化について柔軟に対応できるメリットがある一方で、データベースを逐次的に参照しているのでモデルベース法に比べて推薦時間を有することになる。

モデルベース法は利用者がシステムを使用する以前にあらかじめ「Aさんの好むものはBさんも好む」といった嗜好パターンをモデルとして構築する。そして、システムを利用するときに利用者データベースではなく、構築しておいたモデルに基づいて情報を推薦する。モデルベース法はシステムが利用される前からモデルを構築するので推薦時間を短くできるメリットがある一方で、データベースに変更を加えると一からモデルの構築を行う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ことがある。

利用者間型メモリベース法とアイテム間型メモリベース法

利用者間型メモリベース法の代表的な手法である GroupLens の手法について述べる。この手法は協調フィルタリングの手中を自動化した研究で有り、簡潔な手法であるがその予測精度は高く、多くの研究で改良されている。

お昼ご飯を食べる際にレストランを探す場合を考える。このとき、自分と食べ物の嗜好が似た何人かの意見を聞いて、それらの情報をもとにどこのレストランで食事を行うか決めることがあるだろう。GroupLens の方法はこうした人々のコミュニティで形成された「口コミ」を用いた推薦の過程を次の二段階で行う。

類似度の計算

利用者データベース中の各標本利用者と活動利用者の嗜好の類似度を求める。類似度とは利用者同士の何らかに対する嗜好の度合いを定量化した物である。

嗜好の予測

活動利用者がまだ評価していないアイテムについてそちらのアイテムへの標本利用者のこのみと、その標本利用者との間の類似度に基づいて、活動利用者がどれくらいそのアイテムを好むかを予測する。

GroupLens の利用者間型メモリベース法では評価値行列の行ベクトル、すなわちある利用者の様々なアイテムへの嗜好を表すベクトルの類似度に基づいて他の利用者の評価値を予測した。一方でアイテム間型メモリベース法では、評価値行列の列ベクトル、すなわちアイテムベクトルの類似度を用いる。これは同じような人から同じような評価を受けているアイテム同士は互いに類似度が高いと考え、ある利用者が感心を持っているアイテムの類似アイテムにも、その利用者は感心をもつという仮定に基づいている。

アイテム間型メモリベース法の簡単な方法として、アイテムベクトルのコサイン類似度や単純な共起性、Person 相関などでアイテム間の類似度を算出し、利用者が閲覧中

だったり、買い物かごに入れていたり、その商品を購入したり、直近にそのアイテムを閲覧しているようなアイテムと類似しているアイテムを推薦する。

§ 2.3 レビューの信頼性の判断支援

近年 Amazon や Yahoo!ショッピングなどの 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行っていると報告されている。

また、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商品に対しての購入意欲を下げることになる。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大きく注目されており、レビューの評価点や文章に対するデータマイニング、自然言語処理などを用いた研究も盛んに行われている。

しかしながら、これらの ec サイトにおいてそれらのレビューの価値を利用し、利益を目的としたサクラによ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。サクラとは商品や店の評判を不当に上げたり下げたりすることを目的として、その店の商品の購入を行わないにもかかわらず、レビュースパムという偽のレビューの投稿を行う悪意を持った利用者のことである。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スパムレビュー問題はしばしば問題に挙げられている。

レビュースパムは、そのレビューを読んだユーザーを騙し、誤解を招かせる可能性がある。このことからレビュースパムは最悪の場合、ユーザや店に実害を与える可能性がある。また、一部の ec サイトではレビューすること自体に値段をつけ、商売を行っていることもある。このように悪質なレビュースパムは利益だけでなく不当な損害をもたらす可能性があるものであり、スパムに対する対応は差し迫った問題である。

上記のような問題を解決するために、これまでにレビューの信頼性について様々な研究が行われてきた。しかしながら、各レビューがサクラによって投稿されたか否かを正確に判断することは不可能である。そのため多くの研究では、スパムが持つであろう特徴を検出し、各レビューに対するレビュースパムである可能（スパムらしさ）の算出を行っている。

上記のようにレビュースパムにはいくつかの特徴があり、その特徴からレビュースパムは以下のような3つのタイプが存在すると言われている。

Type1(untruthful opinions)

商品の評判を上げるまたは下げることを目的として、その商品を購入していないにもかかわらず、ポジティブまたはネガティブなレビューを不当に投稿することで他の利用者に対して誤解を与えてしまう可能性があるレビュー

Type2(review on brands only)

特定の商品に対するレビューを行うのではなく、製造者や販売元の利益や損失のためにそれらについて言及しているレビュー

Type3(non-review)

商品に対する感想ではなく、何の意味も含まないただの文章であったり、広告、ただの質問やランダムなテキストからなるレビュー

また、楽天市場における「みんなのレビュー・口コミ」を利用したスパムレビューの特徴の分析を行った研究が行われている。この研究ではスパムレビューには、「未成」「対応」「速い」「連絡」「メール」「ひどい」「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といった傾向があるとしている。逆に、スパムでない一般的なレビューについては、「香り」「コンパクト」「軽い」といった互換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わかった。

上記の3つのタイプにおいてType2とType3についてはその特徴からレビューに付属している文章から容易に判断できるが、Type1については文章からではそのスパムらしさを判断することが出来ない。そのため、Type1のようなレビュースパムを検出することが課題となっている。しかし、この問題を解決する上でそのレビューが本当にスパムであるのかどうかという問いに対する絶対的な答えが存在しないことである。

そこで複製されたレビューに着目してレビュースパムの検出をした研究が行われている。複製、または複製に近いレビューを収集し、レビューの文章を読むことでそれらがスパムとして妥当かどうかを人手により判断してもらい検証を行った。その結果、複製もしくは複製に近いレビューにはType2とType3の特徴が多く含まれていることがわかっていて、また、複製もしくは複製に近いレビューの中でType2,Type3に属さないものは以下の3つのものが多く含まれていた。

- 異なるユーザーIDで同じ商品に投稿された同じあるいは非常に似た本分のレビュー
- 同じユーザーIDで異なる商品に投稿された同じあるいは非常に似た本分のレビュー
- 異なるユーザーIDでこと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分のレビュー

これらは残りのType1である可能性が高いといえる。そのため、複製されたレビューは実際のレビューではなくサクラによって悪意の元で投稿されたスパムレビューだと判断することが妥当であると結論づけている。

また、サクラは協調的にスパムレビューを投稿することで商品に対して不当な評判を与えている可能性があるとし唆している。これは、複数のサクラが特定の店または商品について不当なレビューをすることでその店または商品の価値を自由に操作しようと、サクラが団結してグループでスパムレビューを投稿しているということである。

このようなサクラグループに対策を行うために、データマイニング分野における頻出アイテムセットの抽出の方法を用いて、サクラグループと思われるいくつかの投稿者グループを作成した。そして、その作成された投稿者グループを専門家にサクラグループかどうかを判断してもらった。その結果サクラグループであるかどうかの意見が複数の専門家で一致した。つまり、投稿者をいくつかのグループに分けることでサクラグループであるかどうかの判断が可能になるということである。そして、そのような投稿者のグループ分けは投稿者の投稿時間や投稿履歴、投稿されたレビュー間の文章の比較などから可能である。

つまりは利用者にとって単一のレビューをひとつひとつ判断してもそれはスパムレビューであるかどうかの判断は出来ないが、投稿者の投稿時間や投稿履歴、レビュー文章の比較をシステム側で行いそれらの結果をユーザーに示すことで、それらの情報から判断しスパムレビューの見分けができることが可能になると期待されている。

取得科目の推薦と教材の最適化

§ 3.1 協調フィルタリングからの科目の推薦

過去の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 (User Based Collaborative Filtering: UBCF) を適用し, 学生がまだ履修していない科目についての予測評価値を算出する.

ユーザーベース協調フィルタリング

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.

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

ただし, $\mathcal{Y}_{ax} = \mathcal{Y}_a \cap \mathcal{Y}_x$ 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一つ以下である場合, 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一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あいてむ 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 を評価済みである場合は $\hat{r}_{ay}$ の予測評価値は算出する必要はない. 式 (3.2) の第一項は, 第二項が中間的な評価を取る特徴を補正

するためのバイアス項であり、第二項は $\mathcal{R}_y$ 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 = 学生、アイテム = 学生が履修することのできる科目、評価値 = 各科目に対する学生の成績と置き換えて 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も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」という考えをして UBCF を実装する。

§ 3.2 シラバスからの教材作成

シラバスとは大学における講義内容や授業のスケジュールをまとめた資料である。シラバスには各教科ごとに開講学期、単位数、単位区分、教育目標、授業概要、学生の到達目標、授業計画授業に関連するキーワード、成績評価基準、授業の際の参考資料、履修条件が細かく記載されているので、学生がその授業の単位履修を考える際は、必ずと言っていいほど参考にする資料である。このシラバスは入学時に冊子として配布されるほかにも Web 上で閲覧が可能な Web シラバスが存在する。Web シラバスには検索機能が設けられており、開講年度、授業科目コード、授業科目名、学則科目名、担当教員、科目区分、配当学年、キーワードで検索が可能である。

本研究では上記で述べた Web シラバスからの情報をスクレイピングし、シラバスから得られた情報について再びスクレイピングを行い、その結果をユーザーである学生に提示する。このシラバスからユーザーに提示する情報を教材と呼び、学生はこの教材を利用し自身の成績の向上を図ってもらう。以下では教材作成までの流れの説明を行う。

まず、シラバスから各授業ごとの「授業計画」についてスクレイピングを行う。この授業計画にはそれぞれの授業の第一回目から第十五回目までの講義内容がどのようになっているかが書かれている。それらをスクレイピングし CSV ファイルに出力を行う。次に出力された CSV ファイルを読み込み、第一回目から第十五回目までの授業計画を検索ボックスに入力し、スクレイピングを行う。スクレイピングが完了すると授業回数ごとの関連したホームページの URL とホームページのタイトルが出力されるのでそれを教材として保存する。シラバスから出力された講義内容を図に示す。また、その講義内容からスクレイピングした内容の一部を示す。

シラバスからの教材作成の際には、Python での Web スクレイピングを行う。Python でスクレイピングを行うには Google Chrome や Firefox などのブラウザの操作を自動化できる Selenium と HTML や XML 内を解析して目的としているコンテンツを抽出できる BeautifulSoup4 を使用する。

Selenium

Selenium は元々、Web アプリケーションの UI テストや JavaScript のテストの目的

講義	
第一回	ネットワークの進展
第二回	デジタル伝送技術の基礎
第三回	ネットワークの階層化とローカルエリアネットワーク
第四回	ローカルエリアネットワーク 1
第五回	ローカルエリアネットワーク 2
第六回	ローカルエリアネットワーク 3
第七回	IPネットワーク 1
第八回	IPネットワーク 2
第九回	IPアドレス
第十回	サブネットアドレッシング
第十一回	IPルーティング
第十二回	ルーティングプロトコル
第十三回	TCPとUDP
第十四回	アプリケーション層
第十五回	まとめ

図 3.1: シラバスからの授業計画のスクレイピング結果

HPtitle	HPurl
(3) 情報化の進展	https://www.mlit.go.jp/hakusyo/transport/shouwa61/ind000501/003.html
第1章第1節1. 情報通信ネットワークにおけるデジタル化の進展	https://www.soumu.go.jp/johotsusintokei/whitepaper/ja/h10/html/98wp1-1-1.html
平成19年版 情報通信白書	https://www.soumu.go.jp/johotsusintokei/whitepaper/ja/h19/html/j1112000.html
情報社会の進展と 情報技術 - 文部科学省	https://www.mext.go.jp/content/20200609-mxt_jogai01-000007843_002.pdf
デジタル化・ネットワーク化の進展に対応した柔軟な権利制限 ...	https://www.bunka.go.jp/seisaku/chosakuken/1422075.html
ネットワークの発展	http://kjs.nagaokaut.ac.jp/mikami/MIS/network.htm
高度情報ネットワーク化	http://www.pref.tochigi.lg.jp/newvision/pdf/p034_p038.pdf
ネットワークの発展 - 木暮 仁	http://www.kogures.com/hitoshi/webtext/kj1-network-hatten/index.html
情報ネットワークの進展と地方ローカル鉄道の再生--銚子電鉄 ...	https://ci.nii.ac.jp/naid/110006866547
2. 知識社会・ネットワーク社会及びグローバル化の爆発的進展	https://www.mhlw.go.jp/shingi/2007/03/dl/s0314-15i-02.pdf

図 3.2: シラバスからの授業計画のスクレイピング結果

で開発されていたが、現在では、テスト以外にもタスクの自動化や Web サイトのクローリングなど様々な用途で利用されている。

Beautifulsoup4

Beautifulsoup に HTML を渡すと、その HTML 内をツリー構造で表現したオブジェクトを生成する。そのオブジェクトに対して自分が必要とする情報を検索して情報の抽出を行う。Beautifulsoup4 以前のバージョンは Python3 系には対応していないので今回は Python3 系に対応している BeautifulSoup4 を使用する。

§ 3.3 信頼性を考慮した教材の提示

協調フィルタリングにより予測評価値が算出された科目はシステムを使用するユーザーに対して推薦される。しかし、推薦された科目だけを履修していても GPA 向上にはあまりおおきな影響は及ぼさない。そこで本研究では 3.2 節において Web シラバスをもとに作成した教材を科目と一緒に推薦し、学生にそれらの教材で学習してもらうことによってさらなる GPA の向上を図る。

甲南大学で行われた、a. 教員に質問、b. ディスカッション、c. 予習。復習、d. 発表の準備、e. 卒業論文、f. 期末テスト、g. ノートに書き写す、h. 授業中の私語、i. 遅刻・欠席の 9 つ

の要因のうちどの変数が GPA に影響を与えているかという実験がある．その結果，ダミーデータのモデル 3 において c と g 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・復習は GPA 向上につながるので，科目を推薦すると同時に教材を推薦することは大いに GPA に影響があるといえる．

レビューにおける信頼性をは 2.3 節で述べたように，レビューの特徴からそのレビュー自体の信頼性を判断することができる．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 つの指標を定義する．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 に近いほどスパムらしさが高く，スパムと疑われるものとする．

類似性

複製された，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がある []．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図る指標として類似性スコアを定義する．まずレビュー r_i の文章を Mecab を利用して形態素解析を行う．これによって区切られた単位要素の集合をレビュー r_i を表す要素集合 X_{r_i} とする．次に Jaccard 係数を用いてレビュー r_i と r_j の類似度を式 (3.3) で求める．

$$sim(r_i, r_j) = \frac{|X_{r_i} \cap X_{r_j}|}{|X_{r_i} \cup X_{r_j}|} \quad (3.3)$$

このとき， $|X_{r_i} \cap X_{r_j}|$ は X_{r_i} と X_{r_j} のどちらにも存在する要素数， $|X_{r_i} \cup X_{r_j}|$ は X_{r_i} または X_{r_j} に存在する要素数である．そして，レビュー r_i の類似性のスコアを式 (3.4) のように求める．

$$S_score_{norm}(r_i) = \max_{r_j} (sim(r_i, r_j) | j \neq i, j = 1, 2, \dots, n) \quad (3.4)$$

このとき， n は r_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 (3.5) の方法で類似性スコアを正規化を行う．

$$S_score_{norm}(r_i) = 5 \cdot S_score(r_i) \quad (3.5)$$

協調性

2.3 節で述べたように，レビュースパムを集団で登校することで商品などの評判を不当に操作するサクラグループが存在する．これは同じグループのメンバが同じ商品にたいして投稿を行い，協力して評判を変えるものである．そこでレビューがサクラグループによって投稿された物である可能性を図る指標として協調性スコアを定義する．まず頻出アイテムセット抽出の方法を用いる． t_p 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，トランザクションと呼ぶ．また，ある投稿者グループが出現した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 の協調性度を式 (3.6) で定義する．

$$collaborate(g_c) = support(g_c) \cdot size(g_c) \quad (3.6)$$

そして、レビュー r_i の協調性スコアを式 (3.7) で求める。

$$C_score(r_i) = \begin{cases} In(max_{g_c \in G_{u_{r_i}}}(collaborate(g_c))) & |G_{u_{r_i}}| \neq \emptyset \\ 0 & |G_{u_{r_i}}| = \emptyset \end{cases} \quad (3.7)$$

このとき u_{r_i} は r_i を投稿した投稿者であり、 $G_{u_{r_i}}$ は u_{r_i} が属する頻出投稿者グループの集合である。さらに、式 (3.8) で協調性スコアを正規化する。

$$S_score_{norm}(r_i) = \frac{5 \cdot C_score(r_i)}{\max(C_score(r_i) | j = 1, 2, \dots, N)} \quad (3.8)$$

このとき N は全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

レビュースパムは時間的に集中して投稿されることがわかっている。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に添えて投稿される5段階の値である。

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。これは、時系列データに対してイベント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、特定の「単語」を含んだ Web ページの投稿が急激に増えることがある。このような現象をバーストと呼び、蒸気の手法はこのような現象を検出する際に用いられる。この手法の「単語」を「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、そのときの評価値を持つレビューの集中性スコアとする。

ある店の m 日目のレビュー集合を B_m そち、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 日目のレビュー集合に対して、バースト検知手法を用いて、普段よりも評価値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 $r = \{r_1, r_2, \dots, r_{u+1}\}$ を考える。そして、 r_i と r_{j+1} の投稿時間間隔を x_i としたとき、 r の投稿時間間隔列 $x = \{x_1, x_2, \dots, x_u\}$ を求める。そして、投稿時間間隔列 x に対して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 $r_i \in g_b$ の集中性スコアは、 g_b のレビューの数 $size(g_b)$ を用いて式 (3.9) で求める。

$$T_score(r_i) = In(size(g_b)) \quad (3.9)$$

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、式 (3.10) で正規化を行う。

$$T_score_{norm}(r_i) = \frac{5 \cdot T_score(r_i)}{\max(T_score(r_i) | j = 1, 2, \dots, N)} \quad (3.10)$$

提案手法

§ 4.1 予測評価値からの適切な授業科目の推薦

3.1 章での UBCF の予測評価値をもとに学生に対して科目の推薦を行う。本研究の目的は学生になるべく高い GPA を取得してもらい、就職・進学に有利を進めてもらうというもので、予測評価値 (予測成績の値) が高い科目を優先して学生に推薦を行う。具体的には、予測成績の値が 3 以上、つまりレターグレードが A 以上と予測される科目を優先する。

大学において卒業要件単位というものが存在する。卒業要件単位とはその名の通り卒業するのに必要な単位数のことである。これらの卒業要件単位を満たさないとたとえどれだけ多くの単位を取得していても卒業することはできない。富山県立大学もその例外はなく、卒業要件単位が存在し卒業要件単位を満たさなければ卒業できない。

しかし、卒業要件単位を満たしていても卒業できない場合がある。それは必修科目を取得できていないときである。必修科目とは必ず単位を取得しなければいけない科目のことである。卒業要件単位を満たしていてもこの必修科目を取得できていなければ卒業ということにはならない。卒業するための取得単位の制約は他にもあり、選択必修科目と呼ばれるものがある。選択必修科目は、その単位自他を必ず取得する必要はないが、指定された単位のうちいくつかの単位を取得する必要がある単位である。もちろんこの選択必修単位が足りなければ卒業には届かない。

このように大学には、卒業を考えた際に注意すべき履修条件が膨大であり、予測成績が高い単位だけを取得しているだけでは卒業できない。そこで本研究ではこれらの卒業要件単位、必修科目、選択必修科目のすべてを考慮した上で予測評価値が最も高くなるように学生に科目を推薦する。つまりはここでいう適切とは、「予測成績が高く、卒業要件単位、必修科目、選択必修科目のすべてを満たす」ということである。

§ 4.2 関連資料のアップデート

3.3 章で述べたように科目の推薦と同時に 3.2 章で作成した教材を提供する。しかし、3.2 章で作成した教材では Web ページで検索された上位の結果がそのまま、おすすめの教材として学生に表示されてしまう。そのままの状態が表示されるとたとえ全く科目と関係の無い内容が混ざっていてもそれがおすすめの科目として推薦されてしまう。そこで学生の評価によってためになる教材とそうでない教材がわかるようにするシステムを実装する。

§ 4.3 提案手法のアルゴリズム

Step1 教学データと教材データの蓄積

学生がシステムを使用する前に、学生×科目の行列を数年分、あらかじめ蓄積しておきこれらをデータベースをとして扱う。しかし、2.1 節でも述べたように学生の教学データは個人情報として扱われるので、デモデータを作成し CSV ファイルに保存したものをデータベースとする。また、3.2 節で述べた教材もあらかじめ作成しておき CSV ファイルで保存しておく。

Step2 協調フィルタリングによる予測成績の予測

学生にはシステムを使用する際に学籍番号を入力してもらう。そこで指定された学籍番号の科目情報を抽出する。抽出した学生データを既に卒業した学生データと結合し、指定した学籍番号をターゲットとして Pearson 相関で卒業済みの学生との類似度を計算する。類似度が計算できたら類似度が高い順にソートを行い、類似度の高い上位 10 名を抽出する。抽出した 10 名とターゲットにした学生の科目の成績の加重平均を求めることで予測成績を導出する。

Step3 適切な科目の選択と推薦

すべての科目について予測成績が導出されたら、予測評価値が高い順にソートを行う。次に必修科目を推薦科目一覧に追加する。続いて、選択必修単位数を満たすように選択必修科目の中で予測評価値が高いものから推薦科目一覧に追加していく。最後に卒業要件単位を満たすようにまだ推薦科目一覧に入っていない科目の中で予測評価値が高いものから推薦科目一覧追加を行い、最後に推薦科目一覧を HTML に表示し学生が閲覧できるようにする。この時、学生には卒業研究に集中してもらうべく、なるべくすべての単位を 3 年の後期までに推薦する。4 年の後期以降は落単した単位のみを自発的に履修してもらう。

Step4 推薦科目の教材提供

Step3 で推薦した科目に関する教材をあ提供する。提供する教材は Web ページと Youtube の動画の 2 種類である。推薦した科目名をクリックするとその科目の教材を見ることが可能である。教材は、講義の第一回から第十五回まであり、それぞれの回ごとに Web ページと Youtube の動画の 5 つずつを HTML に表示する。

Step5 教材の更新

数値実験並びに考察

§ 5.1 数値実験の概要

本研究の協調フィルタリングによる最適な科目の推薦が行えているかの確認を行う。協調フィルタリングで情報推薦を行ってもパーソナライゼーションができておらず、同じ内容ばかりを推薦しては推薦システムとして意味のないものになってしまう

そこで一つ目の実験として、まったく同じ科目を履修している学生二名について実験を行う。学生AとBは1年の後期の講義までをまったく同じ単位を履修している。しかし、二人の成績はそれぞれ異なり、取得した単位は同じであるが成績が異なるといった状況を考える。この時の学生AとBに対する推薦科目を確認し推薦内容がしっかりとパーソナライゼーションできているかを確認する。

実験方法は単純なものであり、取得成績が同じであるが成績の異なる二人の学生のデータを準備する。それらについて協調フィルタリングを行い、推薦科目を決定する。その推薦された科目を確認し二人の推薦結果がどれくらいの割合で被覆しているかを5回測定し、その平均を算出し確認を行う。

§ 5.2 実験結果と考察

結果を図に示す。結果としては2年前期の推薦科目の被覆率は53%、2年後期は68%、3年前期は50%、3年後期は52%という結果になった。

考察としてはほとんどの学期で被覆率は50%となっており、全体で見ても56%ほどの被覆率となっているので、たとえ同じ科目を履修したとしてもしっかりとパーソナライゼーションができている結果となった。

おわりに

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工学部電子・情報工学科情報基盤工学講座の奥原浩之教授，António Oliveira Nzinga René 講師に深甚な謝意を表します．最後になりましたが，多大な協力をしていただいた研究室の同輩諸氏に感謝致します．

2022 年 2 月

滝沢 光介

参考文献

- [1] “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
- [2] 鶴田美保子, “大学生の就職活動を成功させる要因”, 金城学院大学論集 人文科学編, 第 15 巻第 1 号, pp. 109-119, Sept 2018.
- [3] 畔津憲司, “九州市立大学経済学部 2012 年度卒業生の卒業後進路及び就職活動実態等に関する調査報告”, 北九州市立大学『商経論集』, 第 49 巻第 1・2 号, pp. 75-120, Dec 2013.
- [4] 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.
- [5] MARCERON, A., BLIKSTEIN, P. and SIEMENS, G. “Learning Analytics: From Big Data to Meaningful Data”, Journal of Learning Analytics, 2(3), pp. 4-8, 2015.
- [6] 稲葉利江子, 山肩洋子, 大山牧子, 村上正行, “発言の自由度を高めたレスポンスアナライザを活用した大学授業の実践と評価”, 日本教育工学会論文誌, 36(3), pp. 271-279, 2012.

