

Media Data Tech Studio

秋葉原ラボ技術報告

4

volume.





Media Data Tech Studio

巻頭言「10年後の8月」

秋葉原ラボ 研究室長 福田 一郎

秋葉原ラボ技術報告 vol.4



ガールズロックバンド・ZONE^{*1}の代表曲「secret base ～君がくれたもの～」の歌詞に『10年後の8月』という印象的な一節がある。本格的な夏を前に本稿をしたためている現在、2011年4月に開設した秋葉原ラボはその『10年後の8月』を迎えようとしている。10年前の2011年は東日本大震災があり、その10年後の今年は新型コロナウイルスの影響が続き、日常生活にはまだ大きな制限がある状態となっている。いずれの年も後世に激動の1年と呼ばれるであろう1年となっている。そんな中、節目の10年を迎えた秋葉原ラボも組織として大きな変革のときを迎えようとしている。秋葉原ラボは改組を予定しており、この技術報告も秋葉原ラボとしては最後の技術報告となるかもしれない。



秋葉原ラボも前身組織から発展する形で成立したので、今回もまた発展を見込んでの改組ではある。本社のある渋谷ではなく秋葉原という地で、ある意味エンジニアの「秘密基地」のような形でやってきた組織も10年経つと時流に合わせて変化しないといけない部分は出てくる。開設した当初は「ビッグデータ」という言葉が流行り始めた時期であったが、いまはデータを活用したサービス提供が当たり前になりつつある。当時は大量のログデータを扱うシステム自体に新奇性があったが、いまやパブリッククラウドのフルマネージドサービスとして提供される時代になっている。それゆえに、これまで研究開発や基盤システム整備に重きを置いていた軸足をやや事業やサービス寄りにすることが必要になってきたわけである。

しかし、ここで研究開発や基盤システム整備の機能を完全になくしてしまうわけではない。これらの機能は、サービス自体の開発組織に比べてその存在が周囲の人々に意識されることは少ないが、組織として持つことによって技術基盤が維持されることになり、それが提供するサービスを発展させるために必要となってくるのだ。

冒頭で紹介したZONEの曲「ボクの側に…」の歌詞に次のようなフレーズがある。

『失って気づいた その大切さを』

失ってからでは遅い、多くの人には気付かれていなくても維持していかなければならない技術がある。その信念をもって今後も技術に対して真摯に取り組んでいきたい。

そんな信念の一端をこの技術報告から読み取っていただければ幸いだ。

^{*1} [https://ja.wikipedia.org/wiki/ZONE_\(%E3%83%90%E3%83%B3%E3%83%89\)](https://ja.wikipedia.org/wiki/ZONE_(%E3%83%90%E3%83%B3%E3%83%89))

目次

巻頭言「10年後の8月」	i
1 管理会計における多種多様なシステムリソースの包括的な収集・管理	1
2 組織間の連携を実現するためのワークフローエンジン	7
3 ABEMA における動画特徴抽出と推薦への応用	15
4 類似画像検索システムの設計と運用	24
5 ピグパーティにおける Orion フィルタの候補キーワード抽出	30
6 国際的質問紙調査の実施に関するテクニカルノート	37
発表一覧	43



Media Data Tech Studio

1

管理会計における多種多様なシステムリソースの包括的な収集・管理

執筆者 飯島 賢志

概要 秋葉原ラボが所属する技術本部では、機械学習を支えるデータ基盤、コンテナ基盤、課金基盤など、またそれらの基盤上で構築されたシステムなども含み 50 以上のシステムが運用されている。現在、技術本部では管理会計のしくみの見直しに取り組んでおり、技術本部で運用するシステムで発生した費用を各事業部や子会社に請求する際、各システムのリソース使用量に基づいた請求額になるように管理会計システムを構築した。本稿では、この管理会計システムの一部であり、多種多様なシステム間のリソース使用量の収集・管理を担う DFAS について紹介する。

Keywords リソース計測，管理会計，配賦基準

1 背景

近年、データに基づく仮説検証や意思決定など各国の企業・組織でデータの活用が活発化している。秋葉原ラボが所属する技術本部においても、データの活用を、事業への貢献だけでなく、組織の経営判断や業務の効率化の面でも推し進めている。技術本部では以前から管理会計のしくみの改善に取り組んでおり、各システムのリソース使用量などのデータに基づきシステム費用などを各事業部に按分する配賦方法を採用していたが、しくみがシステム化されていないことでさらなる抜本的な改善が難しい状況であった。

その背景として、まず技術本部は横軸組織であり、社内の事業に加えてグループ会社からも利用されるさまざまな社内システムを開発している。それからこれらの社内システムで生じたインフラコストや開発費などの費用は配賦基準にのっとり、利用者が社内であれば利用者の事業に配賦し、利用者がグループ会社であれば利用者のグループ会社に請求することになる。しかし、システムの利

用者はこれらだけでなく、ほかの社内システムの場合や、分析業務に携わっている社員および業務委託の受託者などの個人の場合もあり得る。また、あるシステムが別のシステムを経由して事業に利用されるなど利用状況の関係性は複雑であり、人の手作業で表現するのは容易ではなく配賦金額や請求金額を正確に計算することは難しい。

このような問題を解決するために技術本部ではしくみのシステム化に着手し管理会計システムを構築した。本稿では、この管理会計システムの一部であり、システム間のリソース使用量の収集・管理を担う DFAS について紹介する。なお、一般的に会計の用語で、配賦とは発生した費用をある一定の基準でいくつかの部署などに割り振ることをいう。按分とは基準となる数量に比例して割り振る配賦方法を指す。本管理会計システムでは部門やグループ会社などの事業単位を社内外を問わず統一的にとらえるために、社内で生じた費用の割り当てである「配賦」についても、事業単位に対する「請求」と表現する。

2 解決すべき課題

2.1 汎用的なシステム使用量の表現

前述したように、按分の配賦方法を採用することで各事業に配賦するシステム費用を算出していた。しかし、たとえ既存の事業側のシステム利用状況が変わらなくても、新しい事業が増えればその新しい事業にもシステム費用を按分するため、既存の事業への請求額が下がる、またその反対の事象も按分の性質上起こり得る。このように、按分の配賦方法では事業数の変動という外部要因で按分比率が変わり請求額も変動してしまうため、各事業が中長期的に事業計画を立てる上でシステム費用がどれくらいになるかの見積もりにはあまり適さない。そこで、事業の拡大に伴うシステム利用状況に応じてシステム費用を見積もるためには、配賦方法を変更して各事業が利用者としてシステムを使った使用量を元にして、あらかじめ設定したシステム単価を乗じて請求額を決定する方法が考えられる。1つのシステムがある事業単位に対して請求する金額は下記のように表すことができる。

$$\text{請求額} = \sum (\text{システム使用量}_i \times \text{単価}_i)$$

※ i : システム使用量のメトリクス

この場合、システム使用量としてどのようなメトリクスを設定すべきかは各システムの特性により異なるため、各システムごとに異なるさまざまなシステム使用量を扱える必要がある

2.2 さまざまな利用者への対応

システム使用量の利用者を特定して費用を請求する請求先の紐付けも表現できることが必要となる。利用者の特定に対象システムの認証機能を利用できる場合もあるが、そうでない場合は本システムで使用量と利用者を紐付ける必要がある。また、利用者が事業である場合やシステムである場合、任意の部署に所属する個人である場合もあり、これらさまざまな利用者を表現できる構造が必要である。

2.3 多段請求の考慮

システムはほかのシステムの利用者ともなりうる。たとえば、システムがAPIなどの形でほかのシステムを利用している場合などがある。このようにシステム間の利用があるとき、システムを利用するとそのシステムを経由して別のシステムを同時に利用することになる。つまり、システム使用量に応じたシステム費用を正確に算出するためには、事業単位に属する利用者だけでなく、システム側から見た利用者であるほかのシステムへの請求も考慮した多段請求を表現できることが必要になる。図 1.1 に多段請求の例を示す。多段請求を通じた各システムの請求は最終的に各事業単位に対して行うことになる。

3 DFAS の設計

3.1 概念モデル

前節で述べたように、DFAS では汎用的にシステム使用量を表現して、さまざまな利用者に対応し、多段請求を考慮できる設計が必要である。DFAS の設計の概要として概念モデルを図 1.2 に示す。

DFAS では利用者があるシステムを使う事象を利用タイプとその使用量、およびリソース属性で表現できる。利用タイプとは、API リクエスト数や保存バイト数など CPU やメモリ、ディスクなどのシステムリソースに関係するメトリクスのタイプである。使用量は各利用タイプのリソースがどれだけ使われたかの数値を示す。リソース属性には、使われた利用タイプのシステムリソースなどに関する情報を含めることができる。API のリクエストパス、保存したオブジェクトのネームスペース、利用したテーブル名などが例としてあげられる。このようなリソース属性の情報からどのアカウントに費用を請求すべきかの紐付けは後述する `resource_to_account_rule` で表現される。

また、各事業単位や各システム、各個人はアカウントとして表現され、それらは利用者にも請求先にもなり得る。たとえばある事業単位に属する個人がシステムを利用したときは、その個人が利用

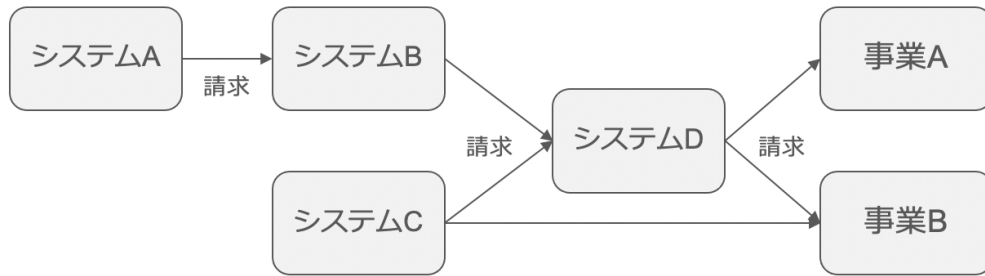


図 1.1 多段請求の例

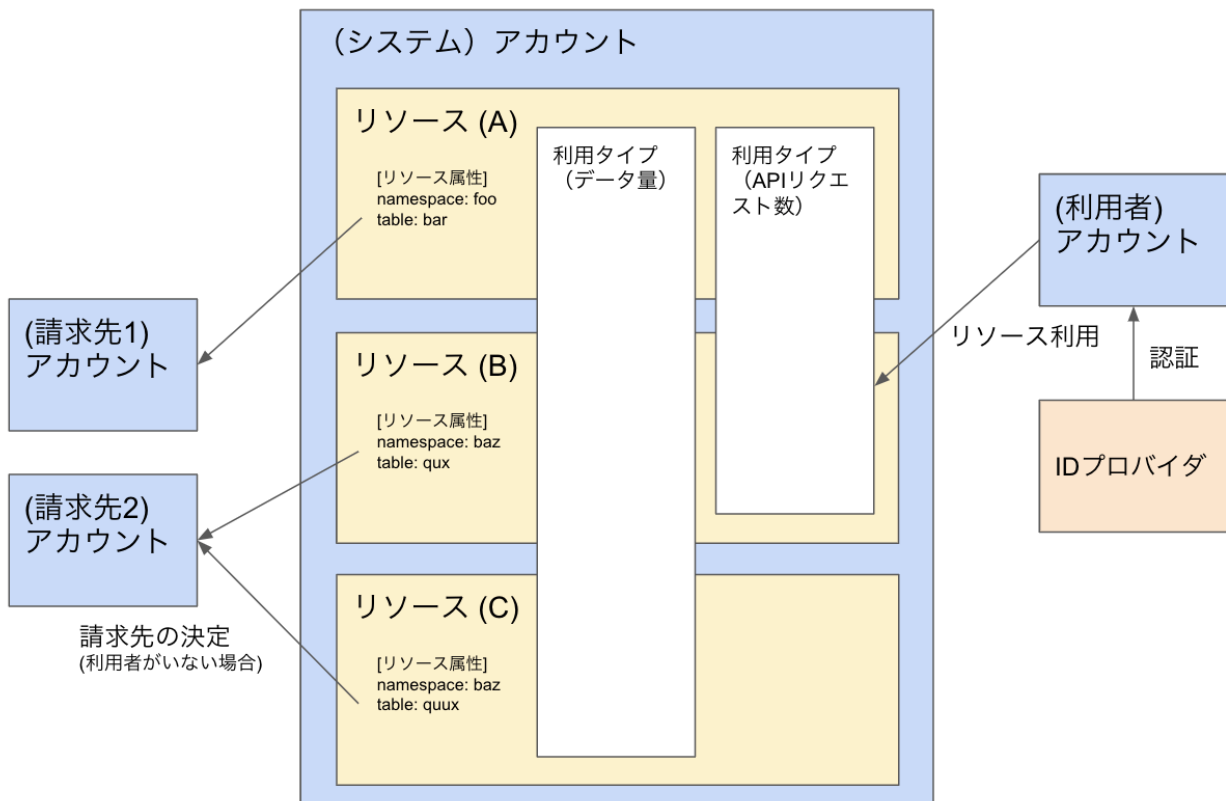


図 1.2 DFAS の概念モデル

者であり、所属する事業単位が請求先となる。一方、多段請求が発生する場合には、システムを利用したほかのシステムが利用者や請求先となることもある。それらを一貫して扱うために、本システムでは利用者や請求先となり得る主体をまとめてアカウントと表現することにした。

ID プロバイダの認証、またシステム自身が認証機能を持つ場合は自身の認証機能を通じて利用者を特定できる。もしシステムが認証機能を持たず、利用者を特定できない場合、使用されたリソース属性から請求先を結び付けることになる。

3.2 データ構造

図 1.3 に DFAS のデータ構造の一部を示す。

利用者や請求先を表す `account` は 1 つの `id_provider` に所属し、`oidc_issuer` (Issuer Identifier) と `provider_user_id` (Subject Identifier) の組として一意に表現できる。

`resource_to_account_rule` では、JSON で表現されるリソース属性が `match_expression` (JSON) の各要素を含むとき、指定された請求先 (Billing Account) に使用量を紐付ける。`match_order` は複数の

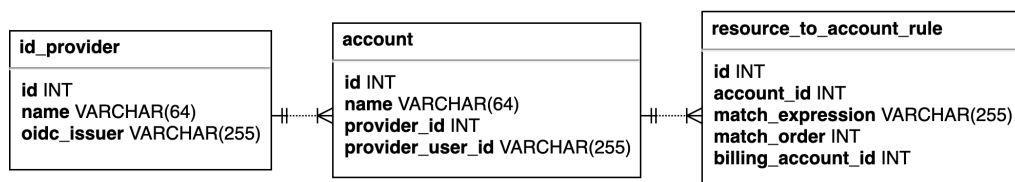


図 1.3 DFAS のデータ構造の一部

resource_to_account_rule に一致する場合に優先する順位である。認証がなく利用者を特定できない場合に対して、このように match_expression と match_order を組み合わせたルールにすることで、柔軟に請求先アカウントを紐付けることができる。

4 DFAS の実装

4.1 管理会計システム

まず、DFAS の位置付けを明確にするために管理会計システムの全体像を図 1.4 に示す。現在、管理会計システムは大きく 3 つのコンポーネントから成り立っている。1 つ目は個人やシステム間の使用量を収集する DFAS、2 つ目に DFAS で集めたシステム間の使用量データを取り込み、そのほかのインフラ費用、人件費なども含めて技術本部内のあらゆる費用を統合的に管理する Macys、最後に Macys から連携された各システムの費用を可視化するダッシュボードの役割を持つ Kittyhawk の 3 つである。Macys では各システム毎に費用を紐付け、収益は DFAS で集めたシステム間の使用量データとリソースの単価を掛け合わせることで算出し、収支を計算する。各システムは多段請求を通じて最終的に各事業単位に対して費用を請求することになる。このように、DFAS は各システムの収支、各事業に請求する金額を計算する上で必要となるシステム間の使用量データを作る役割を果たしている。

4.2 DFAS のシステム概要

システム概要として DFAS のデータの流れを図 1.5 に示す。連携するシステムについては後述する。

まず、DFAS は MetricConverter を通じて、各システムの監視システムに対してシステム使用量に関連するデータを収集し、システム使用量ログ（図 1.6）に変換したうえでデータ基盤である Patriot に保存する。一方、各システムを一意に認識するための情報を ID プロバイダから取得するとともに、DFAS Web でシステム使用量を示すログと請求先を紐付ける情報 resource_to_account_rule を各システム担当者が入力する。この情報とログを掛け合わせて集計したシステム間の使用量データを API を通じて統合会計システムである Macys に受け渡す。これが DFAS が担う一連の機能である。

4.3 連携システム

DFAS はほかのシステムと連携し、それらを有機的に活用することで DFAS が成すべき機能を実現している。ここでは DFAS と連携しているシステムを紹介する。

ID プロバイダ

システムまたは個人を管理するために、社内の ID プロバイダを利用している。社内の複数システムへのシングルサインオンとして利用もでき、社内の正社員や業務委託の受託者などの個人を一意に表現できる ID プロバイダがあり、DFAS では個人を管理する ID プロバイダとして利用している。また、システムを一意に表現できる ID プロバイダとしてサービスアカウント認証基盤である sa-iam を利用している。ほかにも ID プロバイダ機能を有しているシステムであれば、ID プロバイダの 1 つとして扱うこともできる。これらの ID プロバイダを通じて、利用者または費用の請求先としてのアカウント、システムリソースを消費されたシス

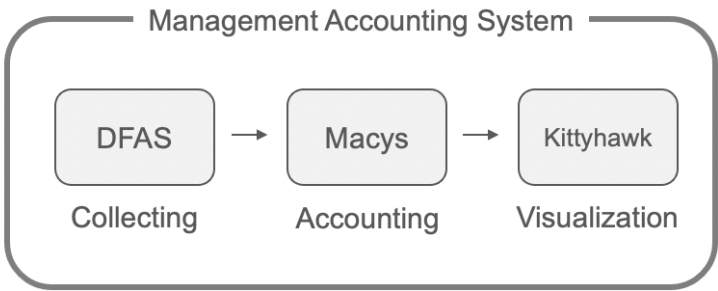


図 1.4 管理会計システムの全体像

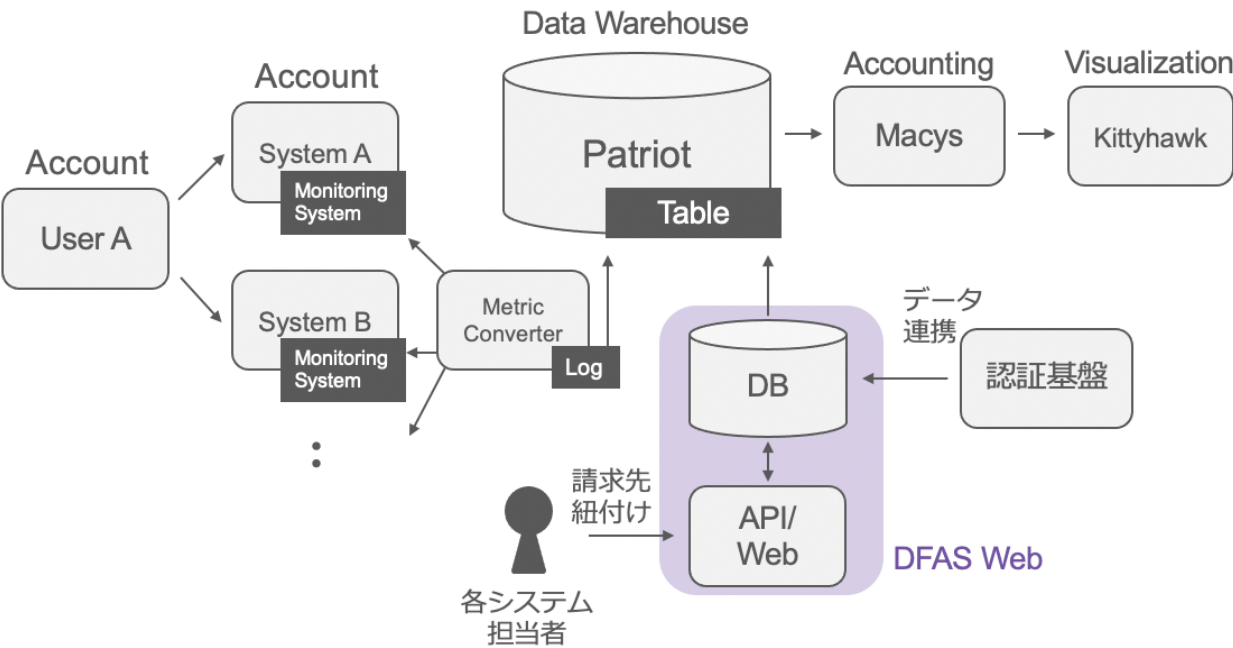


図 1.5 DFAS のデータの流れ (仮)

収集したログの例

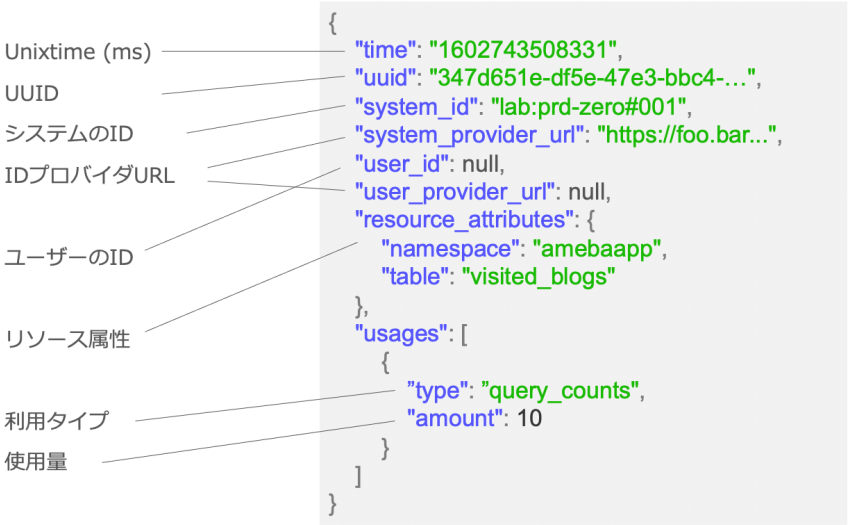


図 1.6 システム使用量ログの例

テムとしてのアカウントを表現できるようになっている。

監視システム

各システムは時系列データベース機能を持つ監視システムに自身の使用量を示すメトリクスを反映する。DFAS は監視システムの API のエンドポイントにアクセスし、データを収集しシステム使用量ログに変換する。監視システムはシステム毎に異なるものであっても対応できるようにプラグブルで柔軟なしくみになっている。

Patriot

DFAS が収集した使用量に関するデータは秋葉原ラボのデータ基盤である Patriot に格納する。また、DFAS がデータを収集するとき、データを保存するときの処理は Patriot が持つ Patriot Workflow Scheduler で管理され実行される。さらに、Patriot からデータを取り出し処理をする FDC API があり、DFAS から Macys にデータを受け渡すとき、この API 機能を使用している。

5 まとめと課題

本稿では、技術本部の管理会計システムの一部であり、システム間のリソース使用量の収集・管理を担う DFAS について紹介した。DFAS を活用することで会計管理システム全体として各システムの使用量に応じた請求・管理が実現できる。

今後の DFAS および管理会計システム全体としての課題として、管理会計システムを活用することで管理会計のしくみの改善、各システムの改善にいかにつなげられるかがあげられる。リソース使用量の単位あたりの金額として単価を定めることで、各システムの相対的な価値を表現できるようになったが、社内の代替システムとの比較や外部サービスとの比較を通じて、システム自体の運用コスト削減の取り組み、機能追加や抜本的な改修など付加価値を上げる取り組みが組織として継続的に実行できるかが重要である。また、本稿の執筆段階では、まだ本番運用が開始されていないため、まずは継続的に管理会計システムが適切に使われるために障壁になっている要素が本番運用開始後にあるようであれば、その解消の優先度が高いと考えられる。まずは運用が軌道に乗るように地道な対応が今後の第一歩となる。



Media Data Tech Studio

2

組織間の連携を実現するためのワークフローエンジン

執筆者 岩城 洋一

概要 秋葉原ラボでは各サービスの定期処理、機械学習モデル作成などのワークフロー実行を支援するための基盤を提供している。機械学習モデルの作成のためには、データの収集や前処理などのさまざまなデータに関連する処理を連携して動かす必要があり、組織の規模が大きくなるほどそれらの処理は各チームまたは各組織によって開発・運用される。このため近年は組織間のさまざまな業務ワークフローを連携させる重要性が高まっている。本稿では組織間の効率的なワークフロー連携を実現することを目的とした現在開発中のワークフローエンジンについて紹介する。

Keywords ワークフローエンジン、バッチ処理、Kubernetes、Operator

1 はじめに

秋葉原ラボでは機械学習のモデル作成を含む各種ワークフローが定期的に実行されている。これらのワークフローは秋葉原ラボ内のシステムやデータと連携するだけでなく、他組織のシステムやデータと連携するものもある。当社では社内におけるデータ有効活用の促進を進めており、他組織とのシステムやデータによる効率的な連携はワークフローにおいても重要性が高まっている。本稿では秋葉原ラボにおける、組織間の効率的なワークフロー連携を実現するためのワークフローエンジン開発の取り組みについて紹介する。

1.1 ワークフローの用語に関する補足

本稿ではワークフローとは「システムやサービスの一連の処理を行うために構成される1つ以上のジョブ／タスクの集合体」を指す。たとえば機械学習モデルを作成する場合は、モデルを作成す

るための一連の流れがワークフローであり、ワークフローの中で実施されるデータ準備、前処理、学習処理などが個々のジョブ／タスクである。

2 背景と課題

秋葉原ラボでは各ワークフローをワーカの OS／ライブラリ環境に依存しない形でワークフローを実行するために、オープンソースのワークフローエンジンである Apache Airflow^{*1}をベースとしたワークフロー実行基盤をラボ向け Kubernetes 環境^{*2}上で運用している。またこのワークフロー実行基盤では秋葉原ラボ内の要件を満たしたり、ワークフロー利用者の学習コストやワークフロー作成コストを軽減するために Airflow のコードベース等に対して機能追加や改修を行っている。

ワークフロー実行基盤によりライブラリ等に依存しないワークフローの管理・運用という当初の目標は達成しているが、以下のような課題がある。

^{*1} <https://airflow.apache.org/>

^{*2} https://d2utiq8et4vl56.cloudfront.net/files/topics/25465_ext_22_0.pdf?v=1607324285

1. 依存コンポーネントが多い：Airflow では Web, Scheduler, Worker というコンポーネントで構成されており、またそれらのコンポーネントが連携するために Redis, データベースも必要となるため依存コンポーネントが多くなる。運用環境では各コンポーネントの冗長化も行っているため、これらを含めた運用コストが高くなる傾向にある。
2. 機能拡張が困難：Airflow をベースとしているため、ワークフロー利用者からの Airflow にない機能があった場合に Airflow の内部実装を理解したうえで設計と実装が必要となる。このため機能追加や機能改善時には Airflow 側のソースコード、データベーススキーマを調査したうえでソースコード改修を行う必要がある。
3. ワークフロー連携のしくみが存在しない：ワークフロー連携とは「Amazon S3^{*3}の特定バケットの配下に新しいオブジェクトが追加された」といったような別ワークフローの実行結果を検知してワークフローを実行するしくみである。Airflow は定期実行によるワークフロー実行には対応しているが、ワークフロー連携のしくみは持たない。

特に3つ目の課題は、機械学習モデル作成などにおける組織間のワークフロー連携の重要性が高まっていること、秋葉原ラボ以外の組織で運用されているバッチ処理システムやワークフローエンジンでもワークフロー連携が課題となっていることも踏まえると、今後のサービス開発や改善の競争力を高めるためにも解決することが不可欠な課題である。

2.1 既存の代替手段の検討

前述した課題を解決するための代替手段の候補として、Kubernetes 上で動作し、ワークフロー定義を宣言的に記述可能なワークフローエンジンである Argo Workflows^{*4}、Tekton^{*5}を検討した。

Argo Workflows と Tekton は Kubernetes Operator ^{*6} として動作するため、ワークフローの管理と制御が Operator に集約されて依存コンポーネントが少なくなるので前節の1つ目の課題を解決できる。しかし、2つ目の課題については機能を拡張するためのプラグイン機能等は提供されていないため解決することはできず、3つ目の課題についても複数のワークフロー連携を実現するためのしくみが提供されていないため解決できない。以下に3つ目の課題を解決できない詳細な理由を挙げる。

- 標準ではワークフロー連携をサポートしない：Argo Workflows と Tekton はともにワークフローとその中のタスクを定義、実行、制御するために必要な Kubernetes CRD (Custom Resource Definition)^{*7} などを提供するが、複数のワークフロー間の連携をサポートするしくみが提供されていない。
- ワークフロー連携を実現する拡張のしくみが複雑：Argo Workflows では Argo Event^{*8}、Tekton では Tekton Trigger^{*9} というイベント検知関連の拡張コンポーネントがあり、これらのアドオンを用いることで様々なデータソースのデータ変化を検知して組織間のワークフロー連携を実現することはできる。ただしこれらの拡張コンポーネントは新たな Kubernetes CRD およびイベントルーティング等のしくみといった複雑性を導入することになるため、ワークフロー利用者がワークフロー連携を実現するため

^{*3} <https://aws.amazon.com/jp/s3/>

^{*4} <https://argoproj.github.io/argo-workflows/>

^{*5} <https://tekton.dev/>

^{*6} <https://kubernetes.io/docs/concepts/extend-kubernetes/operator/>

^{*7} <https://kubernetes.io/docs/concepts/extend-kubernetes/api-extension/custom-resources/>

^{*8} <https://argoproj.github.io/argo-events/>

^{*9} <https://github.com/tektoncd/triggers>

の学習コストがより高くなる。

- ワークフロー連携の管理・運用のしくみがない：前述した Argo Event や Tekton Trigger は基本的にワークフロー実行等のアクションを実施するまでが責務であるため、それ以降の対象ワークフローの実行結果や履歴などを管理するためのしくみは提供されていない。このためこの部分の管理・運用はワークフロー利用者に委ねられることになるため運用の負担が大きくなる。

このように Argo Workflows や Tekton ではイベント検知関連の拡張を用いることでワークフロー連携自体は実現することは可能となるが、システムの依存コンポーネントを増やし、かつワークフロー利用者の学習コストや運用負荷の増大を招くことになるため代替手段としては採用に至らなかった。Argo Workflows と Tekton の課題を解決するための Issue や Pull Request を作成するアプローチも考えられるが、どちらも比較的規模の大きい OSS プロジェクトであるためさまざまな理由により Pull Request がマージされない可能性もあるため見送っている。機能面でワークフロー連携を実現しつつ、管理・運用面でワークフロー利用者の負担を最小限にすることを両立するため、次節で述べるワークフローエンジン開発を行うこととなった。

3 ワークフローエンジンの概要

秋葉原ラボでは上述した従来のワークフローエンジンの課題を解決するために、Wurfrahmen と呼ばれるワークフローエンジンを開発している。Wurfrahmen は既存ワークフローエンジンに不足している機能やしくみを拡張機能として補完することにより各種課題を解決することを目的としたワークフローエンジンである。秋葉原ラボでは各サービスの品質向上などを促すための社内向けシステムを提供しており、各組織のサービスやシステムのワークフローと連携することは極めて重要であるため、特に複数のワークフロー連携のしく

みを拡張機能として導入することが最も重要な目的の1つとして挙げられる。

ワークフローエンジンの主な特徴を以下に示す。

- ワークフローの手動実行、定期実行に加えて、ワークフロー連携を実現するために「各システムやサービスにおけるデータ変化等によるイベント契機実行」に対応する。
- イベントを監視するための条件を、システムやサービスといった監視対象によらず統一的な記法で記述できる。
- Kubernetes CRD の YAML マニフェストとして宣言的に記述、管理することが可能。
- ワークフローの管理・運用を行うための統合的な機能、UI を提供する。（パラメータ指定実行、ワークフロー連携時の実行履歴の管理）
- ワークフローの基本機能、および要件を過不足なく満たすための拡張機能を疎結合に分離することにより、基本機能側の実装を改修せずに機能拡張が可能。

3.1 ワークフロー連携のしくみ

Wurfrahmen のコア機能であるワークフロー連携を実現するための考え方について説明する。各ワークフローを連携するための基本モデルとして、各組織のシステムやデータソースの状態やデータ変化をイベントとして監視し、そのイベントを検知した契機でワークフローを実行することにより組織内および組織間のワークフローを疎結合に連携するものとする。各組織のサービス、システム、データソース等をドメインとして定義し、ドメイン毎のイベントの監視条件を記述することにより、ワークフロー利用者は組織間のドメインをまたいだ疎結合なワークフロー連携を宣言的に記述して管理・運用できる。（図 2.1）

- 対象とする組織、ドメイン、データイベントは1つ以上を指定できる。
- 複数の組織のドメインとデータイベントを監視することにより、複数組織とのワークフロー連携を実現できる。

- ドメインとそのイベントに着目して連携を行うため、既存のバッチ処理システムやワークフローエンジン上で運用されるワークフローとも疎結合に連携することが可能。

3.2 イベント監視モデル

前述したとおり、Wurfrahmen では各ドメインのデータ変化などのイベントの監視と検知を行うことにより、組織内外のドメインのイベントに基づいたワークフロー連携を実現できる。イベントの監視と検知を行うためにワークフロー利用者はイベント監視条件を記述する必要があるが、ワークフロー利用者の学習コストや運用負荷を軽減するためにこのイベント監視条件の記述は統一したかつシンプルな構文で記述できることが望ましい。

これを実現するため Wurfrahmen ではドメインを定義し、そのドメインの中で Object（オブジェクト）と Attribute（属性）のシンプルなモデルを構築する。

- ドメイン：特定のシステムやサービスを指す。たとえばシステムであれば S3 や Hive、サービスであれば ABEMA^{*10}などがドメインとして挙げられる。ワークフロー連携に必要な抽象度に応じて、ドメインとして定義するサービスやシステムを選択する。
- Object：データの実体を指す。たとえば S3 ドメインであれば S3 オブジェクトなどが Object として挙げられる。
- Attribute：Object が持つ属性を指す。たとえば S3 ドメインであれば S3 オブジェクトのバケット名やタグなどが Attribute として挙げられる。

ワークフロー利用者はイベント監視条件の中で Object の Attribute の値やパターンを記述することで監視対象を指定できる。（図 2.2）

- 各ドメインに対して、データの実体を Object、データの属性を Attribute として定義する。

- Attribute は、Key (Attribute Name) と Value (Attribute Value) で構成される。
- Object は 1 つ以上の Attribute と紐付いており、Attribute から関連するすべての Object と紐付いている。
- Object, Attribute のデータ差分を管理することにより、Attribute と紐づく Object の状態および変化を検知可能。（Created, Updated, Removed, Exists）

イベント監視条件に合致する Object を検知した場合、Wurfrahmen ではこれに対応するドメインイベント契機でワークフローを実行し、実行されたワークフローの中では該当するドメインイベントの Object と Attribute を参照できる。これによりワークフロー外部で発生したドメインイベントの内容に応じた処理や制御を行うワークフロー連携を実現できる。

このイベント監視条件のモデルには以下の利点があり、ワークフロー利用者は Wurfrahmen を用いてワークフロー連携を伴うワークフローを記述する際の敷居を下げることに寄与している。

- イベント監視条件を記述するために SQL や、Argo Events 等の複雑なイベントルーティング制御の知識などを必要としないため学習コストが低い。
- ワークフロー利用者が各ドメインのデータの追加、更新、削除などを検知するためのスクリプトやプログラムを実装する必要がない。
- Object と Attribute によるシンプルなモデルを採用しているため、広く普及しているオブジェクト指向モデルの知識をベースにドメインの追加や拡張ができる。
- Object と Attribute の変化パターンからイベント検知を行うため、ユースケース拡張のたびにイベントを増やす必要がない。適切な Object と Attribute が定義されていれば、イベント監視条件でイベントを表現できる。

^{*10} <https://abema.tv/>

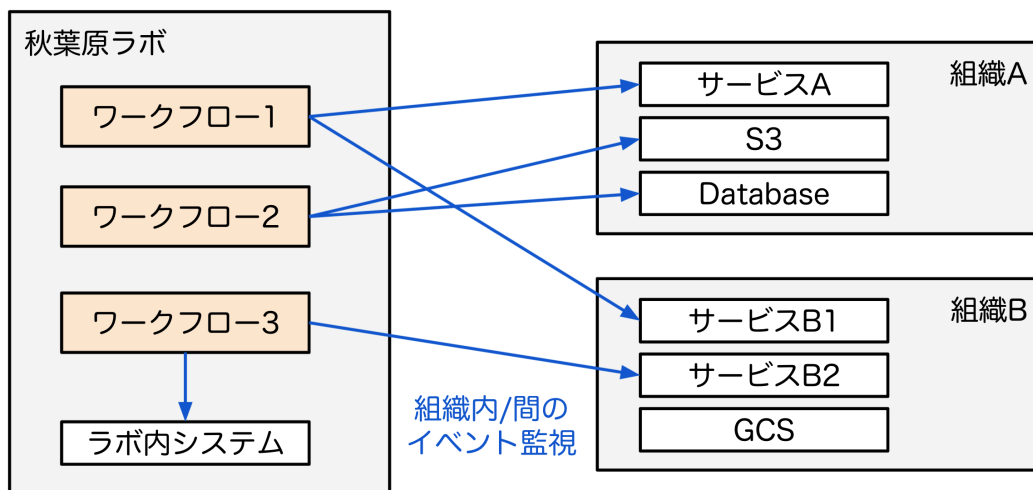


図 2.1 組織内外のワークフロー連携

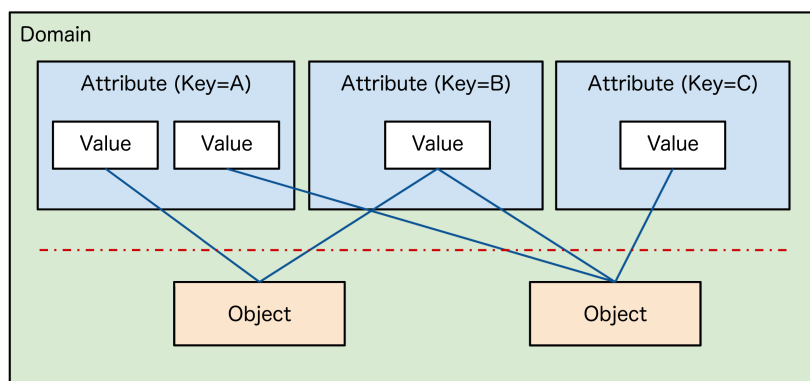


図 2.2 Object/Attribute のイベント監視モデル

イベント監視条件の記述例

イベント監視対象であるドメインおよび接続情報を指定したうえで、以下のイベント監視条件として以下の2つを記述する。

- **watch:** 監視対象の Object を条件文で記述する。“<Object 名>.<Attribute 名>”という形式で監視対象となる Object を指定し、==や>等の比較演算子を用いて Attribute の値と比較する形で記述する。
- **detect:** 監視対象の Object で検知すべきイベント。Created, Updated, Deleted, Exists のいずれかを指定できる。

図 2.3 に AWS S3 および Hive のドメインに対する domain, watch の記述例を示している。なお

この例で登場したドメインは IT サービスの開発や運用に携わる者にとっては広く普及したシステムやサービスなどのドメインだが、ドメインはこれ以外にもたとえば秋葉原ラボ内で提供する機械学習のモデル管理システム等の社内システムやサービスのドメインも定義している。これらのすべてのドメインで前述した Object, Attribute によるモデルに基づいたイベント監視条件を記述できる。

4 新ワークフローエンジンのアーキテクチャ

本節では、前述した Wurfrahmen のしくみや機能を実現するために採用しているアーキテクチャについて述べる。

対象キーのS3オブジェクトが更新される

```
watch: s3object.bucketName == "my_bucket" && s3object.key ==
"my_path/my_object"
detect: Updated
```

対象キー配下にS3オブジェクトが追加される

```
watch: s3object.bucketName == "my_bucket" && s3object.key =~ "^my_path/.*"
detect: Created
```

対象DB/テーブルにHiveパーティション (year > 2018) が追加される

```
watch: partition.databaseName == "my_db" && partition.tableName ==
"my_table" && partition.columns.year > 2018
detect: Created
```

図 2.3 各ドメインの監視条件記述例

Wurfrahmen ではワークフローを宣言的に記述、管理を行うため、Cloud Native Computing Foundation (CNCF) プロジェクト^{*11}に代表されるクラウドネイティブエコシステム、および Kubernetes エコシステムの恩恵を最大限に受けられる Kubernetes Operator パターンを用いて実装している。

前述した通り、Argo Workflows や Tekton 等の代替手段候補のワークフローエンジンは単独では秋葉原ラボの要件を満たすことはできないものの、以下の特徴をもつ Kubernetes CRD として基本機能が提供されている。

- Kubernetes の JSON / YAML の宣言的 API として定義されているため、言語やシステム非依存のインタフェースが提供されている。
- Kubernetes 上のリソースであるため、そのほかの Kubernetes リソースと同様に API 経由で spec や status の設定および監視ができる。
- ワークフロー完了時等に関連する Kubernetes イベントを発行するため、ワークフローの各種イベントも同様に監視ができる。

Kubernetes Operator パターンで実装されている Wurfrahmen は、Kubernetes の CRD リソースやイベントを介して、Argo Workflows や Tekton

と疎結合に連携する形で機能を拡張できる。これによりワークフロー実行制御の基本機能は Argo Workflows レイヤーが提供し、Argo Workflows レイヤーの基本機能では提供できないワークフロー連携などの機能を拡張機能として Wurfrahmen レイヤーが提供することを実現できる。Wurfrahmen, Argo Workflows, Kubernetes の組み合わせによるワークフローエンジンの構成図を図 2.4 に示す。

- Wurfrahmen：Argo Workflows では提供できない拡張機能を提供する。秋葉原ラボ開発のデータ変化監視コンポーネントと連携することにより、各ドメインのデータ変化等のイベント契機でのワークフロー実行を実現する。
- データイベント監視コンポーネント：イベント監視条件を記述することにより、各ドメインにおけるデータ変化等のイベント監視および検知を行う機能を提供する。もともとは Wurfrahmen の機能の一部だったが、ほか用途でも汎用的に利用できるようにするため別コンポーネントとして分離している。
- Argo Workflows：基本的なワークフロー実行および制御を提供する。
- Kubernetes：Kubernetes クラスター。ワークフローに関連するリソースやイベントの管理、制御を行う。

^{*11} <https://www.cncf.io/>

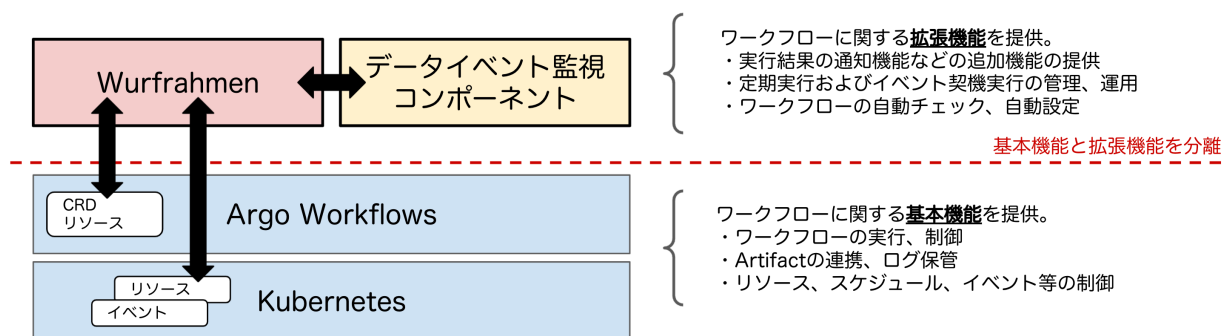


図 2.4 Wurfrahmen の構成

Argo Workflows と Tekton のうち、前者の Argo Workflows を基本機能を提供するワークフローエンジンとして採用した理由は、(1) 機能追加を含めた開発が活発であること、(2) Tekton よりも GCP 等の特定クラウド環境への依存度が低いこと、(3) Kubeflow^{*12}のように Argo Workflows を基本機能として利用している実績があること、(4) CNCF Hosted Project であること、といったものが挙げられる。

Wurfrahmen の拡張機能の提供には 3 つのパターンがあり、これらのパターンを組み合わせ実現している。拡張機能の提供例と合わせてそれぞれのパターンを以下に示す。

1. Wurfrahmen CRD リソースをもとにサーバとして提供：ワークフローのイベント契機実行、パラメータおよび対象タスクを指定可能なワークフロー実行の Web UI などを提供する。
2. Argo Workflows CRD リソースへの変換時に提供：Wurfrahmen CRD から Argo Workflows CRD に変換する時に、社内システムと連携するためのアドオン設定を行ったり、ワークフローを実行する時に必要な Argo Workflows および Kubernetes に関連する煩雑な設定を自動的に行ったりする。
3. ワークフロー関連の状態変化やイベントを契機に提供：Argo Workflows CRD リソースの状

態が変化したり、ワークフローやタスクが完了したりした時に発生する Kubernetes Event を契機に拡張機能が実行される。たとえばワークフロー完了時にその実行結果をワークフロー利用者に Slack 通知を行う。

いずれのパターンも Argo Workflows および Kubernetes のリソースとイベントを介して拡張機能を実現しているため、Argo Workflows や Kubernetes の内部構造や実装の詳細を把握したり依存したりすることなくそれぞれのインタフェースに基づいて機能を拡張できる（図 2.5）。

5 まとめ

本稿では組織間の効率的なワークフロー連携を実現することを目的としたワークフローエンジンである Wurfrahmen の概要、しくみ、アーキテクチャについて紹介した。

今後は (1) 秋葉原ラボの既存のワークフロー実行基盤から Wurfrahmen へのワークフロー移行、(2) 他組織の既存バッチ処理システムやワークフローエンジンから Wurfrahmen へのワークフロー移行を進めていく予定である。これにより混在するワークフローエンジンの Wurfrahmen への集約・統合を行い、組織内外でのワークフロー連携を促進する予定である。

^{*12} <https://www.kubeflow.org/>

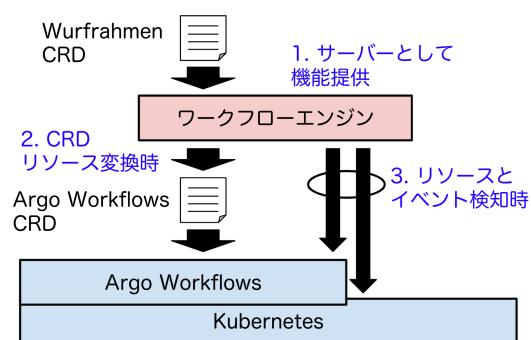


図 2.5 拡張機能提供の 3 つのパターン



Media Data Tech Studio

3

ABEMA における動画特徴抽出と推薦への応用

執筆者 若松 浩平

概要 当社のメディアサービスの一つである新しい未来のテレビ「ABEMA」に向け、推薦用途への応用やコンテンツ理解の促進を目的として AVE(ABEMA Video Embedding) を提供している。AVE は ABEMA のコンテンツである動画に対し機械学習とコンピュータビジョンに関する技術を利用することで分散表現を付与したものである。本項では動画からの特徴量抽出の取り組みである AVE v1 と、推薦への利用によるユーザー体験の向上を目的に行ったアップデートである AVE v2 について紹介する。v2 では v1 の有していた、内容の大きく異なるコンテンツを推薦してしまうという問題に対し、ユーザーの行動分析から得た知見をもとに改善を行った。結果としてオフラインでの検証でランキング指標の向上が確認できたほか、実際にオンラインで行った A/B テストでは一部ジャンルにおけるサービス指標の向上が確認できた。

Keywords 機械学習, コンピュータビジョン, 動画画像処理, 特徴抽出

1 はじめに

当社のメディアサービスの一つである ABEMA は新しい未来のテレビとして展開する動画配信サービスである。ABEMA ではユーザー体験の向上を目的にデータの活用を進めており、機械学習を用いた番組編成を自動化するプロジェクト、データを活用したキャスティング支援ツールの開発、番組のアートワークをパーソナライズするためのアルゴリズムの開発などに取り組んでいる。

本稿では ABEMA におけるデータ活用の一環として取り組んだ AVE(ABEMA Video Embedding) について紹介する。元来 ABEMA の推薦は視聴ログをベースにして行われており、視聴回数の少ないコンテンツに対して有効な推薦がしづらい、いわゆるコールドスタート問題を抱えていた。そこで AVE では動画から特徴量を抽出し、コンテンツベースの推薦へ活用することで前述の問題を解決しユーザー体験の向上を目指した。

AVE では始めにコンテンツベースでの推薦の実現のため、特に動画からの特徴抽出に主眼を置いた v1 を実装した。その後、v1 の運用の結果を踏まえ推薦におけるさらなる精度の向上を目指し、視聴ログによる特徴空間の変換を行った v2 を実装した。

これらの検証では、オフラインにおいて視聴回数の少ない動画に対しベースラインを上回る結果が得られ、コールドスタート問題の改善傾向とコンテンツベースの推薦の有効性を示した。また、オンラインでは一部ジャンルにおいてサービスの指標である CTVR の改善が確認された。

本稿ではこれらの v1, v2 の手法や実装、提供を行うシステム、その評価や検証について説明する。まず 2 節で動画からの特徴量抽出の取り組みである v1 について紹介し、続いて 3 章で推薦へ向けアップデートを行った v2 について紹介する。4 章では AVE を提供するシステムについて、最後に 5 章で本稿を総括して今後の課題について述べる。

2 AVE v1

本節では特に動画からの特徴抽出に主眼を置いて AVE v1 について紹介する。2.1 項で AVE v1 の導入先の推薦システムについて簡潔に述べ、2.2 項でどのように動画から特徴抽出を行うかについて述べる。また、2.3 項では AVE v1 の評価について、2.4 項では課題について説明する。

2.1 背景

ABEMA では、ユーザーに合わせ適切なコンテンツを配信するための推薦のしくみが複数導入されている。このうち、ユーザーが視聴中の動画に対し、次にユーザーが視聴しやすい動画を提示する推薦システムに AVE が導入されている。

推薦システムの概略を図 3.1 に示す。推薦システムはクエリに対し推薦したいコンテンツを候補としてあげる Candidate Generator、それらをユーザーに合わせ視聴されやすい順番に並べかえる Ranker から構成されている。

AVE の開発以前、システムを構成する Candidate Generator は視聴ログに基づいて推薦候補を作成していた。しかし、このような形式の推薦はユーザーが潜在的に興味を持つ動画の視聴数やクエリとなる動画が少ない際に、適切な推薦が行えなくなってしまうコールドスタート問題を有している。コールドスタート問題への対処としては、しばしばコンテンツベースの推薦があげられる。コンテンツから情報を抽出しそれを推薦に活用することで、コンテンツの持つログの量の多寡によらず適切なコンテンツを推薦できる。AVE v1 によって動画から得られた特徴を用いる Candidate Generator を導入することで、コンテンツベースの推薦を実現しコールドスタート問題の改善によるユーザー体験の向上を目指した。

2.2 処理

動画は連続した画像と音声によって構成されており、特に画像は事前学習されたモデルを利用することで比較的容易に品質の高い特徴抽出を行うことができる。そこで v1 では、「動画を構成する

画像からの事前学習モデルによる特徴抽出」、「画像特徴の集約」という 2 ステップで動画の見た目に関する特徴抽出を実現した。

処理イメージを図 3.2 に示す。まず、動画からフレーム画像を等間隔に 20 枚取得する。次に、取得したそれぞれの画像の特徴を抽出する。これには事前学習を行った Inception V3 [1] を用い、最終層の出力である 2048 次元のベクトルを画像の特徴とした。最後に、得られた 20 個のベクトルから各次元について平均値を取ることで 2048 次元のベクトルを作成しこれを動画の特徴量とした。

推薦での提供の際には、データサイズの削減を目的に PCA を利用して 100 次元に次元削減を行う。Candidate Generator として用いる場合は、次元削減後の特徴空間内でクエリとなる動画に対しユークリッド距離が近いコンテンツを推薦候補として提示する。

2.3 評価

オフライン検証

オフライン検証では、従来手法の行動ログベースの Candidate Generator と比較することにより AVE v1 の有効性を検証する。従来手法は (Weighted)MF [2] で得た動画ベクトルのコサイン類似度が高いコンテンツを候補集合とする。

ここでは、ランキング指標である nDCG [3] を用いて評価を行った。評価用のデータには、視聴ログからクエリの動画に対しユーザーが組み合わせて見やすい動画のランキングを作成し利用した。

結果として、nDCG@30 は視聴ログを基とするベースラインの手法 (MF) で 0.396、AVE v1 で 0.370 となった。また、動画の視聴回数に対する平均 nDCG のスコアをグラフにしたものを図 3.3 に示す。

全体の nDCG@30 ではベースラインの手法が v1 より良い結果を示すことが確認された。一方で、図 3.3 より視聴回数 10000 回以下においては v1 がベースラインの手法を上回る結果が確認された。したがって、視聴回数が少ない動画に対しベースラインより指標の改善が確認できることか

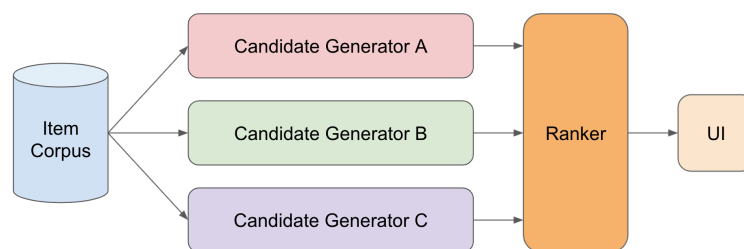


図 3.1 ABEMA のコンテンツ推薦システム

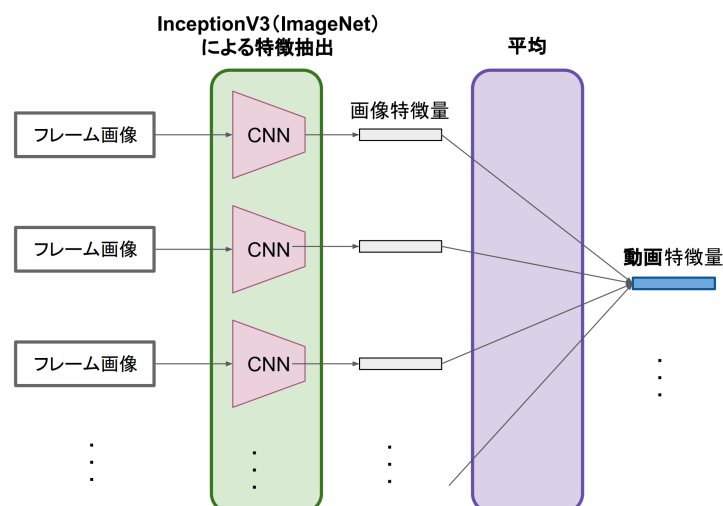


図 3.2 AVE v1 の特徴抽出の処理イメージ

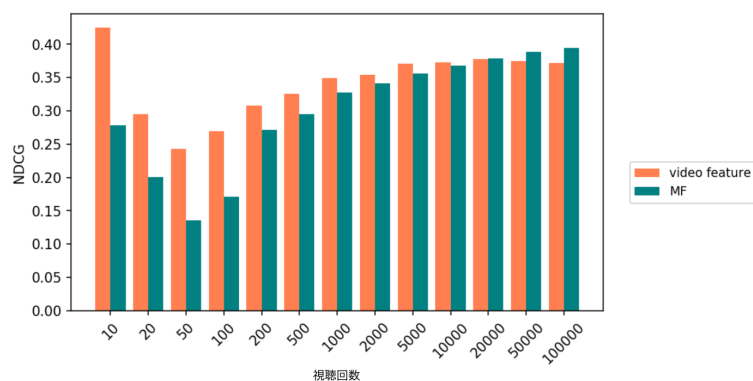


図 3.3 視聴回数に対する nDCG

ら、v1 がコールドスタート問題に対し有効に働く可能性が示唆された。

オンライン検証

オフライン検証の結果を踏まえ、v1 を Candidate Generator の一手法として導入する A/B テ

ストを 2020 年 3 月に実施した。評価指標には、提示された推薦候補がクリックされたのちに視聴に結び付いたかを示す CTVR を採用した。

ジャンルを抜粋してベースラインとの相対値を集計したものを表 3.1 に示す。ドラマ、恋愛番組、バラエティなどのジャンルや全体では CTVR が低

下したものの、アニメ、映画、麻雀ジャンルなどでベースラインに対し CTVR が改善された。オンライン検証の結果、CTVR の改善が確認された 5 ジャンルで v1 が本番環境において用いられることが決定した。

2.4 考察と課題

表 3.1 より、オンライン検証において v1 が有効に働くジャンルとそうでないジャンルに分かれることが確認された。

有効に働いたジャンルとしてアニメや麻雀、趣味があげられる。たとえばアニメでは画風が作品性を表現している場合、麻雀や趣味では動画を象徴するような道具を含むシーンが写るといった場合がある。そのため見た目の特徴抽出によって動画の内容をうまく表現でき、良い結果を示したと考えられる。

一方で、恋愛番組やドラマ、バラエティなどの多くの実写作品では CTVR が低下した。実写ジャンルは多くの場合実在の人間を被写体としており、見た目による違いが現れづらい。そのため v1 の特徴空間内で作品の関連性を表現できず CTVR が低下したと考えられる。

関連性のない動画を提示している可能性について調べるため、ベースライン、v1 の推薦集合 30 件が有する平均ジャンル数を算出した。一部ジャンルの結果を表 3.2 に示す。表から、ほとんどのジャンルでベースラインに比べ v1 の推薦集合の平均ジャンル数が多いことがわかる。特に恋愛番組、バラエティ、ニュースなどの CTVR の低下したジャンルではその傾向が顕著であり、v1 がジャンルの異なるような関連性の低い推薦集合を作成してしまうことがわかった。

3 AVE v2

本章では v1 の課題を踏まえ、距離学習の導入によって推薦での精度向上を目指した v2 について紹介する。

3.1 距離学習

距離学習とは入力とするデータの類似度がユークリッド距離やコサイン類似度などの尺度と対応するような別の空間に埋め込む手法である。

AVE における距離学習による特徴空間の変化のイメージを図 3.4 に示す。見た目の類似度が表現される v1 の特徴空間から、ユーザーにとっての動画の関連性が距離として表現される v2 の空間に埋め込むことで、推薦集合がより良い集合になることが期待できる。

また、特徴空間の補正のイメージを図 3.5 に示す。距離学習によってサンプル間の距離が調整されると、調整対象と同じ特徴を持つサンプルに対しても同時に調整が及ぶ。これにより直接学習の対象とならない動画に対しても距離学習の効果が及び、コールドスタート問題にも有効に働くことが期待される。さらに、今後抽出する特徴を増やした場合にも同様に補正の対象となり、学習後に追加されるサンプルに対しても精緻な関連性の再現が頑健に行われることが期待される。

3.2 視聴ログから教師データへの変換

距離学習をするにあたって教師データとしてコンテンツの関連性や類似度を表す尺度が必要となる。今回は尺度としてユーザーの視聴ログを利用することとした。利用にあたって、関連性のある動画はユーザーが連続して視聴しやすいという仮説のもと、ユーザーの一連の視聴行動の中に任意の動画のペアが含まれる回数を動画間の関連性とみなすこととした。以下の手順で視聴ログを教師データに変換する。

視聴ブロックの作成

ユーザーの一連の視聴行動を視聴ブロックと定義する。視聴ブロックにはユーザーが連続して視聴しやすい動画を含めることを目的とする。

ユーザーが視聴した動画を時系列に並べたものを $\{v_1, v_2, \dots, v_n\}$ (v_i は i 番目に視聴した動画) としたとき、 i 番目の視聴が完了してから $i+1$ 番目の視聴を開始するまでの時間を視聴間隔とよび

表 3.1 v1 導入の A/B テストの結果 (ベースラインとの相対値)

ジャンル	v1 の相対値 [%]
全体	-2.58
アニメ	+13.36
韓流・華流	+21.18
ドラマ	-7.78
趣味	+43.54
麻雀	+6.66
映画	+26.82
ニュース	-3.77
恋愛番組	-17.41
バラエティ	-12.13

表 3.2 推薦集合 30 件が有する平均ジャンル数

ジャンル	ベースライン	v1	相対値 [%]
全体	2.50	3.62	+59.31
アニメ	1.27	1.54	+20.74
韓流・華流	1.12	1.76	+56.71
ドラマ	2.02	2.78	+37.76
趣味	4.65	4.10	-11.67
麻雀	1.06	1.45	+36.91
ニュース	1.19	3.63	+204.63
恋愛番組	1.72	4.04	+134.77
バラエティ	2.39	6.33	+164.62

$t_{i,i+1}$ とする．すべての視聴間隔に対し，ある閾値 T と比較して $t_{i,i+1} > T$ の箇所で動画の集合を区切った部分集合を作成し，これらを視聴ブロックとする．

教師データへの変換

同一視聴ブロック内における重複を排除した 2 つの動画の組み合わせの登場回数を，全ユーザーのすべての視聴ブロックで合計する．この数値の大きい上位の組み合わせを関連がある動画とみなす．

3.3 Triplet Semi-Hard Loss

距離学習を行うにあたり，どんなサンプルを対象にしてどのような距離の補正を行うかを決め

る必要がある．今回は学習の設定のしやすさと効率を踏まえ Triplet Loss + Semi-Hard Negative Selection [4] を採用することとした．

Triplet Loss

Triplet Loss は関連性が高いとされるペアよりも関連性の低いペアが埋め込み空間上で近い場合にペナルティを設ける損失関数である．

任意のサンプルである anchor, anchor と類似度が高いとされる positive, 類似度の低いとされる negative を入力とする．動画 v_i に対する v1 のベクトルを x_i ，そのベクトルに対する変換の操作を $f_\theta(x_i)$ とした場合，式 (3.1) を最小化するようにパラメータ θ を最適化する．なお， $f_\theta(x_i)$ の出力は

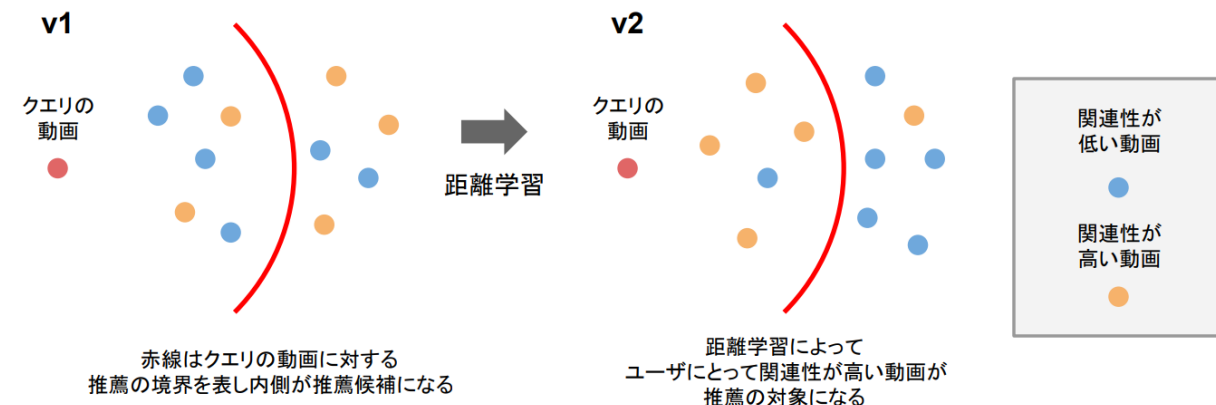


図 3.4 距離学習の前後における特徴空間の変化のイメージ

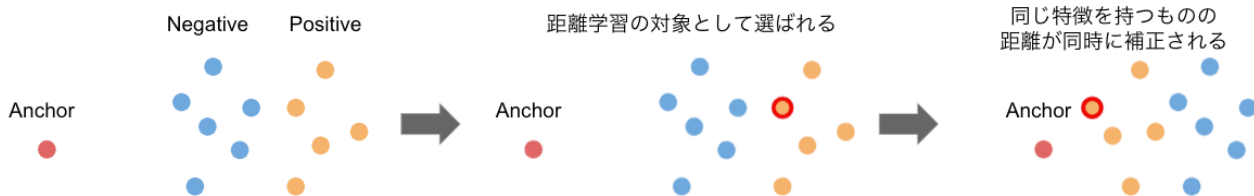


図 3.5 距離学習による特徴量の補正のイメージ

変換を実施するモデルに依存し、任意の次元のベクトルとなる．今回は $v1$ と同じ 2048 次元のベクトルとしている．

$$\sum_{i=1} L(f_{\theta}(x_i), f_{\theta}(x_{i_{\text{positive}}}), f_{\theta}(x_{i_{\text{negative}}})) \quad (3.1)$$

$x_i, x_{i_{\text{positive}}}, x_{i_{\text{negative}}}$ は i 番目のサンプルに対する anchor, positive, negative に対応した $v1$ のベクトルである．ここで L は損失関数であり $f(x_i)$ に対し $f_{x_{i_{\text{negative}}}}$ の方が $x_{i_{\text{positive}}}$ よりも埋め込み後の特徴空間内で近い際にペナルティを設ける関数を設定する．

今回は式 (3.2) のヒンジ損失関数を用いる．

$$L_{\text{hinge}}(x, y, z) = \max(0, \|x - y\|_2^2 - \|x - z\|_2^2 + \alpha) \quad (3.2)$$

今回は、マージンの設定は実験の結果から $\alpha = 0$ としている．最適化されたパラメータ θ を用いて動画 v_i に対応する $v1$ のベクトル x_i から得られた $f_{\theta}(x_i)$ が動画間の関連を表現した分散表現である．

3.4 評価

オフライン評価

オフライン評価では $v1, v2$ に対し、2.3 と同様にランキング指標である nDCG を用いて評価を行った．

結果として、nDCG@30 は $v1$ で 0.2797, $v2$ で 0.2948 となり、 $v1$ に対し 1.52 ポイントの改善が確認できた．また、ジャンルごとに指標の差異について確認すると、 $v1$ が課題としていた実写ジャンルである恋愛番組やバラエティを含む 14 ジャンルで nDCG が改善したことが確認できた．

オンライン評価

2020 年 12 月に $v2$ を Candidate Generator として導入する A/B テストを実施した．ベースラインとの CTVR の相対値をジャンルを抜粋して表 3.3 に示す．全体ではベースラインに対し CTVR が低下したものの、 $v1$ 導入時と比較して相対的に

低下の幅は小さくなった ($-2.72 \rightarrow -2.02$) ほか、格闘や趣味、映画ジャンルで改善傾向が見られた。

v1 で課題となっていた実写ジャンルの恋愛番組については、ベースラインに対しては低下したものの v1 導入時と比較して大幅な改善が見られた ($-17.55 \rightarrow -0.16$)。一方で、依然としてドラマやバラエティでは CTVR が低下する結果となった。

3.5 考察と課題

オンライン検証の結果、新たに改善傾向が見られたジャンルとして「格闘」が挙げられる。既存の配信ロジックが同ジャンルでの急上昇作品やユーザーが視聴中のシリーズの新作作品をレコメンドするのに対し、v2 は同大会の動画といった内容に関連性のあるものをレコメンドでき CTVR が向上したと考えられる。

多くの実写ジャンルでは v1 と同様に CTVR が低下した。音楽ジャンルを例に挙げると、あるグループのライブ映像に対し、v2 ではグループを問わずライブ映像をレコメンドするという結果が見られた。実際に音楽ジャンルのライブ映像を視聴したユーザーについて調べたところ、同じキャストが出演するバラエティジャンルの番組を連続して視聴するといった視聴行動が確認された。ライブ映像とバラエティ番組のような、変換前である v1 の差異が大きいもののユーザーが連続してみやすい動画の組み合わせは距離学習による補正が難しく、CTVR が低下してしまうと考えられる。

v2 で利用する距離学習モデルは入力される特徴を連続した視聴のしやすさに結び付ける。このため、特定のコンテキストに基づいた視聴が、入力する特徴によって説明できない場合にうまく学習ができないと考えられる。今回の場合では見た目に関する特徴を表現した v1 を入力としたため、たとえば音声に関わる視聴行動などは特徴空間で表現できない。この問題を改善するためには、入力である特徴にユーザーの視聴行動を説明できるような情報を増やす必要がある。

特徴量に情報を追加する際には異なるドメインの特徴をどの段階でどのように結合するかがしば

しば課題になる。v2 で距離学習を行うモデルは多次元のベクトルを出力するモデルであれば内部の構造に制限がなく、特徴量を結合する際に起こり得る課題に対し柔軟に対応できる。また、連続した視聴のしやすさを表現するという目的に従ってそれぞれのドメインの特徴から適切に学習することが期待されるため、情報の追加に対して頑健に働くと考えられる。

4 システム

本節では AVE v2 の提供システムについて紹介する。本システムは動画からの特徴量の抽出を行う推論バッチ、モデルを学習する学習バッチから構成され、秋葉原ラボのバッチ処理実行基盤上で実装、提供した。また、バッチ内で利用されるプロセスやデータ、モデルのバージョン管理には同じく秋葉原ラボの機械学習モデル管理基盤である Etna [5] を利用した。

推論バッチの概略を図 3.6、学習バッチの概略を図 3.7 に示す。

バッチの概略図中の丸はプロセスを、四角はデータを表しており、それぞれ Etna を利用してバージョンを管理している。特に生成される特徴量、距離学習モデル、PCA モデルについては Etna が提供する API を通してデータをアップ/ダウンロードしている。ここで PCA モデルは提供時のデータサイズを小さくするために次元削減として導入している。

Etna を採用した理由として、特に推論時のモデル間のバージョンの整合性の管理が挙げられる。PCA への入力には距離学習モデルによって変換された特徴量が想定されるが、変換に利用する距離学習モデルのバージョンが異なる場合、特徴空間も変化するため正しい次元削減が期待できない。そのため PCA のバージョンから利用するモデルを正しく選定する必要がある。実際に A/B テスト実施中に学習のバッチが不具合で動作停止してしまい距離学習モデルのみが更新されるという事態が起こったが、正しく PCA に対応する距離学習モ

表 3.3 v2 導入の A/B テストの結果 (ベースラインとの相対値)

ジャンル	v2 の相対値 [%]
全体	-2.02
アニメ	-0.43
韓流・華流	-0.07
ドラマ	-5.38
格闘	+4.05
趣味	+17.30
麻雀	+0.66
映画	+3.06
ニュース	+3.60
恋愛番組	-0.16
バラエティ	-13.78

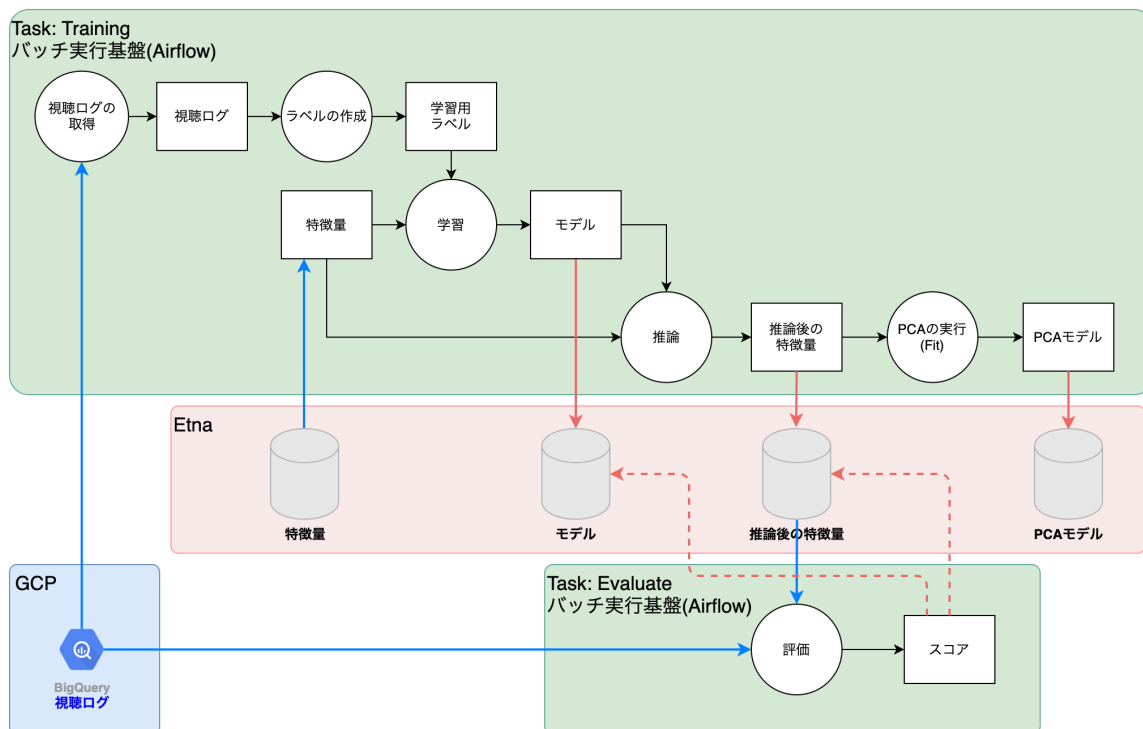


図 3.6 学習用バッチの処理の概略図

デルを読み込むことができ異常な出力を防ぐことができた。

5 まとめと今後の課題

本稿では ABEMA の推薦システムにおけるコールドスタート問題の解決を目的として動画からの特徴抽出を行った AVE について紹介した。動画

からの特徴抽出を主眼に置いた v1 ではオフラインにて視聴回数の少ない動画に対しランキング指標改善の結果が見られ、動画特徴量によるコールドスタート問題改善の傾向を示したほか、オンラインでも一部ジャンルにおいて指標の改善傾向が見られた。距離学習によって推薦における活用での精度向上を目指した v2 では、オフライン指標で

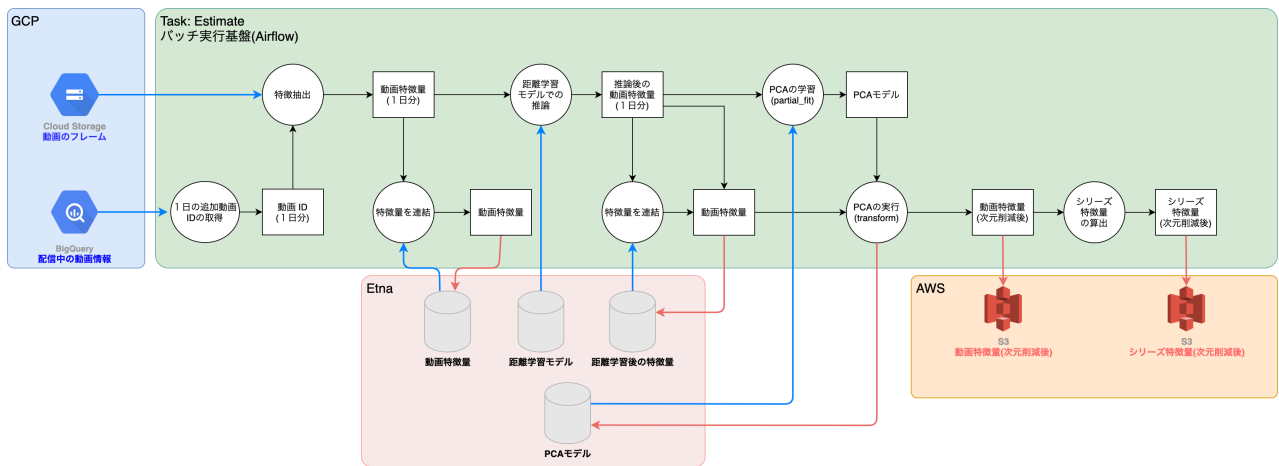


図 3.7 推論用バッチの処理の概略図

v1 を上回る結果が得られたもののオンラインでは画像を基にしたレコメンド自体が有していた課題があらためて確認された。一方で視聴ログを基にした距離学習が今後の特徴量の追加に対して頑健に機能する可能性についても示唆された。

動画からの特徴抽出や推薦への応用はいまだ確立された手法がなく、発展途上の分野である。応用の一例として何らかの参考になる点があれば幸いである。

参考文献

- [1] C. Szegedy, V. Vanhoucke, S. Ioffe, J. Shlens, Z. Wojna. “Rethinking the Inception Architecture for Computer Vision”, In *CVPR*, 2016.
- [2] Y. Hu, Y. Koren, C. Volinsky. “Collaborative Filtering for Implicit Feedback Datasets”, In *IEEE International Conference on Data Mining (ICDM 2008)*, pages 263 – 272, 2008.
- [3] Kalervo Järvelin, Jaana Kekäläinen. “Cumulated gain-based evaluation of IR techniques”, In *ACM Transactions on Information Systems* 20(4), 422 – 446, 2002.
- [4] F. Schroff, D. Kalenichenko, J. Philbin. “FaceNet: A Unified Embedding for Face Recognition and Clustering”, In *CVPR*, 2015.
- [5] 大内裕晃. “データと処理の依存関係を整理する機械学習モデル管理基盤の開発”, CyberAgent 秋葉原ラボ技術報告 Vol. 3, 2019.



Media Data Tech Studio

4

類似画像検索システムの設計と運用

執筆者 松月 大輔

概要 秋葉原ラボは当社メディアサービスから日々生成される大量のデータを活用したサービスの健全化への取り組みやレコメンデーションシステムの開発を行うことでサービスの品質向上に貢献している。本稿ではこれらの取り組みに応用が期待される類似画像検索システムの開発・運用について紹介する。開発した類似画像検索システムは、類似画像検索の課題である検索部分のレイテンシを改善するようなくみに加え、特徴ベクトルの生成を汎用的に行うことができる社内基盤であるフィルタ基盤を使用することにより、複数のサービスに応用しやすい設計を実現した。

Keywords 類似画像検索, 画像特徴量, ANN

1 はじめに

近年、画像を取り扱うインターネットサービスが拡大している。当社においても例外ではなく、ブログサービス「Ameba ブログ*1」やマッチングアプリ「タップル*2」など複数のサービスが画像データを取り扱っている。秋葉原ラボではユーザー投稿画像を含むプロダクトから日々生成される大量のデータを用いてレコメンデーションシステムの開発やサービスの健全化に取り組んでいる。本稿では、画像からレコメンデーションおよび健全化を促進する手法として期待している類似画像検索システムの設計について紹介する。

画像データによりユーザーが行動を選択するケースは多々存在する。たとえば Ameba ブログや ABEMA*3のサムネイル画像を根拠に行動を選択するケースが挙げられる。画像がユーザーの行

動に影響を与える場合、類似画像検索によるレコメンデーションは類似コンテンツの提供が可能であるため有用である。類似画像検索を用いたレコメンデーション手法では、行動ログベースのレコメンデーション手法の課題でもあるコールドスタート問題*4の解決策として期待される。

また、類似画像検索は健全化への応用が期待される。現在秋葉原ラボでは上岡 [1], 岩井 [2] が画像分類技術を用いたフィルタ開発により、ユーザーの投稿画像に対する自動フィルタリングサービスを運用しており、サービスの健全化に大きな貢献を果たしている。画像分類は教師あり学習であるため、サービスにとって不適切なカテゴリを教師データとして用意する必要があることに対し、類似画像検索は教師データを必要としない。そのため、カテゴリを問わず画像の類似性による不適切コンテンツの再投稿防止が可能であり、画像分類

*1 <https://ameblo.jp/>

*2 <https://tapple.me/>

*3 <https://abema.tv/>

*4 行動ログがない新規データでは適切なレコメンドが難しいこと

による手法と異なる観点で健全化への貢献が期待できる。以上の理由から秋葉原ラボでは類似画像検索システムの構築は有意義であると考え、実際に開発・設計を行っている。

本稿では、まず2節で関連研究について紹介し、3節では実際に設計した類似画像検索システムについて記述する。最後に4節でまとめと今後の課題について述べる。

2 関連研究

類似画像検索は一般的に検索したい画像を何らかの特徴ベクトルに変換し検索クエリとする。そして検索対象の画像群を特徴ベクトル集合に変換し、検索クエリとの間で類似度計算を行うことで類似度の高い画像を取得できる。近年、類似画像検索技術は画像認識技術の大幅な改善と高速な近傍探索手法の提案に加え、大量のメタデータを含む画像データがインターネットサービス上で扱われることで注目されている。

1990年～2000年代では、SIFTやHOGなどの特徴量を利用した検索手法が研究されてきた。SIFT(Scale Invariant Feature Transform)特徴量[3]はスケール空間を用いた手法で、画像の回転・拡大縮小・照明変化にロバストな特徴量である。またHOG(Histogram of Original Gradient)特徴量[4]は局所的な画像勾配をヒストグラム化する手法で、定点カメラで撮影した画像のように画像の構成が大きく変化しない場合に有効である。これらの手法はそれぞれのメリットデメリットを考慮したうえで、さまざまな類似画像検索に応用されてきた。近年ではこれらの特徴量にとって代わり、CNN[5]のようなDeep Learningを用いた特徴量抽出手法が一般的である。これは画像分類などの教師あり学習の問題設定でCNNを学習し、その中間層の出力を特徴量として扱うものである。特に距離学習とCNNを組み合わせ

た手法は類似画像検索に非常に有効的で、2015年に発表されたFaceNet[6]では距離学習手法であるSiameseNetwork[7]の特徴量抽出部分にCNNを採用しており、顔画像における類似画像検索の分野で大きな精度改善を実現している。2021年現在においても、顔画像における類似画像検索の分野においての研究では距離学習とCNNの組み合わせがトレンドである[8, 9]。

類似画像検索を行うにあたり、近傍探索の高速化はシステム運用において重要な課題である。近年では高速化のためにANN(Approximate Nearest Neighbor)が広く扱われている。実際にFacebook AI Researchは高速なANNアルゴリズムであるFaiss^{*5}の開発を行っており、ライブラリを公開している。また音楽ストリーミングサービスのSpotifyはAnnoy^{*6}を公開している。Annoyは実際にSpotify内部の音楽レコメンデーションで使用されており、ANNによるアプローチがインターネットサービス上で有効であることを示している。秋葉原ラボではPhoenix^{*7}が近傍探索アルゴリズムを実装しており、グラフ探索による高速なANNの手法であるHNSW(Hierarchical Navigable Small World graphs)[10]アルゴリズムを提供している。今回設計した類似画像検索システムの近傍探索部分においてもPhoenix上でHNSWアルゴリズムを使用した近傍探索を行っている。

3 類似画像検索システムの構築

本節では実際に秋葉原ラボで設計した類似画像検索システムの機能と設計について説明する。3.1項では類似画像検索システムの要件と設計指針について、3.2項では実際に設計した類似画像検索システムの概要について解説する。

3.1 要件と設計指針

類似画像検索システムはいくつかの要件がある。本項では要件を提示し、要件に対する設計指針を

^{*5} <https://github.com/facebookresearch/faiss>

^{*6} <https://github.com/spotify/annoy>

^{*7} 秋葉原ラボが開発している推薦基盤

示す。

低レイテンシで類似画像検索の検索結果が取得可能である

類似画像検索の検索時に低レイテンシで検索結果を取得することは、類似画像検索システムの開発において一般的な課題である。

図 4.1 に設計指針をまとめたフロー図を示す。

検索結果取得におけるレイテンシのボトルネックとなるのは、近傍探索とクエリとなるアイテムの特徴ベクトル生成である。近傍探索の高速化では一般的に ANN を使用する方法が一般的である。設計した類似画像検索システムにおいてもグラフ探索による ANN の手法である HNSW を採用している。特徴ベクトルの生成を行う工程は、一度生成した画像の特徴ベクトルは再度計算処理が行われないようにすることで、特徴ベクトル生成のリクエスト数を削減しレイテンシの改善が可能である。そこで、検索対象データを登録するフローを設計し、その中で特徴ベクトルを保存するデータベースを構築する。図 4.1 に示すように、類似画像検索を行いたいアイテムの情報を基に検索対象データが登録されているデータベースから特徴ベクトルを抽出し、ANN を実行することで検索結果を取得する。反対に、データベースに特徴ベクトルが保存されていない場合は、その都度特徴ベクトルを生成する。こうすることで、検索対象データおよび、検索クエリとなる画像が類似画像検索システムにとって既知である場合、低レイテンシで検索結果を取得可能になる。

検索時に属性情報でフィルタリング可能である

類似画像検索の検索時には、属性情報でフィルタリング可能であると便利である。たとえば、アイテムの同一カテゴリ (ジャンルやタグなど) 内で検索結果を取得したい場合、検索結果の取得後にフィルタリングするのは非効率であり、検索結果がフィルタリングされた状態で取得可能であることが望まれる。

属性情報で類似画像検索結果をフィルタリング

するためには、データベースに属性情報を含めた状態でデータを登録し、ANN を実行する際に必要になるインデックスの構築時のデータ参照で属性情報による絞り込みが実現できればよい。本実装ではランダムアクセスに強い HBase をデータベースとして選択し、スキーマに属性情報を追加する実装を行った。

また、インデックスの構築は検索のリクエストごとに実行すると検索結果取得のレイテンシを悪化させてしまうため、バッチ処理で実行する。ここで行うインデックスの構築は、すでにインデックスに追加済みのデータと更新時に新規で追加したいデータの差分が抽出できると、効率的にインデックスの更新が可能である。そこで、HBase のスキーマにデータ登録時の timestamp をセットすることで、差分抽出を可能にしている。

複数のサービスで利用可能である

当社メディア事業部は複数のサービスを展開している。そのため、秋葉原ラボでは複数サービスで利用可能な汎用的なシステム設計を行っている。類似画像検索システムも例外ではなく、新しいサービス適用時に少ない工数で機能適用が可能であることが望ましい。

類似画像検索システムを提供するサービスにより設計に差分が生じる可能性が高い部分は特徴ベクトルの生成方法である。特徴ベクトルの生成を汎用的に行うために、設計した類似画像検索システムでは特徴ベクトル生成部としてフィルタ基盤 [12] を使用した。フィルタ基盤では、物体検出や特徴量抽出、画像のハッシュ値計算などをそれぞれ単体のフィルタとして実装が可能であり、またそれらのフィルタを任意の組み合わせで定義可能である点から、汎用的な特徴量抽出部分として適していると考えた。

3.2 類似画像検索システムの設計

本項では 3.1 項で示した要件と設計指針の基、実際に設計した類似画像検索システムについて解説する。図 4.2 に類似画像検索システムの全体図を

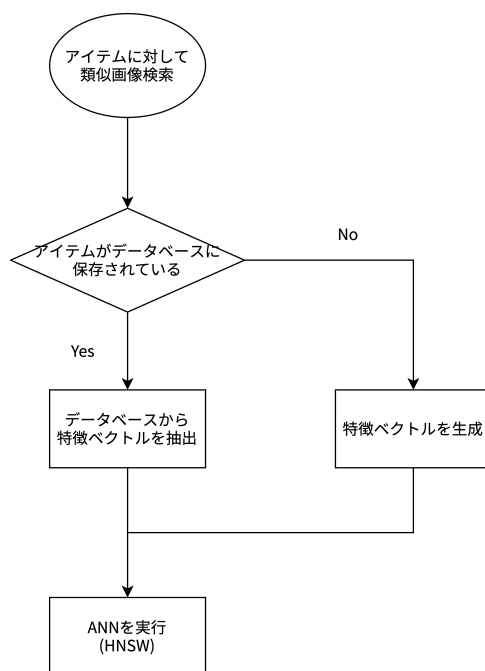


図 4.1 検索結果取得の高速化フロー

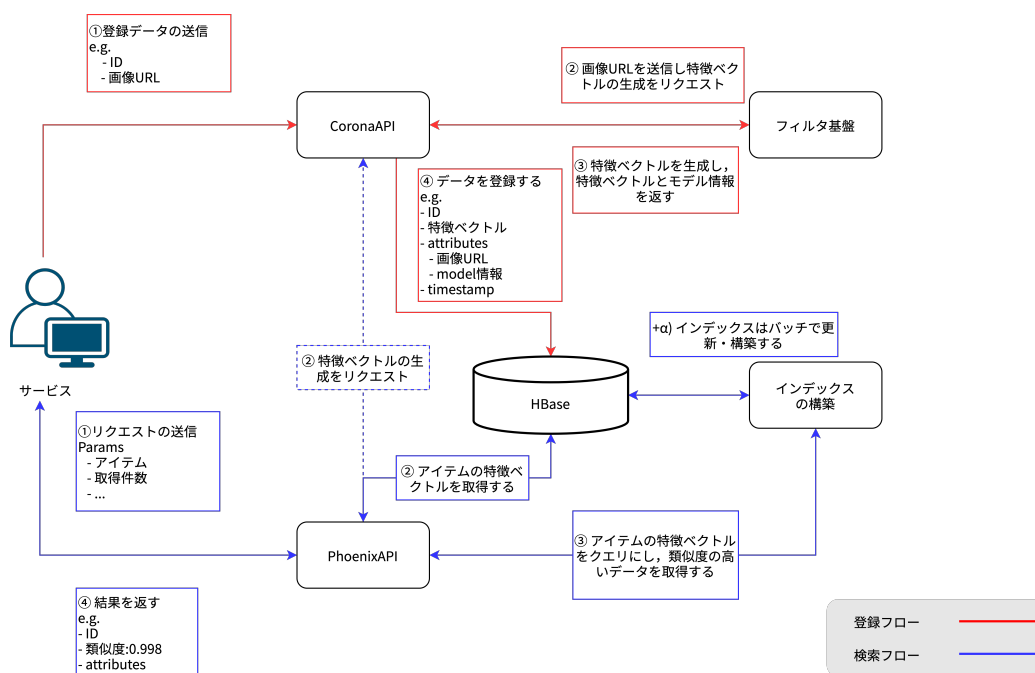


図 4.2 類似画像検索システムの全体図

示す。

類似画像検索システムの全体像は、サービスから検索対象データ送信を受け付けるための CoronaAPI と、特徴ベクトルの生成を行うフィルタ基盤、検索対象データおよび特徴ベクトルを保存するための HBase により構成される登録フローと、

類似画像検索の検索リクエストを受信し ANN を実行する PhoenixAPI および、ANN を実行するために必要なインデックスを構築するコンポーネントにより構成される検索フローの 2 つのフローに分けて設計した。

登録フローは図 4.2 の赤矢印で示している通り、

まずはサービスから登録したい画像 URL を属性情報とともに CoronaAPI に送信する。次に CoronaAPI はフィルタ基盤に画像 URL を送信し、特徴ベクトルを取得する。最後に取得した特徴ベクトルとサービスから受信した属性情報に加え、データ登録を行うタイミングの timestamp を HBase に登録する。

検索フローでは図 4.2 の青矢印で示している通り、サービスから取得したい検索クエリとなるアイテム情報を PhoenixAPI に送信する。次に PhoenixAPI は特徴ベクトルを取得するためにまずは受信したアイテムが HBase に登録済みのデータであるかを参照し、登録済みのデータである場合、HBase から特徴ベクトルを抽出する。登録済みでない場合は CoronaAPI に画像 URL を送信し、特徴ベクトルを生成・取得する。その後、取得した特徴ベクトルをクエリにし類似画像検索結果を取得する。検索結果のフィルタリングを行う場合はここで静的な属性情報を指定することで実現される。

また、これらのコンポーネントのインフラ環境は秋葉原ラボで開発・運用している Kubernetes である lab-k8s [11] 上で動作しており、それぞれのコンポーネントごとにメトリクスの監視を行っている。大まかな流れは以上の通りであるが、以降各コンポーネントについて詳しく説明する。

CoronaAPI

CoronaAPI はサービスが類似画像検索システムにデータを登録するための機能を提供する。具体的には、サービスからの API リクエストで受け取った画像データを基にフィルタ基盤へ特徴ベクトルの生成をリクエストし、生成された特徴ベクトルを HBase に保存する。HBase に登録するデータは、ID、特徴ベクトル、属性情報 (attributes) および timestamp をスキーマに登録しており、属性情報は検索フローにおけるサービスが受信する類似画像検索結果に含まれる。

CoronaAPI ではデータの登録だけではなく、画像 URL を受信し、特徴ベクトルを返すようなエン

ドポイントを提供している。具体的には、図 4.2 の検索フロー青点線矢印で示すように、PhoenixAPI が ANN を実行する際に HBase から特徴ベクトルを参照できない場合には、特徴ベクトルを生成する必要があるというケースに対応している。

フィルタ基盤

フィルタ基盤は CoronaAPI から gRPC で画像 URL を受信し、レスポンスとして特徴ベクトルおよびモデルの詳細を返す。フィルタ基盤では単体のフィルタの設計とフィルタの任意の組み合わせであるチェーンが定義可能である。たとえば、物体の一部分にフォーカスするような特徴ベクトルを生成したい場合は以下のようなチェーンを定義する。

1. 画像ダウンロードフィルタ
2. 物体検出フィルタ
3. 特徴ベクトル生成フィルタ

ここで定義したチェーンは、CoronaAPI から画像 URL を受信するタイミングで指定可能である。

PhoenixAPI

サービスは類似画像検索を行いたいアイテムと、取得したい件数を PhoenixAPI に送信することで、任意の数だけ類似したアイテムを取得できる。ここで類似したアイテムとは類似度、ID および属性情報 (attributes) であり、attributes は CoronaAPI 節で定義した属性情報が取得できる。PhoenixAPI 内部ではアイテムの ID や画像 URL を基に HBase から特徴ベクトルを抽出、抽出できない場合は CoronaAPI に特徴ベクトルの生成をリクエストし取得している。そして取得した特徴ベクトルをクエリとして ANN を実行し類似画像検索結果を取得している。

インデックスの構築

ここでは HBase から特徴ベクトルを参照し、インデックスを構築する。ここで HBase を参照する際に登録されている属性情報によりフィルタ条件

を加えることで、属性情報で絞り込まれたインデックスが構築され、フィルタリングされた検索結果を取得できる。また、インデックスの構築はバッチ処理で実行している。そのため、検索対象データの反映は一定の遅延が発生する。登録データが分間で数件程度であれば、毎秒差分更新を行うコストが小さいため、低遅延で運用可能であるが、登録データが毎秒数百件追加されるようなサービスでは遅延が無視できなくなってしまう。これは今後の改善点として取り組んでいきたいと考えている。

4 まとめと課題

本稿では画像特徴量を用いたレコメンデーションおよび健全化への貢献を目的に開発している類似画像検索システムについて紹介した。設計した類似画像検索システムでは、検索結果を低レイテンシで取得するために特徴ベクトルを HBase に保存し、検索時に特徴ベクトルを参照するような設計を行った。今後は、検索対象データもオンラインで更新されるしくみを開発することに加え、本システムの適用事例を増やし、システムの汎用化のために必要な課題整理および開発を行っていきたいと考えている。

参考文献

- [1] 上岡 将也：“マッチングサービスにおける画像審査の自動化”，CyberAgent, CyberAgent 秋葉原ラボ技術報告 Vol. 3, 2019.
- [2] 岩井 二郎：“NG 画像フィルタの設計から運用まで”，CyberAgent, CyberAgent 秋葉原ラボ技術報告 Vol. 3, 2019.
- [3] D.G. Lowe：“Object Recognition from Local Scale-invariant Features”，IEEE International Conference on Computer Vision (ICCV), pp.1150-1157, 1999.
- [4] N. Dalal, B. Triggs：“Histograms of oriented gradients for human detection”，IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR), pp.886-893, 2005.
- [5] A. Krizhevsky, S. Ilya, *et al.*：“Imagenet classification with deep convolutional neural networks”，Advances in Neural Information Processing Systems (NIPS), pp.1097-1105, 2012.
- [6] F. Schroff, D.Kalenichenko, *et al.*：“Facenet: A unified embedding for face recognition and clustering”，IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR), pp.815-823, 2015.
- [7] J. Bromley, I. Guyon, *et al.*：“Signature verification using a “siamese” time delay neural network”，Advances in Neural Information Processing Systems (NIPS), pp.737-744, 1993.
- [8] X. Zhang, R. Zhao, *et al.*：“Adacos: Adaptively scaling cosine logits for effectively learning deep face representations”，IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR), pp.10823-10832, 2019.
- [9] J. Deng, J. Guo, *et al.*：“Arcface: Additive angular margin loss for deep face recognition.”, IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR), pp.4690-4699, 2019.
- [10] YU. A. Malkov, D.A.Yashunin: “Efficient and Robust Approximate Nearest Neighbor Search Using Hierarchical Navigable Small World Graphs”, IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol.42.4, pp.824-836, 2018.
- [11] 岩城 洋一 佐藤 栄一 松原 周平：“秋葉原ラボの研究開発とシステム運用を支えるクラウドネイティブ環境”，CyberAgent, CyberAgent 秋葉原ラボ技術報告 Vol. 3, 2019.
- [12] Connolly, Juhani：“フィルタ基盤の開発における取り組みと利用されている技術”，CyberAgent, CyberAgent 秋葉原ラボ技術報告 Vol. 3, 2019.



Media Data Tech Studio

5

pigパーティにおける Orion フィルタの候補キーワード抽出

執筆者 上田 紗希

概要 pigパーティではサービスの健全性を維持するために、総合監視基盤システム Orion を導入しユーザーの投稿に不適切なキーワードが含まれていた場合に検知・対処するためのコンテンツフィルタを運用している。しかし日々膨大な投稿が行われる本サービスでは Orion フィルタに登録すべき新しいキーワードの探索が人的または時間的リソースの問題から困難であり、加えて表現の複雑性から監視者が独自にキーワードを推測することもまた困難であるという問題があった。本稿ではこの問題を解決するために導入された Orion フィルタへの追加候補となるキーワードを自動抽出するシステムについて紹介する。

Keywords 機械学習, 自然言語処理

1 はじめに

サイバーエージェントのメディア事業ではさまざまなサービスを提供しており、それらの中にはユーザーによるコンテンツの投稿が可能であるサービスも数多く存在する。たとえば、新しい未来のテレビ「ABEMA^{*1}」ではユーザーは配信されているコンテンツに対しコメントを投稿できる。また、ブログサービス「Ameba ブログ^{*2}」ではテキストや動画画像を含む記事やそれぞれの記事に対するコメントを投稿できる。マッチングアプリ「タップル」ではユーザーは自分の画像やテキストを用いたプロフィールを投稿・公開し、ほかのユーザーとメッセージの交換が行える。

このようなユーザーによって生成されるコンテンツを含むサービスを運用していく上では、しばしば悪意のあるユーザーによる不適切コンテンツの投稿が他ユーザーの利益や利用体験を損なったり、あるいは不適切コンテンツが大量に投稿され

ることによってサービスのネットワークやサーバのリソースを圧迫したりといった問題が発生する場合がある。したがって、サービスの健全性を維持するためにはこれらの不適切コンテンツに対する監視・対応を行うことが求められる。

秋葉原ラボではメディア事業が運営するサービスを健全に保つため、これらの投稿を抽出するためのフィルタとフィルタに基づいた総合監視基盤システム Orion [1] を開発・運用している。Orion はさまざまなサービスに接続可能であり、不適切投稿に備えたフィルタ群を事前に作成しておくことでサービスの不適切投稿を迅速に発見し、それらに対応できるしくみを実現している。たとえば Ameba ブログでは Orion を用いてコピースパムコンテンツと思われるブログ投稿を自動判定し、疑いがあるコンテンツについては Web UI を通じて監視オペレータが内容を確認、必要に応じて非表示にするシステムを運用している。

^{*1} <https://abema.tv/>

^{*2} <https://ameblo.jp/>

本稿ではメディアサービスの一つであるピグパーティ^{*3}において Orion の活用支援として導入されている Orion フィルタの候補キーワード抽出システムについて紹介する。ピグパーティでは日々大量の投稿が行われており、Orion フィルタに登録すべき新しいキーワードの探索が人的または時間的リソースの問題から困難であり、キーワード自体もまた表現の複雑性から推測が困難であるという問題があった。本システムはこの問題を解決するために、過去の投稿データを Orion フィルタに検知された投稿と内容が類似しているかに基づいて分類し、Orion フィルタに検知された投稿を含めたそれらのグループ間でのキーワードの出現頻度の差異に着目することで、フィルタに追加すべき未知のキーワード候補を抽出するものである。本手法によってキーワードの選定にかかる時間を削減し、また監視者が独自に推測できないような複雑なキーワードの発見が可能となった。

2 ピグパーティのサービス健全化に対する取り組み

アバターコミュニティアプリである「ピグパーティ」は、サイバーエージェントのメディア事業が提供するサービスの一つである。ピグパーティではユーザーが自身のアバターを作成することでエリアやお部屋と呼ばれる仮想空間上やサービス内限定のコミュニティで他ユーザーとのテキストチャットや画像投稿によるコミュニケーションを行うことが可能であり、本サービスはソーシャルネットワーキングサービスとしての側面を有している。

しかし、このようなコミュニケーション中心のサービスを運営する上では悪意あるユーザーによる他ユーザーへの誹謗中傷や個人情報の聞き出し、そのほか法令に反する表現を含むコンテンツの投稿等の問題が発生することが考えられる。し

たがってピグパーティでは、それらの行為を利用規約で禁止するとともにユーザーの投稿を監視することでサービスの健全性の維持に努めている。加えてピグパーティは利用者の 60% 超を 10 代が占めているという特徴があるため、前述の問題は未成年の誘い出し被害や児童ポルノ事犯の防止等の観点からも対応は必須である。近年 SNS 利用に起因する事犯の被害児童数が増加傾向にあり [2]、サービス内での未成年者の安全を確保することは企業の枠を超えインターネット全体の健全化を促す上でも大きな課題であると言える。

ピグパーティ内では主に未成年の利用者に向け、サービス内で他ユーザーに対して外部の SNS の ID や連絡先等の個人情報を教えたり実際に会ったりすることで生じる危険性や犯罪やトラブルの被害に遭う可能性について啓蒙活動を行う^{*4}とともに、前述の総合監視基盤システム Orion を通じて実際にそれらの事犯につながる可能性のある投稿をはじめとした利用規約に反する投稿が行われていないか 24 時間体制での有人監視を行っている^{*5}。

2.1 ピグパーティの監視体制・Orion による対策

ピグパーティではユーザーがテキストチャット等に投稿したテキストに対してキーワードマッチによるフィルタを適用することによって監視を行っている。監視体制について図 5.1 に示す。

キーワードフィルタには「投稿 NG ワードフィルタ」と「抽出ワードフィルタ」の 2 種類が存在する。前者の「投稿 NG ワードフィルタ」はピグパーティがサービス内で独自に有しているフィルタであり、健全性を維持する上で不適切と定められた投稿禁止キーワードが登録されている。これによって作成した投稿にそれらのキーワードが含まれる場合、ユーザーは投稿の送信を行うことができない。後者の「抽出ワードフィルタ」は Orion のフィルタであり、ユーザーの投稿が「投稿 NG

^{*3} <https://lp.pigg-party.com/>

^{*4} <https://lp.pigg-party.com/rules>

^{*5} ピグパーティではメッセージの作成時に、必要に応じて会話の内容を確認している旨のポップアップが出現する

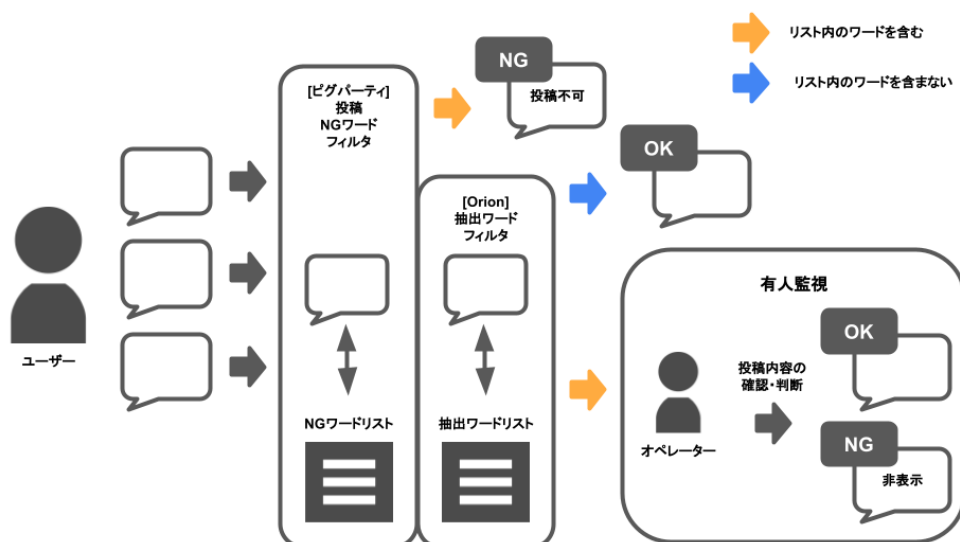


図 5.1 ピグパーティの監視体制

ワードフィルタ」を通過した後に適用される。このフィルタには、通常の投稿にも含まれることがあるが不適切な投稿の中で用いられる可能性が高いため監視対象となっている抽出キーワードリストがあらかじめ設定されており、それらのキーワードリストのうちいずれかを含む投稿が行われたことが検知されると、監視オペレータがその投稿内容を確認し、必要に応じて非表示等の対応を行う。多くの場合キーワードの不健全性はその投稿の文脈に左右されるため、「投稿 NG ワードフィルタ」と「抽出ワードフィルタ」を比較したとき含まれるキーワード数は「抽出ワードフィルタ」の方が多い傾向にある。

2.2 Orion フィルタの課題

Orion フィルタで監視対象となるキーワード群を適切に設定し定期的に更新することはフィルタを継続的かつ有効に運用していく上で必要不可欠な作業であり、ピグパーティでも新たに不適切なキーワードが発見された場合、適宜フィルタの更新を行っている。

しかし、偶発的な発見を除いてフィルタに追加すべき新しいキーワードの探索は運用者の手作業によるものとなっており、大きく 2 つの問題が存在した。

1 つ目はキーワードの探索に掛る人的・時間的リソースの不足である。ピグパーティにはユーザーによって日々大量の投稿が行われている。そのため、フィルタによって検知されていない投稿を含めた全投稿を対象にそれらを目視で確認しながら新しいキーワードを探索することは多くの人手または時間が必要となり、実行することは困難である。

2 つ目はキーワードの複雑性である。ユーザーの中には自身の投稿が非表示になる等によりフィルタによって監視されていることを察知すると、それ以降は伏字や婉曲な表現等を含む別のキーワードを用いることによってフィルタによる検知を回避しようとするケースも散見される。したがって不適切な投稿の中で用いられるキーワードは日々変化しており、加えて多くの場合においてユーザーが考案した造語や比喻等の複雑な表現となる傾向にある。そのためそれらのキーワードを Orion の運用者が観測することなく独自に思い付くことは困難であると言える。

しかし、新しいキーワードを監視対象に含められない状態が続くことは本来監視すべき投稿を見逃している状況であり、サービスの健全性を維持するためには新キーワードの探索とフィルタへの追加を継続的に行っていくことは必須の課題であると言える。したがって、人的・時間的リソー

スを抑え、また運用者の推測のみに頼らないキーワード探索方法が必要である。次節ではこの問題を解決するために実際にピグパーティ導入されているシステムについて紹介する。

3 Orion フィルタのキーワード候補推薦システム

ピグパーティでは Orion のフィルタを適切に更新するために、投稿ログからフィルタに追加する候補となるキーワードを機械学習を用いて自動で抽出するシステムを構築している。システムの概略図を図 5.2 に示す。

本システムではまず、過去に Orion フィルタとオペレーターのチェックにより不適切と判定された投稿（以降 NG メッセージ）とそれ以外のすべての投稿（以降 OK メッセージ）を用い、OK メッセージの中から NG メッセージに内容が類似する投稿（以降 NG 類似メッセージ）を検出する。その後これらのメッセージ群を比較することで、NG メッセージや NG 類似メッセージに特徴的に現れ、かつ Orion フィルタに未登録であるキーワードを抽出する。これらの処理は図 5.2 のフィルタキーワード候補抽出に該当する。

次に、抽出されたキーワードは実際にそれらのキーワードが含まれる会話ログのサンプルとともにアノテーションツールである OrionAnnotator [3] へ送信される。オペレータは OrionAnnotator 上で会話ログを確認しながら各キーワードが Orion で監視されるべきものであるかアノテーションを行い、フィルタキーワードとして採用されたキーワードについては Orion の抽出キーワードリストに追加する対応を行う。

これらの一連の処理は定期的に実行されるため、オペレータは大量の会話ログを確認することなく自動的に抽出されたキーワード群から新しく Orion のフィルタに追加すべきキーワードを精査・更新できる。次節ではシステム中のフィルタキーワード候補抽出手法についてより詳しく紹介する。

3.1 フィルタキーワード候補抽出

本節で説明するフィルタキーワード候補抽出の処理は以下のような 6 つのステップに分かれている。

1. ログを基に過去一定期間の投稿を NG メッセージと OK メッセージとに分類する
2. キーワード辞書を作成し各メッセージをキーワードに分割する
3. 各メッセージから文章ベクトルを作成する
4. OK メッセージを NG 類似メッセージと NG 非類似メッセージに分類する
5. 各メッセージ群内でのキーワードの出現頻度の違いを基に各キーワードの NG スコアを計算する
6. NG スコアがあらかじめ定めた条件に当てはまるキーワードをフィルタキーワード候補として抽出する

ステップ 1 ではキーワード候補抽出の元データとなるメッセージデータセットを作成するために NG メッセージと OK メッセージを投稿ログから作成する。ピグパーティ上の投稿は多くの場合相手ユーザーとの会話形式でやりとりされる。そのため各投稿の文字数は短い傾向にあり、1 つの投稿のみでは以降の処理を行うための十分なキーワードが含まれない可能性が高い。この問題を解決するために、ある投稿を基準としたときその前後で発信された 5 投稿のテキストをその投稿のテキストに結合することで 1 つのメッセージとした。NG メッセージについては Orion で不適切と判定されたメッセージを基準にこの処理を行い、OK メッセージについては Orion フィルタをクリアした（不適切と判断されていない）投稿を NG メッセージと同数サンプリングした後同様の処理を行った。

ステップ 2 ではステップ 1 で作成したメッセージデータからキーワード辞書を作成し、辞書を基に各メッセージをキーワードに分割した。本ステップではキーワード辞書の作成に Sentence-

*6 <https://github.com/google/sentencepiece>

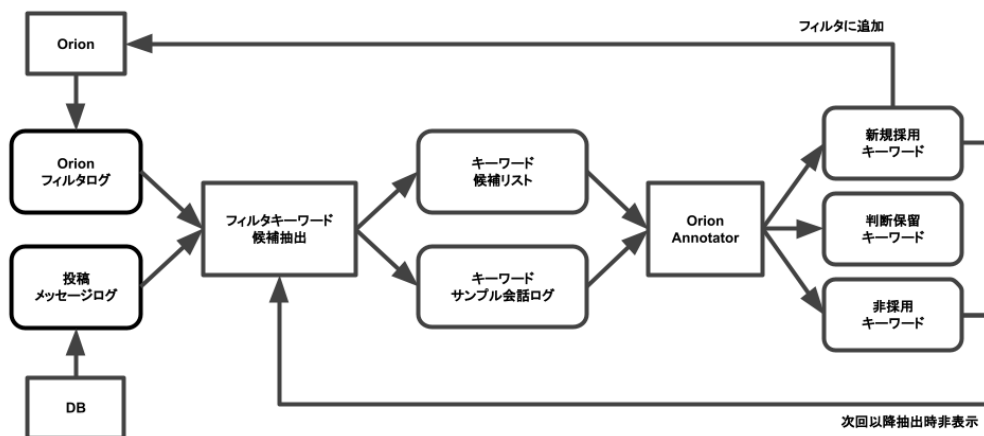


図 5.2 キーワード候補推薦システム概略図

Piece*⁶を用いた。SentencePiece はテキストから直接キーワード分割を行う手法である。ピグパーティがコミュニケーションアプリケーションであるという特性上ユーザーの投稿には口語調の語が多く含まれており、既存の辞書ベースの分割手法では適切なキーワード分割が困難であった。また、前述した通り不適切なキーワードについてもユーザーが考案した造語等が含まれている可能性があるため、既存の辞書ではなく SentencePiece を用いてメッセージから直接キーワード辞書を作成することで一般の辞書には含まれないような語彙の抽出を目指した。

ステップ 3 ではステップ 2 でキーワードに分割した各メッセージから文章ベクトルを作成する。文書ベクトルの作成には文章の分散表現を得られる手法である Doc2Vec [4] を用い、各メッセージ d_n に対して Doc2Vec のアルゴリズムである PV-DM と PV-DBOW を適用してそれぞれ 50 次元のベクトルに変換し、式 5.1 の通り 2 つのベクトルを連結して 100 次元のベクトル v_n とした

$$v_n = \text{concat}(v_{\text{pv-dm}}(d_n), v_{\text{pv-dbow}}(d_n)) \quad (5.1)$$

ステップ 4 では作成した文章ベクトルを用いて OK メッセージと NG メッセージの類似度を計算することで、OK メッセージから NG メッセージ

と内容が類似する投稿 (NG 類似メッセージ) を検出する。これは既存の NG キーワードを含んでおらず従来では Orion フィルタを通過していた OK メッセージであっても Doc2Vec によるメッセージの分散表現が NG メッセージと類似していれば、そのメッセージは潜在的な NG キーワードを有しているという仮定に基づいている。類似度の計算には式 5.2 で表されるユークリッド距離を用い、OK メッセージ i と NG メッセージ j の全組み合わせについて距離 u を計算する。

$$u(i, j) = \sqrt{(v_{i_1} - v_{j_1})^2 + \dots + (v_{i_{100}} - v_{j_{100}})^2} \quad (5.2)$$

各 OK メッセージはいずれかの NG メッセージとの距離 u が閾値以下であった場合に NG 類似メッセージと分類され、いずれの NG メッセージとの距離も閾値を上回った場合 NG 非類似メッセージとして分類される。

ステップ 5 では NG メッセージ、NG 類似メッセージ、NG 非類似メッセージの 3 つのメッセージ群でのキーワードの出現頻度の偏りを利用して各キーワードの NG スコアを計算する。キーワード w が NG メッセージ群 D_{NG} で出現する確率を $P(w|D_{\text{NG}})$ 、NG 類似メッセージ群 D_{SIM} で出現する確率を $P(w|D_{\text{SIM}})$ 、NG 非類似メッセージ群 D_{OK} で出現する確率を $P(w|D_{\text{OK}})$ としたとき、

NG スコアは以下の式 5.3, 5.4, 5.5 で与えられるように 3 通りの方法でそれぞれ計算される。なお、式中の ϵ は任意の十分小さい数とする。

$$\text{Score}_A(w) = P(w|D_{\text{NG}}) \log \frac{P(w|D_{\text{NG}})}{P(w|D_{\text{OK}}) + \epsilon} \quad (5.3)$$

$$\text{Score}_B(w) = \left| P(w|D_{\text{SIM}}) \log \frac{P(w|D_{\text{SIM}})}{P(w|D_{\text{NG}}) + \epsilon} \right| \quad (5.4)$$

$$\text{Score}_C(w) = P(w|D_{\text{SIM}}) \log \frac{P(w|D_{\text{SIM}})}{P(w|D_{\text{OK}}) + \epsilon} \quad (5.5)$$

式 5.3 で計算される Score_A は各キーワードが NG 非類似メッセージ群よりも NG メッセージ群での出現頻度が高いほど大きくなる。式 5.4 で計算される Score_B は各キーワードの出現頻度が NG メッセージと NG 類似メッセージで同等であるほど小さくなる。式 5.5 で計算される Score_C は各キーワードが NG 非類似メッセージ群よりも NG 類似メッセージ群での出現頻度が高いほど大きくなる。

ステップ 6 ではステップ 5 で計算した NG スコアを用い、以下の条件のいずれかに該当するキーワードを抽出する。

- $\text{Score}_A > 0$ かつ Score_A 上位 K 位以内のキーワード (K はあらかじめ定めた任意の数)
- $\text{Score}_A > 0$ かつ Score_C かつ Score_B 下位 L 位以内のキーワード (L はあらかじめ定めた任意の数)

抽出されたキーワードから重複を除いたキーワード群がフィルタキーワードの候補となる。

3.2 OrionAnnotator の活用

前述の手法で抽出されたキーワードは実際にこれらのキーワードが含まれる会話ログのサンプルとともに OrionAnnotator にインポートされ、オペレータによって実際に Orion フィルタに登録すべきであるか精査される。OrionAnnotator は秋葉原ラボが開発したアノテーションシステムであり、当社データ基盤と連携しながらさまざまなタスクに対するアノテーションをインタフェース上で

可能にしている。本タスクでは候補となったキーワードごとにアノテーションのエントリが作成され、作業者はアノテーション画面上でキーワードが含まれる実際の会話ログを確認しながら、そのキーワードを Orion フィルタに採用するかどうか、「採用」、「非採用」、「保留」の中から該当するラベルを選択していくことで効率的に決定できる。また、アノテーション結果は当社のデータ基盤に保存され、そのデータを次回以降のキーワード抽出時に利用できる。

4 実運用と結果

現在ピグパーティではキーワード候補推薦システムを用いて隔月で約 500 個の候補キーワードの精査を OrionAnnotator 上で行っている。候補キーワードを 500 個に絞ることで運用担当者 1 名でも精査可能な作業量となっており、これによって定期的なフィルタ更新が従来よりも容易となった。このシステムは 2019 年から継続的に運用されており、現在まで一度の抽出で平均約 10 個のキーワードが新規にフィルタに登録されている。

5 おわりに

本稿ではピグパーティにおける Orion フィルタの候補キーワード抽出システムについて紹介した。本手法では過去の投稿データを Orion フィルタに検知された投稿と内容が類似しているかに基づいて分類し、Orion フィルタに検知された投稿を含めたそれらのグループ間でのキーワードの出現頻度の差異に着目することでフィルタに追加すべきキーワード候補の抽出を行う。これによってオペレータは大量の会話ログを確認することなく自動的に抽出されたキーワード群の中から新しく Orion のフィルタに追加すべきキーワードを選別し、フィルタを更新できる。また、キーワードの精査に OrionAnnotator を用いることでより効率的な作業を可能にしている。

今後の課題としては、不適切なキーワードの傾向は日々変化していることから、定期的なキーワー

ド抽出を行う本システムよりもさらにリアルタイム性の高いキーワード検出システムの開発が挙げられる。冒頭で述べたとおり不適切なコンテンツへの対応はサービスや企業の枠を超えたインターネット全体の健全化を促す上での大きな課題であり、今後も継続的に取り組んでいくことが必要であると考えられる。

参考文献

- [1] Y. Fujisaka, Orion: An Integrated Multimedia Content Moderation System for Web Services, The Annual ACM International Conference on Multimedia Retrieval (ICMR), Industrial Session, 2018.
- [2] 警察庁生活安全局少年課, 令和2年における少年非行, 児童虐待及び子供の性被害の状況, 2021, <https://www.npa.go.jp/publications/statistics/safetylife/R2.pdf>
- [3] 上辻 慶典, 大内 裕晃, 角田 孝昭, 数見 拓朗, 善明 晃由, Orion Annotator: 機械学習を支えるアノテーションシステム, ソフトウェアエンジニアリングシンポジウム 2019, 2019.
- [4] Le, Quoc, Tomas Mikolov, Distributed representations of sentences and documents, International conference on machine learning. PMLR, 2014.



Media Data Tech Studio

6

国際的質問紙調査の実施に関するテクニカルノート

執筆者 森下 壮一郎

概要 本稿では、パーソナルデータの利活用の社会的受容性に関する研究のために実施した質問紙調査における実際的な留意点についてテクニカルノートとして述べる。具体的には、質問紙調査自体がパーソナルデータに相当する可能性に着目して行った匿名でのデータ収集や本人対応の実践について説明する。さらに研究倫理の観点からのインフォームドコンセントについて、共同研究を実施する各研究機関の倫理規定や、論文投稿先として想定される学会の倫理綱領を念頭に置きながら実施したことを述べる。

Keywords 質問紙調査, 研究倫理, 個人情報保護, GDPR, インフォームドコンセント

1 はじめに

人々の生活の多くがデジタルプラットフォームの上でも営まれるようになり、さらに日々のデータ収集および分析技術の発達によって、世界的にパーソナルデータの利活用が進んでいる。2000年代初頭にはすでにE-コマース分野や、広告分野における利活用が一般的になっており [1, 2], 昨今では金融分野（フィンテック）や、医療分野における利活用も盛んである [3, 4].

金融関係のデータは金融分野に、医療関係のデータは医療分野で活用するのが素朴な形であるが、必ずしもデータの種類と利用先のドメインとは対応していない。たとえば医療情報を広告のために共有した事例は同意を得ていなかったこともあり社会問題となった [5]。日々の技術の発達により、たとえばソーシャルネットワークングサービスの“like”から性的指向などのセンシティブな属性を推知するなどの思いも寄らないデータ分析法が生み出されている [6]。仮にこのような技術を応用した事業を行っても、利用者に受容されないであろう。

なお人工知能技術の社会的影響の研究として、江間らは多様なステークホルダーへの質問紙調査を実施しており、情報技術を活用する主体（国や企業、研究者など）に応じて悪用の意図や管理能力についての観念が異なることが示されている [7]。データ利活用においても、データの種類と利用先のドメインに加えてデータを利活用する主体に応じて社会的受容性は異なり、それぞれのパターンについての利用者の考えを知ることは、社会的に受容される態様でデータを利活用する上での指針となるはずである。

以上の考えの下に秋葉原ラボでは、パーソナルデータを利活用する主体と利用目的に応じた社会的受容性について Web アンケート調査を実施して、分析結果の对外発表を行った [8–10]。いずれもオープンアクセスであるので興味関心のある向きには適宜ご参照いただきたい。

本稿ではこれらの成果発表の補遺として、質問紙調査の実施における実際的な留意点をテクニカルノートとしてまとめる。

2 調査の枠組み

デジタルプラットフォームは国家をまたいで展開されているものも多くあり、データを利活用する主体が国内の事業者か海外の事業者かによって受容性は異なるであろう。また国内で開発および展開されて受容された事業でも文化や法制的異なる国では受容されなかったり、あるいは逆に海外で受容されていることが国内で受容されないということも考えられる。このような国家間の関係についても知見を得るために、文化や法制的の違いを考慮に入れて、日本だけでなく海外の国家でも質問紙調査を実施した。具体的には日本に加えて、アメリカ、イギリス、フランスを対象に調査を行った。

2.1 個人情報保護法制に関する検討

調査に着手する前に、質問紙調査自体がパーソナルデータの取得に相当する可能性を考慮に入れながら、各国の個人情報保護法制に関する検討を行った。具体的には日本の個人情報保護法に加えて、EU 一般データ保護規則（GDPR: General Data Protection Regulation）とカリフォルニア州消費者プライバシー法（CCPA: California Consumer Privacy Act）を遵守しながら実施する手続きについて検討した。

なお本調査の目的では最終的に統計情報として結果が得られればよいので、データは匿名でかまわない。もし匿名でデータを収集できれば、そのデータはパーソナルデータに相当しないので、結果的に個人情報保護法や GDPR、CCPA に準拠することになる。しかしながら、素朴なデータ収集の方法ではデータが匿名であることの要件を満たすことができない。たとえば GDPR では、仮に個人を特定できなくとも、Cookie や IP アドレスなどの端末を識別する情報が含まれる場合でもパーソナルデータに該当する（すなわち匿名ではない）と見なされる。インターネット上でブラウザを介した通信を行う以上、実施者に取得意思はなくとも、サーバのログなどとしてこれらの情報を取得

してしまうことは避けられない。

この問題についてはカリフォルニア大学バークレー校の資料 [11] を参照しながら対処した。これはクラウドソーシングプラットフォームを活用することで、端末を識別する情報が調査実施者の意図に反して取得できてしまう問題を回避するものである。次項で詳細を述べる。

2.2 調査の手続き

サードパーティーのクラウドソーシングプラットフォームとアンケート収集プラットフォームを介してデータ収集を行う。具体的には次の3つの手順からなる。

1. 当社がクラウドソーシングプラットフォームに対して、国ごとに絞り込んだアンケート募集を依頼する。
2. クラウドソーシングプラットフォームは依頼を受けて登録者にアンケート募集の連絡（アンケート収集プラットフォームへの誘導）をする。
3. クラウドソーシングプラットフォームからの募集に基づいて、回答者は国ごとに作られている回答ページで質問紙調査に回答する。

割り当てはクラウドソーシングプラットフォームにすでに登録されている情報に基づいて行われるが、その割り当てをするのはクラウドソーシングプラットフォームなので、当社はその情報にアクセスしない。アンケート収集プラットフォーム上では、国ごとにアンケートページを用意することで、国ごとに収集された質問紙調査の結果が得られる。

以上のように、割り当てをする事業者と、アンケートを収集する事業者、アンケート結果を受領する事業者とがそれぞれ別々になっていることで、実施者が受け取るデータはアンケートの回答結果のみとなる。

前述のバークレー校の資料の“2. External HITs”の項目で、クラウドソーシングプラットフォームに Amazon Mechanical Turk (MTurk)

*1 <https://www.mturk.com/>

*1 を採用した場合における匿名とみなせる要件 a, b, c, d が挙げられている。

- a. アンケートで MTurk の識別子 (ID: identifier) を取得しない
- b. アンケート内容で個人を特定しない
- c. 外部サイトで IP アドレスが収集されない
- d. 報酬支払いのための完了コードをユーザー毎にユニークにしない

ここで「外部サイト」とはアンケート収集プラットフォームを指す。

要件 a については、アンケートで MTurk の ID の入力を求めないことで対応した。

要件 b については、アンケートの質問項目に個人を特定するものを含まないようにしたうえで、MTurk から個人や端末を特定する情報（氏名、住所、メールアドレスも含めて、IP アドレス、Cookie 等）を受領しないことで対応した。

要件 c については、外部サイトから IP アドレスを受領しない設定で調査を実施することによって対応した。本調査では、アンケート収集プラットフォーム（外部サイト）として SurveyMonkey^{*2}を採用しており、このサイトでは「匿名回答」の設定を施すことで「作成者」^{*3}からは IP アドレスを参照できなくなる。なお、このサイトはプライバシーポリシーに示されている範囲で Cookie を収集しているが、Cookie の情報は「作成者」に提供されることはない。

要件 d については、以下のように完了コードの発行方法を工夫することで対応した。MTurk からの報酬の支払いは、アンケート収集プラットフォームから発行されるコードを回答者が入力することで行われる。このとき回答者について一意なコードを発行していると、このコードが識別子として働いて個人を識別できてしまう。これを防ぐために回答者について一意とならないコードを発行する。ところで本調査のようなオンライン調査では、報酬のみを得るために問題文を読まずに回答する

という行動が見られることがあり、このような行動はサティスファイスと呼ばれる [12]。これに対しては、サティスファイスを判別するための問題文（アテンションチェック [13]）を用意しておき、アンケート収集プラットフォーム上で判別を行って、やはり回答者について一意とならないように報酬が無効になるコードを発行することで対処した。

以上により GDPR が要求する匿名でのデータ収集の要件を満たす方法でデータ取得を行った。

3 本人対応

収集したパーソナルデータに対する本人による開示および利用停止の請求や苦情等に対応することを「本人対応」と呼ぶ。本人対応についての一定の義務は、OECD 8 原則の「公開の原則」や「個人参加の原則」の規定が根拠となっている [14]。また、個人情報保護法でも保有個人データについて開示や利用停止に応じる義務があり、さらに苦情処理や紛争解決に関する努力義務が定められている。

一方で、当該研究で収集したデータは法令上の個人データではなく、また前述の手続きにより名実ともに匿名で調査を実施している以上、自社で完結した本人対応を行うことはできない。しかしながら、当社の事情にかかわらず本人からの開示および利用停止の請求や苦情等は発生しうる。

当該研究で実施した調査の手続きにより調査が匿名が行われることは前節で述べたとおりであるが、説明と同意文（後述）においてここまで詳細な説明をしているわけではないので、問い合わせが生じること自体は避けられない。

なお当該調査では問い合わせへの対応が必要になる事態を想定していたので、次に示す手立てを事前に講じていた。

- 本人対応ができないことを説明と同意文に示す
- 事前に可能な手段について検討しておく

以下、それぞれの手立てについて詳述する。

^{*2} <https://www.surveymonkey.com/>

^{*3} SurveyMonkey のプライバシーポリシーの用語。アンケートを集める主体。本ケースでは当社。

3.1 本人対応ができないことを説明と同意文に示す

本人とデータとを対応付けることが不可能であることを本人が納得しているならば、そのようなデータを開示や利用停止できないことで本人が不利益を感じる事態は考えにくい。また質問紙調査への参加は任意であるので、事前に疑念がある場合は参加しないということで問題はない。しかしながら当該研究の調査ではアテンションチェックをパスした場合に限って報酬を支払っているのに、本当は真摯に質問に答えていたにもかかわらず、ケアレスミスなどの何らかの理由でアテンションチェックをパスしなかった場合は、報酬を受け取れないことで事後に不満が生じてしまう。

もし、本人とデータとを対応付けることが可能ならば、実際に行われた回答がケアレスミスなのかサティスファイス回答なのか、個別の回答を精査すればある程度は判断可能である。本人がそうのように考えて回答の精査を求めることは考えられる状況である。それが実際には不可能であることについて事前の理解を得るために、説明と同意文には次の文言を入れた。

なお、この調査は匿名で行われますので。どの回答を誰が入力したかを知る方法がありません。したがって、個別の回答についてのご要望にはお応えできないことをあらかじめご了承ください。

3.2 事前に可能な手段について検討しておく

説明と同意文の内容を承知したうえで回答したとしても、アテンションチェックをパスしなかったことにより実際に報酬が支払われない事態に直面すると、納得のいかない心境に回答者が至ることは十分に考えられる。このような状況で回答者から回答の精査を要求された場合は、精査を申し出た回答者による回答を特定する必要が生じる。

回答者の募集はクラウドソーシングプラット

フォーム上で行っているため、クラウドソーシングプラットフォーム上で回答者を識別して、回答者が入力したコードと紐付けることは可能である。また、受領したデータには本調査のみがスコープとなる回答者 ID が付与されており、アンケート収集プラットフォーム上で個人を識別することは可能である。しかしながら前述した匿名でのデータ収集の要件は、クラウドソーシングプラットフォーム上の ID とアンケート収集プラットフォーム上の ID とを紐付けられないようにすることであり、本調査もこの要件を満たしながら実施している以上、調査の過程で両者を紐付けられる情報を当社は得られない。回答時間などの傍証から絞り込むことは可能かもしれないが、回答を一意に特定できるとは限らない。

ただし、本人からアンケート収集プラットフォームに問い合わせたアンケート収集プラットフォーム上での ID を知ることができれば、その ID によって受領した回答が一意に特定できる。調査の匿名性が破られることを回答者に説明したうえで、このような対応を依頼するという手段は考えられる。ただしアンケート収集プラットフォームがそのようなリクエストに必ず応えるとは限らない。

4 インフォームドコンセント

過去に医療分野において、ヒトを対象とした実験で著しく非倫理的な実験が行われたことから、1964 年に採択されたヘルシンキ宣言によりインフォームドコンセントの原則などの研究倫理における基本理念が示された。

インフォームドコンセント (informed consent) は一般に「説明と同意」と訳され、日本では特に医療行為に対して、医師から十分な説明を患者が受けて、その医療行為を受けるか否かを選択して同意することを指す。研究倫理の文脈では、広くヒト (人) を対象とした実験^{*4}で、研究者から十分な説明を実験協力者が受けて、実験に協力するか否

^{*4} 学術用語としては「ヒト」と「人」とは使い分けされる。生物としての種を指すときは「ヒト」、社会的な存在としてみなすときは「人」と表記される。

かを選択して同意することを指す。

本研究をマーケティング調査と見なすと、インフォームドコンセントの必要までは感じられないかもしれない。またヒト（人）を対象とする研究でも侵襲性が低い場合では、研究倫理の観点においても、必ずしも医療分野で必要とされる水準を満たす必要はないかもしれないという議論もある [15]。しかしながら本研究は国際的な調査を行うものであり、また外部研究機関との共同研究として実施することから、各研究機関の倫理規定や、論文投稿先として想定される学会の倫理綱領に満たす形で実施することが望ましかった。また前述したように匿名で調査を実施していることと、そのために可能な本人対応に制限があることを通知する必要があった。

以上に基づいて、説明と同意文を和文（日本向け）と英文（アメリカ、イギリス、フランス向け）とでそれぞれ用意した。なおセクションバイアスを軽減することを目的に、英語と日本語とで説明の粒度を同じ水準にした。

5 おわりに

本稿では、パーソナルデータの利活用の社会的受容性についての国際的な質問紙調査の実施にあたって、質問紙調査自体がパーソナルデータに相当する可能性に着目して、匿名でのデータ収集や本人対応の実践を示した。特に匿名でデータを収集しているにもかかわらず、本人対応が必要となる状況に直面したことと、その相反する要求を満たすために考えられる手立てについて検討した結果とを、調査研究の実践としてここに報告することには十分な意義があると考える。

さらに研究倫理の観点からのインフォームドコンセントについて、共同研究を実施する各研究機関の倫理規定や、論文投稿先として想定される学会の倫理綱領を念頭に置きながら実施したことを述べた。私企業は研究倫理の観点で倫理審査を行う部署を持たないことが多いので、一貫した基準

を設けることが難しく、共同研究先や論文投稿先などの文脈に依存しがちである。なお近年では研究機関が社会貢献の一環として、技術的なもの以外の諸々の課題の解決のために産業界をサポートする取り組みがなされている。たとえば大阪大学 社会技術共創研究センター^{*5}は、倫理的・法的・社会的課題（ELSI: Ethical, Legal and Social Issues）に関する教育を産業界に展開していくことを表明している。社会に受け入れられる事業を営むためにも、私企業の研究開発組織においては、今後も新技術の研究開発において学术界と連携しながら諸問題に取り組んでいくことが必要であろう。

参考文献

- [1] Ben Schafer, Joseph Konstan, and John Riedl. E-Commerce recommendation applications. *Data Mining and Knowledge Discovery*, Vol. 5, No. 1, 2000.
- [2] Lee Sherman and John Deighton. Banner advertising: Measuring effectiveness and optimizing placement. *Journal of Interactive Marketing*, Vol. 15, pp. 60–64, December 2001.
- [3] 堀内進之介. 予測アルゴリズムに基づく与信管理の功罪—ライフチャンスへの影響とその対策の有効性の検討を中心に—. *社会情報学*, Vol. 8, No. 2, pp. 169–185, 2019.
- [4] Wullianallur Raghupathi and Viju Raghupathi. Big data analytics in healthcare: promise and potential. *Health Information Science and Systems*, Vol. 2, No. 1, February 2014.
- [5] Madhumita Murgia and Max Harlow. How top health websites are sharing sensitive data with advertisers, November 2019. <https://www.ft.com/content/0fbf4d8e-022b-11ea-be59-e49b2a136b8d>, 2021 年 7 月閲覧.

^{*5} <https://elsi.osaka-u.ac.jp/>

- [6] Michal Kosinski, David Stillwell, and Thore Graepel. Private traits and attributes are predictable from digital records of human behavior. *Proceedings of the National Academy of Sciences*, Vol. 110, No. 15, pp. 5802–5805, 2013.
- [7] 江間有沙, 秋谷直矩, 大澤博隆, 服部宏充, 大家慎也, 市瀬龍太郎, 神崎宣次, 久木田水生, 西條玲奈, 大谷卓史, 宮野公樹, 八代嘉美. 運転・育児・防災活動, どこまで機械に任せるか: 多様なステークホルダーへのアンケート調査. 情報管理, Vol. 59, No. 5, pp. 322–330, 2016.
- [8] 森下壮一郎, 高野雅典. 個人データ利活用における利用主体と利用目的に応じた社会的受容性. 人工知能学会全国大会論文集(JSAI2020), 3N5OS11b01, 2020.
- [9] 森下壮一郎, 高野雅典, 武田英明, 高史明, 小川祐樹. 個人データ利活用の類型に応じた社会的受容性の質問紙調査. 人工知能学会全国大会論文集(JSAI2021), 2C4OS9b03, 2021.
- [10] Soichiro Morishita, Masanori Takano, Hideaki Takeda, Faiza Mahdaoui, Fumiaki Taka, and Yuki Ogawa. Social acceptability of personal data utilization business according to data controllers and purposes. In *13th ACM Web Science Conference 2021*, pp. 262–271, 2021.
- [11] UC Berkeley Committee for the Protection of Human Subjects. Mechanical Turk (MTurk) for Online Research, January 2020.
- [12] 三浦麻子, 小林哲郎. オンライン調査モニタの satisfice に関する実験的研究. 社会心理学研究, Vol. 31, No. 1, pp. 1–12, 2015.
- [13] Weiping Pei, Arthur Mayer, Kaylynn Tu, and Chuan Yue. Attention Please: Your Attention Check Questions in Survey Studies Can Be Automatically Answered. In *Proceedings of The Web Conference 2020*, pp. 1182–1193, Taipei Taiwan, April 2020.
- [14] OECD 理事会. OECD 理事会勧告 8 原則 (昭和 55 年 (1980 年) 9 月), Sep. 1980. https://www.soumu.go.jp/main_sosiki/gyoukan/kanri/oecd8198009.html, 2021 年 7 月閲覧.
- [15] 大谷卓史, 大澤博隆, 壁谷彰慶, 神崎宣次, 久木田水生, 西條玲奈. 意思決定支援としての研究倫理 ～AoIR 倫理ガイドラインの原理と倫理分析～. 信学技報, Vol. 121, No. 119, pp. 182–189, Jul. 2021.



2021 年（7 月迄）

書籍・解説記事

- 高野 雅典：“「その他」を研究する”，人工知能，特集「編集委員からの抱負と提言 2021」，Vol.36, No.3, p.322, 2021.
- 鳥海 不二夫（編著），石井 晃（著），岡田 勇（著），上東 貴志（著），小林 哲郎（著），榊 剛史（著），笹原 和俊（著），高野 雅典（著），瀧川 裕貴（著），常松 淳（著），三浦 麻子（著），水野 貴之（著），山本 仁志（著），吉田 光男（著）：『計算社会科学入門』，丸善出版，2021.

論文誌・国際会議（査読付き）

- Masanori Takano, Fumiaki Taka, Soichiro Morishita, Tomosato Nishi, and Yuki Ogawa: “Three clusters of content-audience associations in expression of racial prejudice while consuming online television news,” PLOS ONE, Vol.16, No.7, e0255101, 2021.
- Soichiro Morishita, Masanori Takano, Hideaki Takeda, Faiza Mahdaoui, Fumiaki Taka, and Yuki Ogawa: “Social acceptability of personal data utilization business according to data controllers and purposes,” The 13th International ACM Conference on Web Science in 2021 (WebSci’21), pp.262–271, 2021.
- Masanori Takano, Yuki Ogawa, Fumiaki Taka, and Soichiro Morishita: “Effects of incidental brief exposure to news on news knowledge while scrolling through videos,” IEEE Access, 2021.
- Kenji Yokotani and Masanori Takano: “Differences in Victim Experiences by Gender/Sexual Minority Statuses in Japanese Virtual Communities,” Journal of Community Psychology, Vol.49, Issue 6, pp.1598–1616, 2021.
- Masanori Takano and Kenichi Nakazato: “Difference in communication systems explained by balance between edge and node activations,” Journal of Physics: Complexity, Vol.2, No.2, 025013, 2021.
- Kenji Yokotani and Masanori Takano: “Social Contagion of Cyberbullying via Online Perpetrator and Victim Networks,” Computers in Human Behavior, Vol.119, 106719, 2021. 西 朋里, 小川 祐樹, 高 史明, 高野 雅典, 森下 壮一郎, 服部 宏充：“ネットテレビのニュース番組に投稿される視聴者コメントの道徳性に基づく分析”，人工知能学会論文誌，Vol.36, Issue 1, 2021.

国内学会/セミナー

- Masanori Takano: “Racism and News on the Internet: An Example of Japanese Online Television,” Prejudice, Discrimination, and Stereotypes in Northeast Asia, Featured Symposia on Asian Association of Social Psychology (AASP), 2021.
- 森下 壮一郎: “データサイエンスの実践と法・倫理 ～ アバターコミュニティアプリを一例として ～”, 技術と社会・倫理研究会 (SITE), 2021.
- 森下 壮一郎, 高野 雅典, 武田 英明, 高 史明, 小川 祐樹: “個人データ利活用の類型に応じた社会的受容性の質問紙調査”, 第 35 回人工知能学会全国大会, 2C4-OS-9b-03, 2021.
- 高野 雅典, 小川 祐樹, 高 史明, 森下 壮一郎: “ニュース動画のリニア配信とオンデマンド配信における利用スタイルの分析”, 第 35 回人工知能学会全国大会, 1D2-OS-3a-01, 2021.
- 小川 祐樹, 高野 雅典, 森下 壮一郎, 高 史明: “Twitter におけるニュースツイートの閲覧と動画視聴の関連性”, 第 35 回人工知能学会全国大会, 1D4-OS-3c-03, 2021.
- 涌田 悠佑, Michael Mior, 善明 晃由, 佐々木 勇和, 鬼塚 真: “時刻変化するワークロードのための NoSQL スキーマのオフライン最適化”, 情報処理学会第 83 回全国大会, 4L-04, 2021.
- 数見 拓朗, 白井 徳仁, 善明 晃由: “構造化データの自動抽出に向けた変換処理フレームワークの提案”, 第 13 回データ工学と情報マネジメントに関するフォーラム (DEIM), E33-4, 2021.
- 涌田 悠佑, Michael Mior, 善明 晃由, 佐々木 勇和, 鬼塚 真: “時刻変化するワークロードのための NoSQL スキーマの最適化”, 第 13 回データ工学と情報マネジメントに関するフォーラム (DEIM), B14-1, 2021.
- 大内 裕晃: “【技術報告】データと処理の依存関係を整理する機械学習モデル管理基盤の紹介”, 第 13 回データ工学と情報マネジメントに関するフォーラム (DEIM), E11-5, 2021.

2020 年

書籍・解説記事

- 水上 ひろき, 熊谷 雄介, 高野 雅典, 藤原 晴雄: 『データ活用のための数理モデリング入門』, 技術評論社, 2020.

論文誌・国際会議 (査読付き)

- 西口 真央, 鳥海 不二夫, 高野 雅典: “メタデータを利用したソーシャルメディア内グループのネットリスク検知”, 情報処理学会論文誌, Vol.61, No.10, pp.1639–1646, 2020.
- Teruyoshi Zenmyo, Noriyuki Watanabe, Yusuke Wakuta, and Makoto Onizuka: “Schema Mapping between Logical and Internal Layers for NoSQL Applications,” The 1st Workshop on Distributed Infrastructure, Systems, Programming and AI (DISPA), 2020.
- Masanori Takano and Kenichi Nakazato: “Emergence of the tradeoff law of social relationships in artificial societies driven by dual memory mechanisms,” Artificial Life (ALIFE2020), 2020.
- 涌田 悠佑, 善明 晃由, 松本 拓海, 佐々木 勇和, 鬼塚 真: “Secondary Index を活用する NoSQL スキーマ推薦によるクエリ処理高速化”, 情報処理学会論文誌: データベース (TOD), Vol.13, No.1, pp.20–32, 2020.
- Makoto Takeuchi: “Epidemic modeling of viral music diffusion,” NetSciX (Poster; Abstract) 2020.

- Masanori Takano and Kenichi Nakazato: “A balance between edge- and node-excitation mechanisms realizes the difference of communication systems”, NetSciX (Poster; Abstract) 2020.

国内学会/セミナー

- 飯塚 隆介, 鳥海 不二夫, 西口 真央, 高野 雅典, 吉田 光男: “デマの訂正が社会的混乱に与える影響の分析”, 第 18 回データ指向構成マイニングとシミュレーション研究会 (DOCMAS), 2020.
- 小川 祐樹, 高野 雅典, 森下 壮一郎, 高 史明: “オンラインニュース動画の視聴と政治的知識の関連性”, 第 18 回データ指向構成マイニングとシミュレーション研究会 (DOCMAS), 2020.
- 高野 雅典: “アバターコミュニケーションによって現実の社会を補間する”, 第 11 回横幹連合コンファレンス, A-4-5, 2020.
- 横谷 謙次, 高野 雅典: “バーチャルコミュニティが性的少数者の性的不平等と精神的健康に及ぼす影響”, 日本認知・行動療法学会 第 46 回大会, P026, 2020.
- 高野 雅典: “サイバーエージェントにおける計算社会科学”, 第 34 回人工知能学会全国大会, インダストリアルセッション, 2020.
- 森下 壮一郎, 高野 雅典: “個人データ利活用における利用主体と利用目的に応じた社会的受容性”, 第 34 回人工知能学会全国大会, 3N5-OS-11b-01, 2020.
- 武内 慎: “感染症モデルを用いた音楽拡散現象の分析”, 第 34 回人工知能学会全国大会, 2E6-GS-5-05, 2020.
- 高野 雅典, 小川 祐樹, 高 史明, 森下 壮一郎: “インターネットテレビ局におけるニュースチャンネルのユーザ体験が政治関心・ニュース知識に与える影響”, 第 34 回人工知能学会全国大会, 1L5-GS-5-04, 2020.
- 森下 壮一郎: “データサイエンスの実務とビジネス・倫理”, 技術と社会・倫理 (SITE) 研究会 「AI と倫理」シンポジウム (招待講演), 2020.
- 善明 晃由: “AbemaTV におけるコンテンツメタデータマネジメントに関する取り組み”, 第 12 回データ工学と情報マネジメントに関するフォーラム (DEIM), 2020.
- 新平和礼, 善明 晃由, 斎藤 貴文: “ストリーム処理を考慮したレイヤ化されたデータ転送管理システム”, 第 12 回データ工学と情報マネジメントに関するフォーラム (DEIM), 2020.
- 高野 雅典, 小川 祐樹, 高 史明, 森下 壮一郎: “インターネットテレビにおけるニュースへの偶発的接触が政治関心とニュース知識に与える影響”, 第 4 回計算社会科学ワークショップ (CSSJ2020), 2020.

CyberAgent 秋葉原ラボ 技術報告 Volume 4

発行日 2022年1月1日 初版
発行所 CyberAgent 秋葉原ラボ 技術報告編集委員会
〒101-0021
東京都千代田区外神田1丁目18番13号
秋葉原ダイビル 13階

編集 上辻 慶典
上田 紗希
善明 晃由
橋爪 友莉子
福田 鉄也
森下 壮一郎
數見 拓朗
松井 美帆

デザイン 柴 尚子
新聞 絵美
横山 恵
Design Factory



Media Data Tech Studio

