

第7章 研究報告

7.1 小野 謙二（応用データ科学研究部門 教授）

二次元 Porous Disk モデルの性能検証と応用解析

1. はじめに

再生可能エネルギーへの移行が進む中、風力発電所（Wind Farm）の大規模化が急速に進展している。しかし、上流風車が形成する wake によって、下流風車の発電量低下および乱流増加による疲労負荷増大が発生し、風力発電所全体の効率・信頼性向上を妨げている。

従来手法には以下の課題が存在する。

手法	特徴	課題
JENSEN/GAUSSIAN モデル	高速	複雑な wake 干渉を表現困難
高忠実度 CFD	高精度	計算コストが極めて高い

特に、大規模風車配置最適化（WFLO: Wind Farm Layout Optimization）では、多数回評価が必要となるため、3D CFD は現実的ではない。本研究では、その中間的手法として Porous Disk モデルを採用し、実験との比較検証、複数風車 wake 干渉解析を実施している。

2. 解析手法

風洞実験は、九州大学 応用力学研究所（RIAM）の大型境界層風洞で実施された。条件は以下の通り。

- 風洞サイズ $15.0 \times 3.8 \times 2.0$ m
- ロータ径 $D = 1$ m
- 流入速度 $U_0 = 7.4$ m/s
- Reynolds 数 $Re = 2.9 \times 10^5$
- 流入条件 乱流格子による乱流生成

乱流格子を利用することで、実環境に近い wake 分布計測を行っている。

3. 数値解析

数値解析には、Lattice Boltzmann Method（LBM）、Cumulant Collision Modelを採用した。風車モデルには Porous Disk モデルを利用し、ロータによる運動量損失を流れ場へ導入している。また、流入乱流は NREL の Turbsim により生成され、乱流強度（TI）は 8% に設定した。解析は以下の 2 段階で実施した。

- 1) 単一風車 wake 検証
 - 下流 $X/D=3$
 - wake 速度分布測定
 - CFD と風洞実験比較

2) 二基風車干渉解析

- 第二風車を $X/D=5$ 下流に配置
- Y 方向へ横移動
- 各位置で発電量評価

上記のケースで wake 干渉と発電性能の関係を調査した。

4. 結果

1) 単一風車 wake 検証

図1(112ページ)では、風洞実験と数値解析結果の比較が示されている。結果として、CFD結果と実験結果が良好に一致。wake 速度欠損を適切に再現でき、側方加速流も再現できていることが確認された。図2では、時間平均流れ場が示されており、ロータ直後の大きな速度低下があり、両側に形成される加速流が明瞭に可視化されている。

2) 二基風車干渉解析

本研究で特に興味深い点は、「wake 側方の加速流」に着目した点である。研究では、「下流風車を加速流領域へ配置すれば、発電量が増加するのではないか」という仮説を検証している。図3の結果からは、下流風車を横方向 $1D \sim 2D$ に配置すると発電量が最大化されることが示された。特に、単独風車以上の出力、wake 回避だけでなく加速流利用という興味深い現象が確認されている。これは従来の単純 wake モデルでは見落とされやすい複雑流動現象であり、本 PD モデルの有効性を示している。

5. 結論

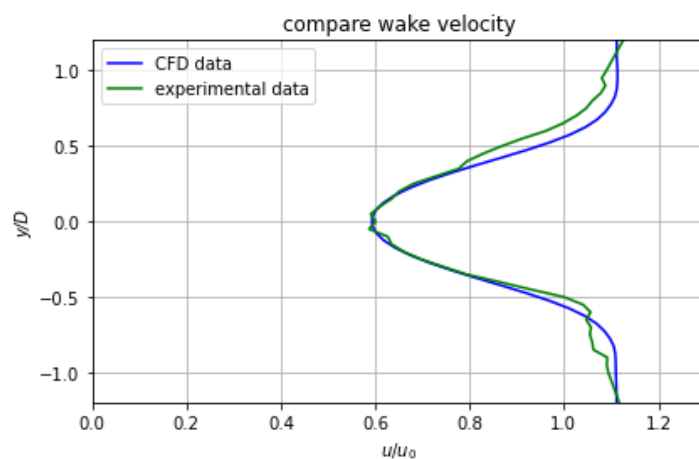
本研究では、二次元 Porous Disk モデル、LBM ベース CFD、風洞実験比較を通じて、wake 解析モデルの性能検証を実施した。主な成果は以下の通り。

- 単一風車 wake を高精度に再現
- 実験結果との整合性を確認
- 側方加速流を適切に表現
- 二基風車干渉を解析
- 特定配置で発電量向上を確認

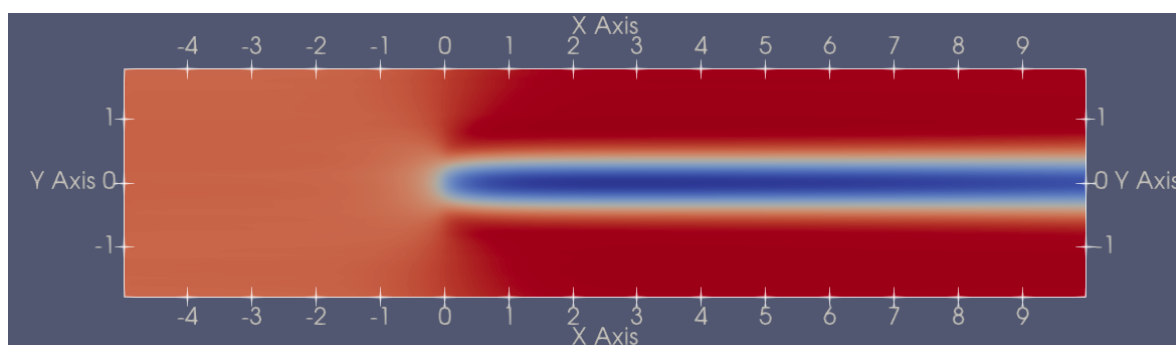
今後は、多数風車への拡張、実地気象データ利用、実風力発電所への適用、配置最適化問題への展開を目指している。

参考文献

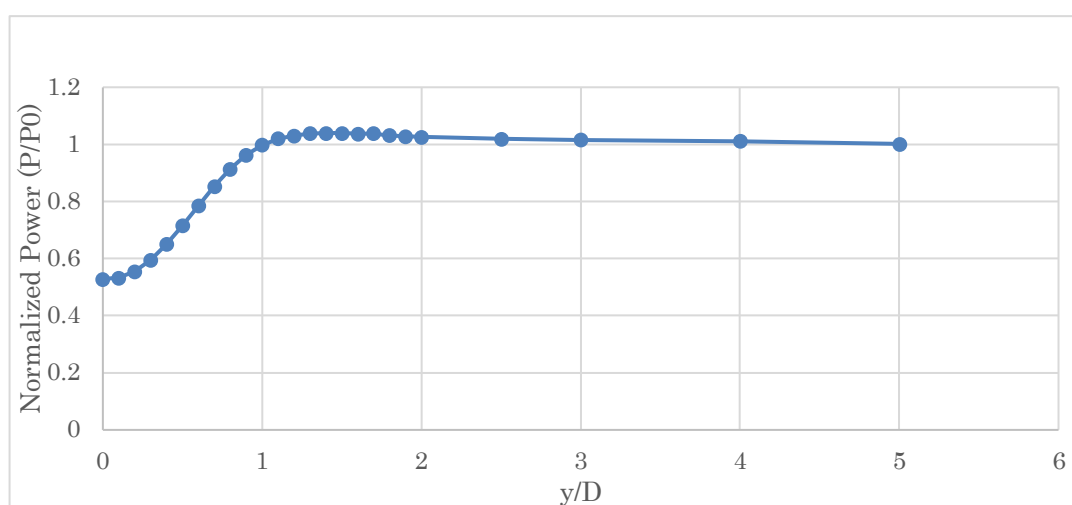
- [1] Shibuya K., Uchida T., Muramoto Y., Watanabe K., Ohya Y., “風車後流中の流動現象解明のための風洞実験ならびに cfd その1：複数風車後流の相互干渉による影響,” 風力エネルギー利用シンポジウム, vol. 42, pp. 124–127, 2020.
- [2] Asakura M., Ichikawa H., Uvhide T., “エジプト・スエズ湾風力発電所における風車ウェイク相互干渉現象の解明に向けた研究開発,” 風力エネルギー利用シンポジウム, vol. 46, pp. 145–148, 2024.



(図1) Comparison of experimental and simulation results



(図2) Time-averaged flow field at T=6000



(図3) Power generation and wind turbine location

7.2 伊東 栄典（応用データ科学研究部門 准教授）

海賊版マンガ対策のためのマンガ画像分類

1. はじめに

近年、著作権を無視した海賊版マンガサイトや最新話マンガのスライド動画投稿が問題になっている。特にここ数年では、マンガのページをそのまま転載したものだけでなく、著作権侵害を判定する AI を出し抜くために画像を意図的にずらしたり、一部分だけを切り出して連続的に表示する形式で転載するものも出現している。後者のような無断転載画像も著作権侵害であると正しく認識するには、画像内のキャラクター情報を拾い集めて作者やタイトルを断定し、さらにその登場順や登場頻度を整理することで巻数とページ数を特定する必要がある。

海賊版マンガサイトおよびマンガ動画の無断投稿を抑止するため、我々は今までマンガのページ画像からマンガのタイトル・巻数・話数を推定する判別手法を研究してきた [1]。ページ画像だけでなく、切り取った一部分のマンガ画像からでも、切り抜き元マンガの情報を機械的に抽出したい。マンガ原作を多数作成した小池一夫氏が述べるように、マンガにおいて登場キャラクターは非常に重要である [2]。我々は有名なマンガのキャラクターであれば、一目見ただけで、それが誰であるか識別できる。

本研究では、マンガのページ画像からキャラクター顔部分検出と、キャラクター顔画像が誰なのかを判定する手法を検討した。まず用いたデータを説明し、キャラクター顔部分検出に適用した手法と検出精度、顔画像からのキャラクター判別に適用した手法と検出精度を報告する。その後、これらの手法を用いて実際に市販のマンガ内のキャラクターについてキャラクター判別を行うことができるのか検証する。最後に、この手法がキャラクター判別以外の情報特定にも利用できるのかを検証する。

2. 用いたデータ

マンガページからのキャラクター顔部分判別機と、キャラクター顔画像からのキャラクター判別機を機械学習で作成するためには、なるべく沢山の「マンガページ画像」と登場する「キャラクター顔画像」、その画像の「キャラクター名」を持つデータを用意する必要がある。機械学習用のデータとして、本研究では Manga109 データセットを用いた。Manga109 は、日本漫画のメディア処理の学術研究使用のために、東京大学の相澤・山崎・松井研究室によりとりまとめられたデータセットである [7, 8]。Manga109 は、日本のプロ漫画家により描かれて出版された 109 作品（タイトル）について、1 冊ずつ選ばれた計 109 冊のマンガ画像を保持する。109 冊のマンガは 1970 年代～2010 年代に出版された漫画で、幅広い対象読者層やジャンルを網羅している。画像に加え 1 冊毎に作品のメタデータがあり、かつ全ページ画像に注釈 (annotation) が付与されている。メタデータには作品タイトル・作者名・出版社・出版年に加え、登場する全キャラクターのリスト、各ページ（見開き）のリストも含む。各ページの注釈は、登場キャラクターの顔領域、キャラクターの全身領域、印字文字のテキスト、手書き文字のテキスト、コマ枠の領域情報を含む。メタデータと注釈は XML 形式で記述されている。アノテーションデータから欲しい対象物の座標データを取得し、画像データから該当部を取り出すことで任意のオブジェクトの切り抜き、保存が可能である。表 1 に XML 注釈の統計を示す。

(表1) Manga109 のXML 注釈数

項目	数値	項目	数値
巻数	109	顔	118,593
ページ数	10,602	全身	157,234
注釈	10,130	台詞	147,887
		コマ	103,850

3. マンガページ画像からのキャラクタ顔部分検出

顔部分の検出に使うモデルとして、以下の3つを検証した [3,5]。animeface (Viola-Jones) モデルは OpenCV の Viola-Jones 物体検出フレームワークに基づいたカスケード分類器の一種で、nagadomi 氏により作成・提供されている [9]。animeface (CascadeClassifier) モデルは、シンプルで高速なアニメ顔検出に適したモジュールである。OpenCV の CascadeClassifier をベースに開発されており、アニメ顔の特徴を学習したカスケード分類器を用いて、画像から顔を検出する [10]。Faster R-CNN (Detectron2) は、画像内の物体を検出する真相学習モデルの一種で、特にリアルタイムな物体検出において高い性能を発揮する。Detectron2 は、Facebook AI Research が開発した、PyTorch ベースの物体検出ライブラリである。以降、animeface(Viola-Jones) を af-VJ、animeface(CascadeClassifier) を af-CC、FasterR-CNN を F_R-CNN とする。

3 種類の検出器について、マンガキャラクタに対する検出率を比較した。Manga109 データセットの images を入力として用いて、検出器が出力するキャラクタ顔の位置情報と、XML 注釈が含む顔の位置情報を比較し、その重なり具合 (IoU:Intersection over Union) が指定値を超えていたものを検出成功例とする。検出結果を表2に示す。

(表2) 判別精度 (IoU=0.5)

Method	Precision	Recall	F1	AP
af-VJ	0.9813	0.1502	0.2648	0.1865
af-CC	0.9927	0.3387	0.5051	0.4937
F_R-CNN	0.7661	0.8306	0.7971	0.8112

公開されていた af-VJ と af-CC はどちらも Precision は高いけれど Recall は低く、検出結果はほぼ顔部分であるけれど、ページに登場する顔のすべてを検出できていないことを示す。一方、F_R-CNN は Precision と Recall のどちらも高い値を出している。F_R-CNN は、真正面の顔だけでなく横顔や領域の半分がフキダシで隠れている顔、小さく描写されている顔まで検出している。さらに F_R-CNN は顔検出の正確さを確信度で数値化する。今回の実験では確信度 50% 以上として判別精度を算出した。確信度を調整することで、顔でない部分に誤反応している検出候補の除外が可能かもしれない。

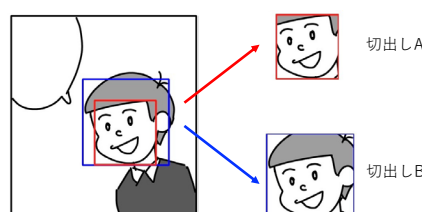
4. キャラクタ顔画像からのキャラクタ判別

次に、ある顔画像が、マンガの誰であるかの判別を試みた [4,6]。対象データは Manga109 データセットに登場するキャラクタである。XML 注釈を用いて抽出できる顔部分画像の解像度が低い（領域が小さい）キャラクタや、顔画像が少ないキャラクタは、学習や分類に適さない。そこで顔領域が 10000px 以上で、かつ 1 人の顔画像が 50 個以上存在するキャラクタの顔画像のみに限定した。表 6 に、学習に用いた切り出し画像の数、キャラクタ数、などを示す。

（表 3）学習データ数

項目	数値
顔画像の総数	5,348 個
キャラクタ数	61 人
キャラクタ事の平均画像数	87.67 個
作品あたりのキャラクタ数	3.59 人

人間がマンガのキャラクタを判別する場合、髪型は重要である。そこで、XML 注釈通りの顔画像と、髪の毛を含むように拡張した顔画像を用意した。注釈 XML の顔座標で切出した画像群を A（切出し A）、上左右を拡張して顔部分を切出した画像群を B（切出し B）とする。切出し B では上に 20%、左に 10%、右に 10% のピクセルを増やして切り出した。図 1 に A と B の切出し方法を示す。



（図 1）顔画像の切出し方

マンガに登場するキャラクタの顔画像を学習させるための CNN アーキテクチャとして、画像認識の分野で高性能を発揮する ResNet と DenseNet の 2 つを用いた。ResNet では、ある層に与えられた信号を、上位層の出力に追加するスキップ接続により深いネットワークを訓練できる。DenseNet は ResNet から開発されている。DenseNet は前の各層からの出力すべてが後の層への入力として用いられる。ResNet には層の数により ResNet34, ResNet50, ResNet101 などがあり、DenseNet も DenseNet121, DenseNet169, DenseNet201 などが有る。今回 ResNet50 と DenseNet121 のみを試した。

学習モデル構築のプログラムは Python 言語を用いて作成した。機械学習および画像判別精度の出力部分は、生成 AI の ChatGPT（無料版）が回答したコードを援用して実装した。ライブラリとして Tensorflow と Keras を用いている。

ResNet50 で切出し A と B の顔画像の分類を学習させたモデルを R-A, R-B とする。同様に DenseNet121 で学習させたモデルを D-A, D-B とする。これらの学習モデルを表 4 に示す。

(表4) 学習で得たキャラクタ顔画像判別モデル

画像群	ResNet50	DenseNet121
切出し A	R-A	D-A
切出し B	R-B	D-B

学習で得た4つのモデルの判別結果の Recall, Precision, F1-Score, Accuracy を表5に示す。表5を見ると、切出しAとBのどちらでも DenseNet121 の精度が高い。次にキャラクタの顔部分の切出し方について、座標通りの切出しAと、上左右に拡張した切出しBを比較すると、DenseNet121を使用したモデルの両方で精度が高い。

ResNetよりDenseNetの方が、切出しAよりBの方が判別精度が高い理由を調べるため、4つのモ

(表5) キャラクタ顔画像判別の精度

モデル	Precision	Recall	F1	Accuracy
R-A	0.8705	0.8393	0.8462	0.8575
R-B	0.8981	0.8826	0.8864	0.8904
D-A	0.9053	0.8921	0.8951	0.9050
D-B	0.9284	0.9181	0.9205	0.9269

デルが、キャラクタ顔画像のどの部分に着目しているかを調べた。各モデルの attention-heatmap を作成した。図2に attention-heatmap を示す。図2を見ると、DenseNetはResNetと比べてキャラクタの目に attention が集めている。髪の毛部分も含むようにした切出しBの画像群を学習したD-Bは、主に前髪と目を含めた部分に attention を集めている。髪部分を含むようにしたことで、キャラクタの目に加えて髪を特徴として取り入れることができ、判別精度が精度が向上したのであろう。

5. おわりに

本論文では、マンガページからの顔部分検出とマンガに登場するキャラクタの顔部分画像からのキャラクタの判別を試みた。Manga109 データセットに含まれるキャラクタの顔画像を機械学習させて、判別した。顔部分検出では、Faster R-CNN による検出が最も精度が高い。本稿では数値を述べていないものの、未学習のマンガでも高い精度でキャラクタの顔部分を検出できている。キャラクタ判別では、キャラクタ画像に髪の毛部分を含めるように拡張し画像で学習した DenseNet121 モデルの精度が高い。学習済みの既知のキャラクタであれば高精度で判別できる。

今後は、未学習キャラクタの顔部分検出と、キャラクタ推定を進めて行きたい。またキャラクタの属性（性別、年齢世代、アクセサリ等）の数値化によるキャラクタ類似度算出や、マンガの作者ごとの絵柄の違いを判別するモデルを検討したい。

参考文献

- [1] 山仲一颯, 伊東栄典: 大規模マンガ画像判別手法の検討, 情報処理学会火の国情報シンポジウム 2024, A4-1, 2024.
- [2] 小池一夫: 小池一夫のキャラクタ新論 ソーシャルメディアが動かすキャラクタの力, ゴマブックス, 2017.
- [3] 霜出 秀太, 添田 碧人, 伊東 栄典: マンガキャラクタ画像に基づく作者および作品の機械的判別に関する考察, 人工知能学会第 133 回知識ベースシステム研究会 (SIG-KBS133), pp.33-38, 2024.
- [4] 添田 碧人, 伊東 栄典: マンガ分析のための類似度に基づくキャラクタ顔画像クラスタリングの検討, 人工知能学会第 133 回知識ベースシステム研究会 (SIG-KBS133), pp.39-44, 2024.
- [5] 霜出秀太, 伊東栄典: マンガ顔画像からのキャラクタ判別に関する検討, 2024 年度 (第 77 回) 電気・情報関係学会九州支部連合大会, 09-1P-10, pp.188-189, 2024.
- [6] 添田碧人, 伊東栄典: マンガページ画像からのキャラクタ顔検出に関する検討, 2024 年度 (第 77 回) 電気・情報関係学会九州支部連合大会, 09-1P-09, pp.186-187, 2024.
- [7] Manga109, <http://www.manga109.org/ja/index.html> (アクセス日 05-Aug-2024).
- [8] K. Aizawa, et.al.: Building a Manga Dataset “Manga109” with Annotations for Multimedia Applications, IEEE MultiMedia, vol.27, no.2, pp.8-18, 2020.
- [9] nagadomi: lbpcascade animeface, https://github.com/nagadomi/lbpcascade_animeface (アクセス日 05-Aug-2024).
- [10] <https://github.com/opencv/opencv/tree/master/data/haarcascades> (アクセス日 05-Aug-2024).

7.3 内林 俊洋（応用データ科学研究部門 准教授）

GTFS を使用した地域公共交通向けバスロケーションシステム強化の検討

1. はじめに

近年、日本の地域公共交通は、特に地方部において人口減少や高齢化の影響を受け、バス路線の廃止や縮小による交通空白地域の拡大が深刻な問題となっている。このような状況に対応するため、多くの地方自治体が自らを運行主体とするコミュニティバスなどの地域公共交通の運行を進めているが、限られた予算や人材の中での運行管理が求められている。一方で、情報通信技術の発展により、公共交通の利便性を向上させる新たな手法が注目されている。特に GTFS（General Transit Feed Specification）などの標準フォーマットを活用したデジタル技術の導入が進められているが、地域公共交通ではその活用が十分でない。

そこで、これまでに地域公共交通の課題に対処するため、「地域公共交通データの整備支援」と、「ICT を活用した地域公共交通の支援」を行ってきた。特に、「ICT を活用した地域公共交通の支援」で行なっているバスロケーションシステムは、モバイル端末のアプリケーションを用いてリアルタイムでバスの位置情報を収集し、乗客や運行管理者に提供することで利便性を向上させるものである。また、同じアプリケーションを用いる乗降客数調査システムは、運転手がアプリを通じて乗降者数を記録し、そのデータを分析することで運行ルートやダイヤの最適化を図ることを目的とする。これらのシステムは、地域公共交通の利便性向上だけでなく、自治体の運行管理業務の効率化にも貢献することが期待される。

しかし、こうしたシステムの導入に際しては、データ管理の安全性や改ざん防止の課題が存在する。特に GTFS のような公共交通データは、乗客の移動情報に直結するため、不正アクセスや改ざんのリスクが懸念される。本研究では、ブロックチェーン技術を活用し、GTFS データの一部をブロックチェーン上に保存することで、データの信頼性向上と改ざん耐性の強化を行う。ブロックチェーンは分散型台帳技術の一種であり、データの改ざんを防ぎつつ透明性の高い管理を可能にする。

2. 関連技術と課題

拠点である福岡県内の地方自治体が運行する地域公共交通を中心に支援活動を行ってきた。特に、「ICT を活用した地域公共交通の支援」では、タブレットを使ったアンケート調査、モバイル端末のアプリケーションから得た位置情報を使用したバスロケーションシステム、モバイル端末のアプリケーションを使用した乗降客数調査、車内モニタを使用した案内表示システム、そして主要バス停に設置したモニタを使用したデジタルサイネージを開発運用している（表1）。これらの支援は、我々が開発しているモバイル端末で動作する SHINGU アプリと複数の AWS のサービスで構成される ACE クラウドで実現している。ACE クラウドは、利用者向けのバスロケーション情報を表示するマップを提供する EC2 と、「SHINGU」アプリと通信するための REST API を含むサーバーレス環境の2つで構成される。SHINGU アプリは、位置情報をサーバへ送信することによるバスロケーション機能、そして運行路線ごとに乗降客数を記録する機能がある。この SHINGU アプリをインストールしたスマートフォンをコミュニティバス内に設置し、位置情報を収集する。コミュニティバス車内から USB 経由で給電しているが、夜中などの運行時間外にバッテリーを消費しないように営業時間外は自動的に電源をオフにするように設定している。位置情報を収集するためには、SHINGU アプリの設定画面で事業者 ID とバス番号を設定する。その後、アプリの起動時または運行時の乗降客人数入力画面にて3秒ごとに ACE クラウドへ位置情報として、現在時刻、事業

者 ID、バス番号、バッテリー残量、そして GPS 座標を送信する。乗降客数を集計するために、運転手はアプリ内にて運行開始時に自身の運行する路線を選択し、遷移した専用画面にて乗降客数を記録する。アプリはクラウドへ現在時刻 30 分前後に運行を開始する路線のリストを要求し、返ってきたリストを表示することでこれから運行する路線の選択をスムーズにする。

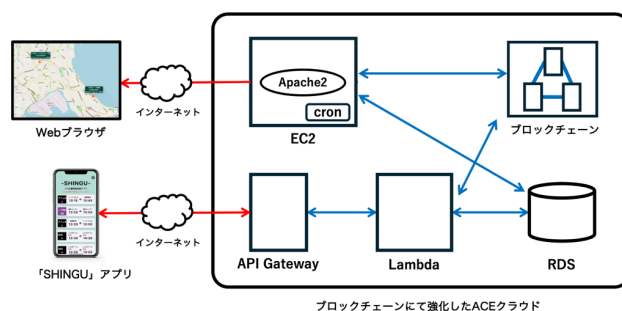
これらの路線情報は、ACE クラウドの RDS に保管されている GTFS のデータを利用している。GTFS はオープンデータとして公開されていることが多く機密性は低いが、我々のシステムのように運行に直結するような使い方をする場合には、サーバへの不正アクセスやサーバ管理者のミスなどによるデータの改ざんの脅威がある。そこで、ブロックチェーンを変更の追跡が可能なセキュアなデータベースとして、GTFS の保管に利用する方法を提案する。

(表 1) ICT による支援を行なっている自治体

福岡県	芦屋町	古賀市
	福津市	荏田町
	新宮町	宗像市
	久山町	嘉麻市
	須恵町	小都市
	築上町	
沖縄県	うるま市	

3. 提案手法

ブロックチェーンは、分散型台帳技術の一種で、データを安全に管理し取引を透明かつ改ざんできない形で記録するためのシステムである。ブロックチェーンは、基本単位であるブロック、ブロックが連続的に繋がっていくチェーン、複数のノードで構成される分散ネットワーク、ネットワーク内のコンピューターでブロックチェーンのデータを保持し取引の検証や新しいブロックの追加を行うノード、そして特定の条件が満たされたときに自動的に実行されるスマートコントラクトで構成される。ブロックチェーンにて強化した提案システムの構成を図 1 に示す。これまでの RDS に加えて、3 台のノードで構成するブロックチェーンがシステムに含まれる。RDS からブロックチェーンに移動するデータは GTFS の一部のみとする。GTFS は、9 種類のファイルで構成されている。この中で改ざんによる運行への影響が大きいと思われる、agency、routes、shapes、stops、そして trips をブロックチェーンに保管する。これらのデータは管理者のみ追加・変更・削除が可能とする。RDS からの変更点としては、以前は直接データベースに書き込んでいたデータは、ブロックチェーンのスマートコントラクトで処理が行われた後に、ブロックが生成されチェーンとして繋がることで保管が完了するため、データを追加・変更・削除するまでに時間がかかると考える。



(図 1) 提案システムの構成

4. 考察

従来のシステムでは管理者のミスや不正アクセスによるデータ改ざんのリスクがあったが、新システムではブロックチェーンに過去のデータ履歴がすべて記録されるため、不正な変更が容易に検出できるようになる。一方で、データの追加・更新にかかる時間が増加するという課題も明らかになった。これは、スマートコントラクトを介したデータ処理が必要となるためであると考えられる。

また、システム全体の運用速度、データの整合性、改ざん耐性について議論する。ブロックチェーン導入による利点として、データの変更履歴がブロックチェーン上に記録されるため、不正なデータ改変を防ぐことが可能となった。そして、運行情報の変更が記録として残るため、自治体や住民がデータの正当性を確認しやすくなった。さらに、特定の管理者に依存しない分散型管理により、サーバ障害や管理者のミスによるデータ損失のリスクが低減した。一方で、ブロックチェーンへのデータ書き込みは即時反映されるデータベースと比較して遅延が発生するため、リアルタイム処理が求められる場面では適用範囲を検討する必要があるというデータ処理速度の低下の問題があった。データの信頼性向上という観点では一定の成果を得られたが、リアルタイム処理の遅延や運用コストといった課題も明らかになった。今後は、運行データの一部のみをブロックチェーンに記録するハイブリッド運用や処理速度を向上させる手法の検討が必要となると考える。

5. おわりに

日本の地域公共交通における課題を踏まえ、ICTを活用したバスロケーションシステムおよび乗降客数調査システムの開発と運用を行っている。本研究は、GTFSデータの管理においてブロックチェーン技術を導入し、データの信頼性向上と改ざん防止を目指した。

公共交通データの安全性確保のため、GTFSの一部をブロックチェーン上に保存する手法を提案した。従来のRDSによる保管では、不正アクセスや管理ミスによるデータ改ざんのリスクが存在していたが、ブロックチェーンを活用することでデータの履歴を透明に管理し、信頼性を高めることができた。一方で、データの追加や更新にかかる時間が増加するという課題も明らかになった。

今後は、より多くの自治体と連携しシステムの適用範囲を広げるとともに、技術面での改善を図ることで地域公共交通の持続可能な発展に貢献していこうと考える。

7.4 殷 成久（教育情報基盤研究部門 教授）

ラーニングアナリティクス：教育ビッグデータの利活用に関する研究

研究背景

教育は、産業発展に伴う社会的ニーズに応じて進化してきた。歴史的に見ても、Industry 1.0 時代の大規模な基礎教育から、情報化時代におけるデジタルリテラシーや創造力の重視に至るまで、各時代の教育変革は新たな教育理念と人材育成目標を伴ってきた。現在の第四次産業革命に対応する Education 4.0 もその延長線上にあり、AI、機械学習、IoT などの技術を活用し、教育の効率化と個別最適化された学習支援を目指している。

その中核的な目標は、従来の知識伝達型教育から、より柔軟で適応性が高く、学習者中心の教育へ転換することである。今後の教育には、単なる知識の伝達にとどまらず、学習者の全人的な成長と持続的な発展能力を重視することが求められる。

こうした背景を踏まえ、本研究では新たな教育理念として「Learning Success」を提唱する。Learning Success は、学習成果だけでなく、学習過程における成長、能力形成、人格的発達、そして社会変化への持続的な適応力を重視する理念である。テクノロジーと学習者中心の教育設計を通じて、学習者が未来社会や産業のニーズに対応できる力を育成することを目的としている。

Learning Success を実現するためには、学習者一人ひとりの理解状況や学習状態を的確に把握し、それに応じた学習支援を提供することが重要である。特に、生成 AI や学習分析、視線計測、ロボットなどの先端技術を活用することで、教材の個別最適化、学習過程の可視化、学習状態に応じたリアルタイム支援が可能になる。これらは、学習者が主体的に学び続け、知識を深く理解し、変化する社会に適応する力を育成するうえで重要な基盤となる。

以上の問題意識に基づき、2025 年度は Learning Success の実現に向けて、主に二つの研究を実施した。第一に、LLM と RAG を活用し、講義資料や参考文献に基づいて信頼性の高い教材を生成するモデルを構築した。第二に、学習者の状態をリアルタイムに推定し、ロボットによって動的に介入する学習支援システムの研究を進めた。これらの研究は、学習者中心の教育を支える技術的基盤を構築し、個別最適化された学習支援を実現するための具体的な取り組みである。以上の背景を踏まえ、Learning Success の実現に向けて、2025 年度は以下の研究を実施した。

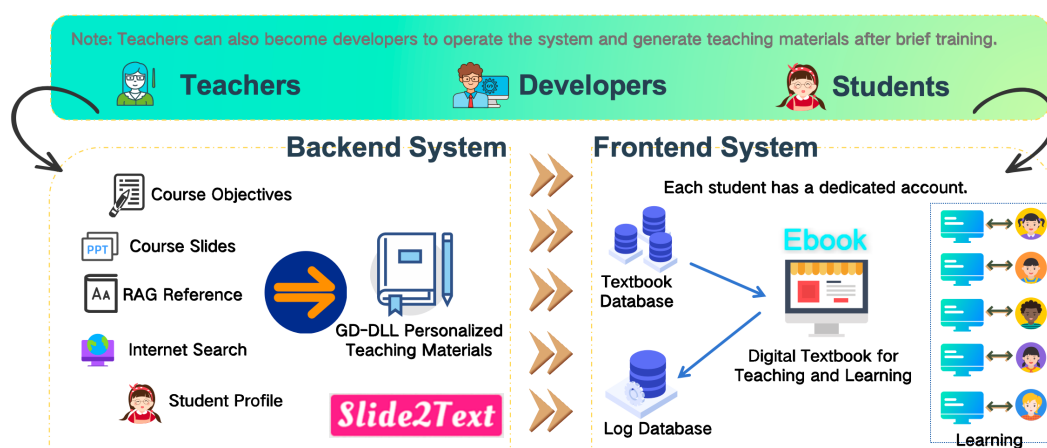
研究実績

1) LLM と RAG による教材生成モデルの構築

2025 年度は、LLM と Retrieval-Augmented Generation (RAG) 技術を活用し、講義スライドと参考文献に基づいて詳細な教材を生成する手法を検討・構築した。従来の講義資料は、キーワードや簡単な説明を中心としたスライド形式が多く、学生が内容を深く学ぶための補足的な学習リソースが不足しているという課題があった。

この課題に対応するため、講義資料からトピックごとに内容を抽出・整理する仕組みを構築するとともに、LLM を用いて講義内容と参考文献に基づく詳細教材を動的に生成する方法を開発した。また、LLM におけるハルシネーションの問題に対応するため、外部情報検索を統合した RAG を導入し、生成教材の正確性と信頼性の向上を図った。

さらに、研究の拡張として、ウェアラブル型視線計測装置を用いて学習者の視線データを収集し、学習（視線）パターンの分析を実施した。今後は、視線データを活用した教材生成手法の検討を進める予定である。

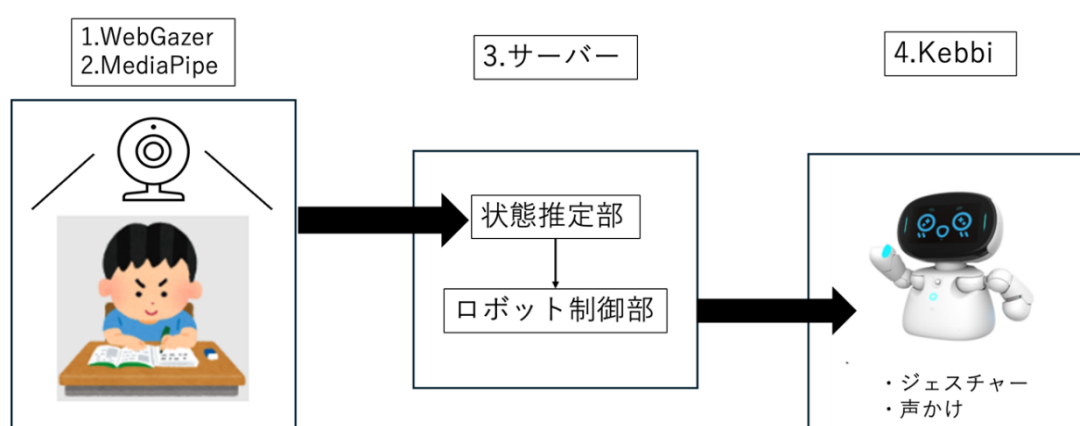


(図1) LLM と RAG による教材生成モデルの構築

2) 状態推定に基づくロボットの動的介入に関する研究

近年、オンライン学習やビデオ教材の普及により、学習者は時間や場所に制約されず柔軟に学習できるようになった。一方で、従来のビデオ教材は一方的であり、学習者の集中力低下や学習内容からの逸脱に対して、リアルタイムに介入することが難しいという課題がある。

そこで本研究では、学習者の状態をリアルタイムで検知し、ロボットを通じて適切な支援を行う学習支援システムを提案する。人間の講師が常に適切なタイミングで声かけやジェスチャーを行うことは難しいが、ロボットは安定した動作とタイミングで介入できる。本システムでは、コンピュータビジョン技術により学習者の視線や姿勢を検知し、集中力の低下が見られた場合にフィードバックを行うことで、学習者の注意を促し、自律的な学習継続を支援する。



(図2) ロボットの動的介入

発表論文等

- [1] Shi, W., Sakaue, Y., Yamashita, Y., Takei, R., Yin, C., & Okada, Y. (2026). Application of interactive 3D content in quality Japanese history education and intangible cultural heritage preservation. *Educational Technology & Society*, 29(1), 268–290. [https://doi.org/10.30191/ETS.202601_29\(1\).SP04](https://doi.org/10.30191/ETS.202601_29(1).SP04)
- [2] Liu, H., Shimada, A., Minematsu, T., & Yin, C. (2025). Building a scaffolding framework with field environment digest system in high school agricultural education. *Interactive Learning Environments*, 1–18. <https://doi.org/10.1080/10494820.2025.2591862>
- [3] Zhao, F., Hwang, G.-J., & Yin, C. (2027). Learning success in the age of generative AI: A new paradigm and theoretical model for education technology. *International Journal of Mobile Learning and Organisation*, 21(2), xxx–xxx. <https://doi.org/10.1504/IJMLO.2027.10077554>
- [4] Su, W.-S., Hwang, G.-J., Yin, C., & Chang, C.-Y. (2025). Exploring gender and educational background effects on 21st-century competencies in robot-integrated STEM education: The role of learning attitudes and satisfaction. *Educational Technology & Society*, 28(4). [https://doi.org/10.30191/ETS.202510_28\(4\).SP11](https://doi.org/10.30191/ETS.202510_28(4).SP11)
- [5] Chang, C.-Y., Lin, H.-C., Yin, C., & Yang, K.-H. (2025). Generative AI-assisted reflective writing for improving students' higher order thinking: Evidence from quantitative and epistemic network analysis. *Educational Technology & Society*, 28(1), 270–285. [https://doi.org/10.30191/ETS.202501_28\(1\)](https://doi.org/10.30191/ETS.202501_28(1))
- [6] Hashiyada, K., Shi, W., & Yin, C. (2025). A framework for using LLMs and RAG to realize the automatic generation of learning materials from lecture slides. In *Proceedings of the 1st International Conference on Learning Evidence and Analytics (ICLEA 2025)*, Fukuoka, Japan, September 5–7, 2025.
- [7] Liu, H., Minematsu, T., Yin, C., Liu, S., Xiong, S., & Shimada, A. (2025). Integrating scaffolding strategies with environmental monitoring systems to enhance learning and practical skills in agricultural education. In *Proceedings of the 1st International Conference on Learning Evidence and Analytics (ICLEA 2025)*, Fukuoka, Japan, September 5–7, 2025.
- [8] Li, H., Ma, B., & Yin, C. (2025). Examining metacognitive difficulties in learning programming: Analysis of student behavior and strategy. In *Proceedings of the 1st International Conference on Learning Evidence and Analytics (ICLEA 2025)*, Fukuoka, Japan, September 5–7, 2025. (Best Short Paper)
- [9] Bai, Y., Zhao, F., Wang, W., & Yin, C. (2025). Multidimensional feature-based textbook difficulty assessment index. In *Proceedings of the 1st International Conference on Learning Evidence and Analytics (ICLEA 2025)*, Fukuoka, Japan, September 5–7, 2025. (Best Poster Paper)
- [10] Su, Y., Xu, Z., Wang, X., Yin, C., & Zhou, J. (2025). An approach to evaluating the role of reading map in supporting self-reflection in e-book environments. In *Proceedings of the 1st International Conference on Learning Evidence and Analytics (ICLEA 2025)*, Fukuoka, Japan, September 5–7, 2025.

- [11] Yoshida, K., Ohno, A., Tani, I., Yin, C., & Kumamoto, E. (2025). A study of a method for role estimation in discussions using large language models. In Proceedings of the 19th International Conference on Innovative Computing, Information and Control (ICICIC 2025), Kitakyushu, Japan, August 26–29, 2025.
- [12] Liu, H., Minematsu, T., Yin, C., Xiong, S., & Shimada, A. (2025). Exploring the relationship between system operation behaviors and learning achievement in agricultural education. In Proceedings of the 33rd International Conference on Computers in Education (ICCE 2025), Chennai, India, December 1–5, 2025.

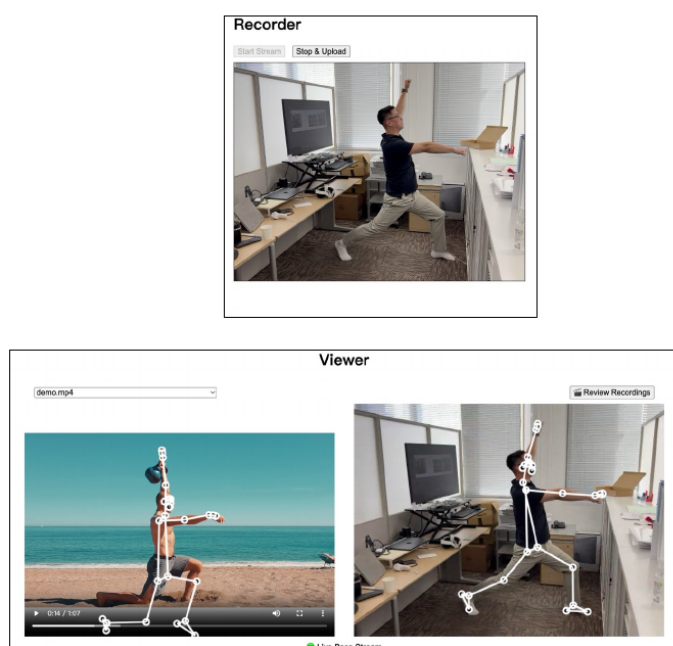
7.5 石 偉（教育情報基盤研究部門 助教）

Computer Vision と 3D 技術を活用した教育支援システムの研究開発

近年、Computer Vision、3D/VR、生成 AI などの技術は、学習者の行動や理解過程を可視化し、個々の状況に応じた教育支援を実現する基盤として注目されています。2025 年度、私は主に、骨格推定を活用したスポーツ学習支援システムの開発、およびインタラクティブ 3D 教材による日本史教育と無形文化遺産継承支援に取り組みました。

1. リアルタイム骨格推定を用いたスポーツ学習支援システムの開発

スポーツのオンライン学習では、学習者が身体の向きを変えた際に指導者映像を見失いやすく、また衣服や身体部位の重なりにより正しいフォームを把握しにくいという課題があります。これに対し、MediaPipe によるリアルタイム骨格推定を用いて、指導者と学習者の動作を骨格情報として可視化する学習支援システムを開発しました。本システムは、サーバ、録画用 Web アプリ、閲覧用 Web アプリから構成され、スマートフォンで撮影した学習者映像と指導者映像をブラウザ上で比較できます。また、VRゴーグルで閲覧することで、身体の向きを変えても指導者映像を常に正面に表示できます（図 1）。

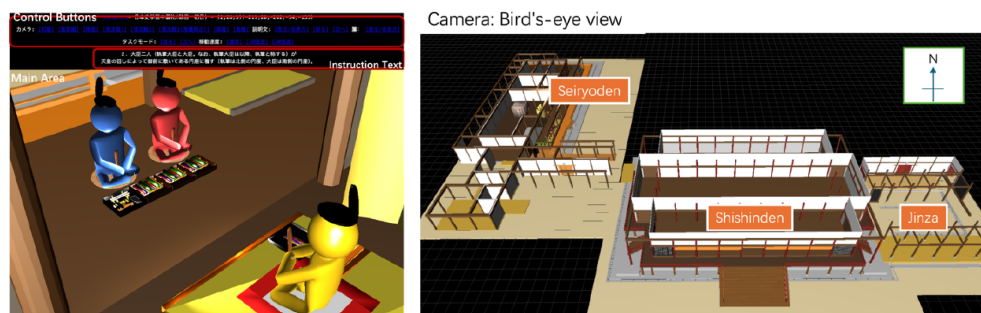


（図 1）骨格推定を用いたスポーツ・ダンス学習支援システムの画面例

12 名の大学生を対象とした先行評価では、従来のオンラインスポーツ学習に対する満足度が低い一方で、本システムの有用性、使いやすさ、学習時間短縮への期待が高く評価されました。今後は、複数カメラによる 3D 骨格推定、多視点録画、自動姿勢評価アルゴリズムを導入し、従来型動画教材や VR 教材との比較実験を進める予定です。

2. インタラクティブ 3D 教材による日本史教育と無形文化遺産継承支援

日本史教育において、儀礼や作法などの無形文化遺産は文字資料だけでは人物配置、道具、動作順序、空間関係を理解しにくいという問題があります。そこで、インタラクティブ 3D 技術を統合した歴史学習コンテンツ開発フレームワークを提案し、平安時代の宮中儀礼「除目」を題材とした 3D 教材を開発しました。Blender で作成した 3D モデルとアニメーションを、HTML5、JavaScript、Three.js を用いて Web 教材として実装し、学習者が視点を自由に切り替えながら儀礼の流れを確認できるようにしました(図 2)。



(図 2) 日本史学習用インタラクティブ 3D 教材の画面例

教育効果を検証するため、九州大学の学生 47 名を対象に比較実験を行いました。参加者は、3D 教材を用いる実験群 23 名と、テキスト型教材を用いる対照群 24 名に割り当てられ、事前と事後のテストとアンケート、インタビューにより評価しました。結果として、3D 教材を用いた実験群は、事後テスト得点において対照群を有意に上回りました。特に、人物配置や空間関係に関する問題で大きな差が見られ、インタラクティブ 3D 表示が複雑な歴史的場面の理解に有効であることが示されました。また、実験群は教材満足度や学習への関心、自己効力感が高く、精神的負荷が低い傾向を示しました。本研究成果は、Educational Technology & Society 誌に掲載されました [2]。今後は、学習ログや視線、行動データの分析、生成 AI を用いた教材作成支援、他地域や他文化の歴史教材への展開を進める予定です。

参考文献

- [1] W. Shi, W. Xiong and Y. Okada, "A System Supporting Sport and Dance Learning via Real-Time Skeleton Estimation," 2025 IEEE International Symposium on Consumer Technology (ISCT), Denpasar, Bali, Indonesia, 2025, pp. 327-331, doi: 10.1109/ISCT66099.2025.11297391.
- [2] W. Shi, Y. Sakaue, Y. Yamashita, R. Takei, C. Yin, and Y. Okada, "Application of interactive 3D content in quality Japanese history education and intangible cultural heritage preservation," Educational Technology & Society, vol. 29, no. 1, pp. 268-290, 2026. doi: 10.30191/ETS.202601_29(1).SP04.

7.6 李 慧勇（教育情報基盤研究部門 助教）

プログラミング教育における自己調整学習支援システムの研究開発

近年、大学教育においてプログラミング能力の重要性が高まる一方、学習者一人ひとりに対して適切なタイミングで個別支援を行うことは大きな課題となっている。プログラミング教育では、自動採点ツールや大規模言語モデル（LLM）を用いた学習支援が急速に普及しているが、単に正誤判定や解答例を提示するだけでは、学習者の問題解決力、メタ認知能力、自己調整学習能力の育成には十分でない。

本研究では、大学におけるプログラミング初学者を対象に、学習者がどのように問題を理解し、コードを作成し、エラーを修正し、自らの学習を振り返るのかに着目した。2025年度は、まずプログラミング学習における学習行動とメタ認知的困難を分析し、その知見をもとに、Moodle上の自動採点ツールCodeRunnerと連携するAI支援システム「CodeRunner Agent」の設計および開発を行った。これにより、学習分析、自己調整学習支援、AIフィードバックを統合したプログラミング学習環境の実現を目指した。

1. プログラミング初学者のメタ認知的困難に関する研究

本研究では、プログラミング初学者が課題解決の過程でどのような困難を抱えるのかを明らかにするため、プログラミング演習における学習行動と自己調整学習方略を分析した。実験では、学生がMoodle上の教材を参照しながら、自動採点ツールCodeRunnerを用いてプログラミング課題に取り組む学習環境を構築した（図1）。学習中のプログラム作成、解答確認、エラー修正などの操作ログを収集するとともに、学習後のインタビューを通じて、計画、自己モニタリング、振り返りに関する学習方略を調査した。



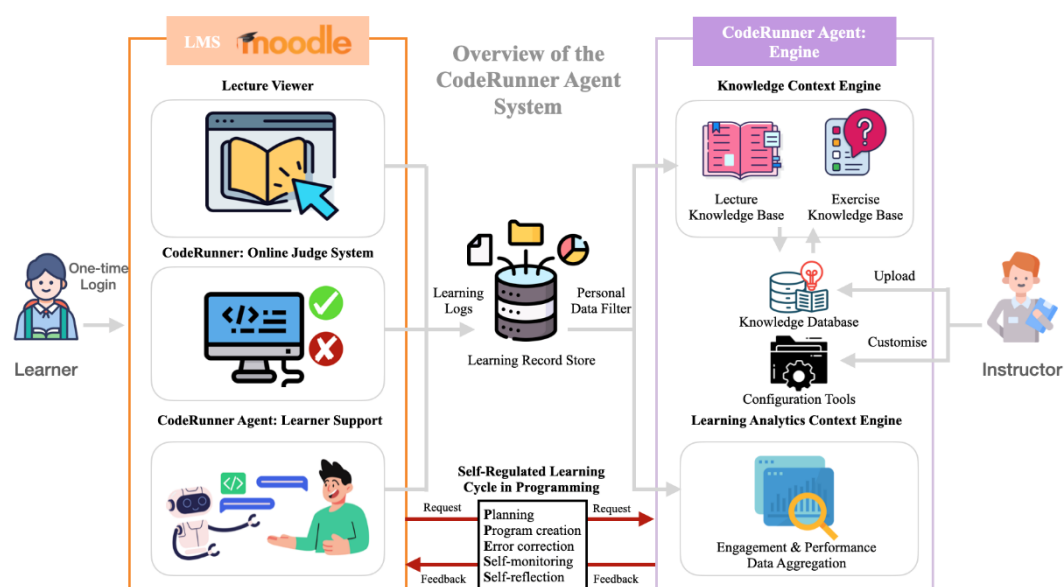
（図1）Moodle上の自動採点ツールCodeRunnerを用いたプログラミング課題例

分析の結果、初学者は経験者に比べて、問題理解、プログラムの論理構成、コード改善などの自己調整学習方略を十分に活用できていないことが明らかになった。特に、初学者は不正解のまま解答確認を繰り返す傾向があり、問題文の理解、エラー内容の解釈、自身の知識不足の把握、時間管理などに困難を抱えていた。一方、経験者はより多様な方略を用いて課題に取り組んでおり、問題の構造化やコードの見直しを通じて効率的に学習を進める傾向が見られた。

これらの結果は、プログラミング教育において、正誤判定を中心とした従来の自動採点支援だけでは不十分であり、学習者の計画、自己モニタリング、振り返りを促す支援が必要であることを示している。本研究成果は ICLEA 2025 で発表し、Best Short Paper Award を受賞した [1]。

2. 自動採点ツールと AI フィードバックを統合した自己調整学習支援システムに関する研究

上記の分析で明らかになったメタ認知的困難を支援するため、Moodle 上の自動採点ツール CodeRunner と連携する AI 支援システム「CodeRunner Agent」を開発した（図 2）。本システムは、学習者のコード、採点結果、講義資料、課題内容、学習ログを活用し、学習状況に応じた文脈依存型の AI フィードバックを提供することを特徴としている。



（図 2）AI 支援システム CodeRunner Agent の全体像

CodeRunner Agent では、自己調整学習理論に基づき、プログラミング学習過程を「計画」「プログラム作成」「エラー修正」「自己モニタリング」「振り返り」の五つの段階に整理した。各段階に応じて、問題理解、論理構成、エラー解釈、進捗確認、コード改善などを促すフィードバックを提示する設計とした。また、直接的な解答を提示するのではなく、学習者自身の思考を促すヒント型のフィードバックを重視し、AI への過度な依存を抑制する工夫を取り入れた。

学生および教員へのインタビューでは、CodeRunner Agent が問題理解、エラー修正、コード品質向上、自己省察を支援する可能性が確認された。特に、講義資料や課題情報と連携しながら Moodle 内で支援を提供できる点は、一般的な対話型 AI ツールにはない利点として評価された。また、教員がフィードバックの内容を授業目標や課題内容に合わせて調整できる点も、教育実践への応用可能性を高める要素である。

一方で、AI フィードバックの信頼性と透明性、利用頻度の制御、学習者の過度な依存の防止、教員によるフィードバック設計の支援、学習状況を可視化するダッシュボードの整備など、今後改善すべき課題も明らかになった。今後は、実際の授業環境において CodeRunner Agent の効果検証を進め、学習者の自己調整学習能力、プログラミング理解、学習継続性の向上に寄与する AI 支援システムの発展を目指す。

参考文献

- [1] Huiyong Li, Boxuan Ma, Chengjiu Yin. (2025). Examining Metacognitive Difficulties in Learning Programming: Analysis of Student Behavior and Strategy. Proceedings of the 1st International Conference on Learning Evidence and Analytics (ICLEA2025).
<https://library.apsce.net/index.php/ICLEA/article/view/5520> (Best Short Paper Award)
- [2] Huiyong Li, Boxuan Ma. (2025). CodeRunner Agent: Integrating AI Feedback and Self-Regulated Learning to Support Programming Education. Proceedings of the 33rd International Conference on Computers in Education (ICCE2025).
<https://library.apsce.net/index.php/ICCE/article/view/5975>

7.7 岡村 耕二（先端サイバーネットワーク研究部門 教授）

組織規模の脆弱性情報から脅威を把握する手法の実現

概要：ASM（Attack Surface Management）は、組織が外部へ公開する資産を継続的に発見し、脆弱性や設定不備等の指摘を収集して管理する取り組みである。資産増加により結果が膨大となり、一覧提示のみでは把握と優先度判断が難しい。本研究は XCockpit ASM の Risk Factor 各 1 行をイベントとし、Asset、severity、type、subject で集約して HeatMap と TreeMap で俯瞰可視化する。九州大学が管理するホストサーバ群が共通的に使用している IP アドレスの資産群を用いて比較し、type 構成と severity 分布の差を短時間で把握できた。外部観測ノイズや濃淡比較の制約に留意し、除外規則と subject 統合を今後の課題とする [1]。

1. はじめに

組織の成長に伴いインターネット上のサービスは多様化し、外部へ公開される資産（ドメイン等）も増加している。大学のように部局・プロジェクト単位で個別運用される環境では、管理主体が分散し、把握漏れや意図しない公開により攻撃対象領域が拡大しやすい。

従来は台帳管理や監査、内部ログ監視により資産とリスクを把握してきたが、クラウド利用やサービス分散により公開範囲が継続的に変化し、内部情報のみで外部露出資産を一貫して把握することは難しい。結果としてリスク発見が遅れ、重大なインシデントに発展する可能性がある。

この課題に対し、外部視点で公開資産を継続的に発見し、脆弱性や設定不備等の指摘を収集・管理する ASM が注目されている。本研究では CyCraft 社 XCockpit ASM を対象とし、得られる検出結果（イベント）を分析に用いる。一方で ASM の結果は膨大でノイズも含み得るため、一覧提示のみでは全体把握や優先度判断が難しい。

そこで本研究は、XCockpit ASM の検出結果を整理・要約し、資産群と severity、type の分布と偏りを把握できる HeatMap / TreeMap を設計・実装して、運用上の意思決定を支援することを目的とする。

2. ASM イベント分布に基づく資産群比較の可視化手法

2.1 手法の全体像

本章では、XCockpit ASM が出力する検出結果を、組織規模の観点から俯瞰可能にし、優先度判断につながる形で整理・要約・可視化する手法を提案する。本研究は ASM プロセスのうち、検出結果を運用判断へ接続するための集約と提示に焦点を当てる。まず、XCockpit ASM の検出結果を入力とし、Asset 表記のゆらぎを抑える前処理として空白整形を行い、URL スキーム、パス、クエリ、ポート番号等を除去してドメイン形式へ正規化する。ついで、Asset、severity、type、subject を軸にイベント件数を集約し、分析に適した形式へ変換する。最後に、可視化として HeatMap により資産群ごとの severity 分布と偏りを俯瞰し、TreeMap により subject 単位の構成比と重点領域を把握する。

2.2 対象システムと分析単位

2.2.1 XCockpit ASM における資産と指摘

本研究が対象とする XCockpit ASM では、外部観測に基づく資産を Domain、Service、Account、IP Address 等のタイプとして扱い、資産に紐づく指摘を Risk Factor として整理する。Risk Factor は、一覧画面において Severity（Color & Number）、Time、Subject、Category 等の項目として提示される。

2.2.2 event の定義

本研究では、外部公開資産（Asset）は主にドメインを指す。XCockpit ASM の Risk Factor 一覧の 1 行を分析単位（event）とし、各 event は Severity（severity）、Event Type（type）、Subject（subject）等の属性を持つものとして集約・可視化を行う。XCockpit ASM が提示する Exposure Profile の 1 行を、分析上の最小単位であるイベント（event）として扱う。各 event は、観測対象資産（Asset）に紐づき、分類軸として type、subject、severity 等の属性をもつものとする。以降の集約と可視化は、event の属性に基づき行う。

2.3 分類軸の定義

2.3.1 type の定義と一覧

本研究では、XCockpit ASM User Guide に示される分類体系に基づいて ASM 出力を整理する。type は、ASM が提示する Risk Factor の大分類であり、本研究では次の 11 種類を用いる。本研究で用いるデータでは、type は Event Type 列に対応する。

- Weakness
 - Certificate Weakness
 - Cipher Weakness
 - Email Weakness
 - DNS Weakness
 - Website Weakness
 - Network Weakness
 - Software Weakness
- Misconfiguration
 - Website Security
 - Network Security
 - Certificate Security
- Vulnerability
 - Exploited Vulnerability

2.3.2 subject の定義と一覧の扱い

本研究では、subject を type に属する下位項目として定義する。User Guide では、各 type に対して Description が与えられており、その中で列挙される具体的な指摘内容（例：missing SPF record 等）を、原則としてカンマ区切りで分割して得られる要素として subject を定める。本研究で用いるデータでは、subject は Subject 列に対応する。なお、Description 中の列挙は文章表現であり、分割規則によって粒度が変化し得るため、本研究では表記ゆれの正規化を併用し、subject を分析用ラベルとして整備する。

2.3.3 severity の取り扱い

XCockpit ASM では、Risk Factor の severity が数値と色を伴って提示され、通知条件としても severity の閾値（例：6～10 等）を設定可能である。severity は 1～10 で表され、1～6 を Low、7～8 を Medium、9～10 を High として区分される。本研究では、event に付与された severity を危険度指標として用い、type および subject との組合せで分布と偏りを可視化する。

2.4 資産群 (Group) の定義

本研究では、組織内の資産を2群に分けて比較する。Group Aは特定のIPアドレスに割り当てられているドメイン群、Group Bはそれ以外の資産群である。この2群に対して、type、subject、severityの分布を同一の手順で可視化し、群間差を俯瞰する。

2.5 可視化設計

2.5.1 HeatMapによる分布と偏りの俯瞰

HeatMapは、行にAsset (ドメイン)、列にseverityを配置し、各セルに当該ドメインにおける当該severityのイベント件数(counts)を割り当てた行列として構成する。本研究では、type (例: Certificate Weakness、Email Weakness等)ごとに色相を割り当て、さらに同一type内で件数に応じて濃淡を変化させる。これにより、特定typeのイベントが「どのseverity帯に集中し、どのドメイン群に偏在しているか」を俯瞰できるようにする。また、Group A/Bで同様の可視化を作成することで、群間の分布差を比較できる。

2.5.2 TreeMapによる重点領域の俯瞰

TreeMapは、subjectを矩形として配置し、矩形面積にイベント件数を割り当てることで、全体に占める構成比と重点領域を同時に俯瞰するために用いる。本研究では、typeを上位階層、subjectを下位階層として階層構造を構成し、矩形面積に件数、色に危険度指標としてseverityに基づく値を割り当てる。これにより、「件数が多く、かつ危険度が高い領域」を優先的に発見できる表現とする。

3. 資産群 (Group A/B) に対するイベント分布比較

3.1 比較の目的

本研究の目的は、ASMが出力する膨大な検出結果を、組織全体を俯瞰できる形へ要約し、運用上の意思決定 (重点領域の把握、優先度判断) に接続しやすくすることである。本章では、提案するHeatMapおよびTreeMapが、(1) 分布と偏りの把握、(2) 重点領域の特定、(3) 資産群間の差分理解、に有効であるかを確認する。

また評価として、IPアドレスに特定のホストIPアドレスを含むドメイン群 (Group A) と、含まないドメイン群 (Group B) を比較し、可視化表現が群間差の説明に寄与するかを検討する。

3.2 実験データ

実験データは、XCockpit ASMから取得したRisk Factor一覧である。本研究では、Risk Factor一覧の1行を1件のeventとして扱い、Asset、Severity、Event Type、Subjectを用いて解析する。

3.3 比較条件

3.3.1 可視化対象と集約単位

可視化の集約単位と割当 (HeatMap: Asset $\times$ severity、TreeMap: type $\rightarrow$ subject 階層、面積 = 件数、色 = severity 代表値) は、先に定義した仕様に従う。本章では、Group A/Bに対して同一仕様の可視化を生成し、分布形状と重点領域の差分に着目して比較する。

3.3.2 2群比較の定義

比較評価では、特定のIPアドレスを含むドメイン群をGroup A、含まないドメイン群をGroup Bとして2群に分割し、両群に対して同一仕様の可視化を生成する。IPアドレスはDNS解決等によりドメインへ付与したIPv4アドレスの集合であり、複数IPを持つ場合は集合に含まれるか否かで判定する。

3.4 実験手順

- (1) XCockpit ASM からイベント一覧を取得する。
- (2) 入力データを前処理し、必要な列（Asset、Severity、Event Type、Subject 等）を抽出する。
- (3) 集約処理を行い、HeatMap 用および TreeMap 用の可視化データを生成する。
- (4) 可視化データを JSON へ変換し、HTML から読み込んで HeatMap と TreeMap を描画する。
- (5) IP アドレスの条件により 2 群へ分類し、両群の可視化結果を比較する。

3.5 評価観点

提案表現の有効性を確認するため、以下の観点で評価する。

- ・ 分布把握：高 severity（例：7 以上）のイベントが集中するドメイン領域を短時間で特定できるか。
- ・ 重点領域把握：TreeMap により、イベント数が多い subject、または高 severity が多い subject を俯瞰的に把握できるか。
- ・ 比較容易性：Group A と Group B において、分布の差分や重点領域の差分を直感的に説明できるか。

補助的な定量情報として、総イベント数、severity 別件数、高 severity（例：7 以上）のイベント件数、および該当ドメイン数を算出し、可視化結果の解釈を支援する。

3.6 データ概要

本研究では、ASM 上のイベント数を集計の基礎として扱う。Group A の ASM 上のイベント数は 812、Group B は 2831 である。また、本章で示す counts は観測されたイベントの頻度（繰り返しを含む）を反映する集計値であり、ユニークな弱点数を直接表すものではない点に留意する。

3.7 可視化結果

3.7.1 HeatMap (Group A)

図 1 に Group A の HeatMap を示す。行は Asset（ドメイン）、列は severity であり、各セルは当該ドメインにおける当該 severity の観測件数（counts）を表す。本可視化では type ごとに色相を割り当て、同一 type 内で counts に応じて濃淡を変化させた。このため、severity の高低だけでなく、「どの type が、どの severity 帯に、どのドメイン群として現れているか」を俯瞰できる。

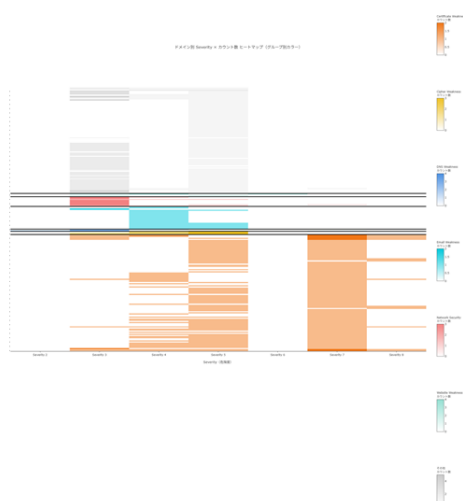
Group A では、Certificate Weakness に対応する領域が相対的に多く観測され、とくに高 severity 側（例：severity=7 付近）にまとまって出現する傾向が見られる。

なお、本研究では証明書関連 subject の意味づけや原因の深掘りは行わず、定義・判定基準は XCockpit ASM UserGuide に従うものとする。

3.7.2 HeatMap (Group B)

図 2 に Group B の HeatMap を示す。Group B では、中 severity 帯（例：severity 4-6 付近）を中心とした分布が目立ち、Group A で観測された高 severity 側のまとまりとは異なる傾向が確認できる。とくに Website Weakness や Cipher Weakness 等において、中 severity 帯の指摘が複数ドメインへ広く分布している様相が観察される。このことは、Group B では「一部ドメインに高 severity が集中する」というよりも、「複数ドメインに中 severity の設定不備が分散している」状態を示唆する。

また Group B 側には IP アドレス形式の Asset も含まれるが、本研究ではこれらを除外せずに扱う方針とする。ただし、ドメイン形式の Asset と IP アドレス形式の Asset は同一視しないよう注意を要する。



(図1) HeatMap(Group A)



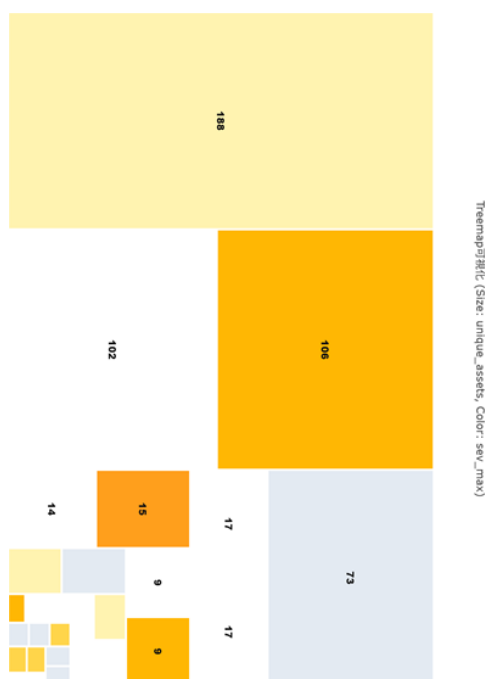
(図2) HeatMap(Group B)

3.7.3 TreeMap (Group A)

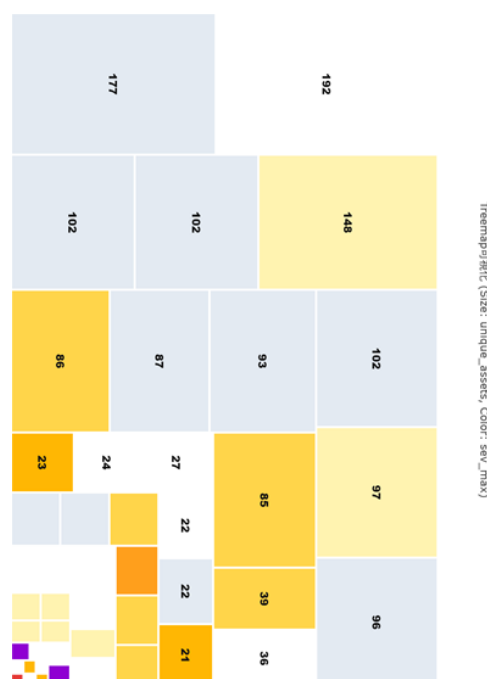
図3にGroup AのTreeMapを示す。本図は、矩形面積に unique assets（当該 subject が観測された資産数）、色に危険度を表す指標（実装では sev max）を割り当てている。Group Aでは、大きな矩形が少数存在し、それらが全体の面積（出現範囲）を占める一方で、その他は小規模な矩形が点在する構成となっている。このことは、特定の subject が相対的に多くの資産で観測され、重点領域が比較的限制される可能性を示唆する。

3.7.4 TreeMap (Group B)

図4にGroup BのTreeMapを示す。Group Bでは、中規模以上の矩形が複数存在し、全体として矩形が広く分散して配置されている。これは、複数の subject が一定規模の資産範囲で観測され、重点領域が分散する傾向を示す。また、色の分布から、危険度指標が高い領域（暖色系）の subject が複数存在することが視覚的に確認できる。



(図3) TreeMap(Group A)



(図4) TraaMap(Group B)

3.8 可視化の結果

本章では、IP アドレス条件により資産群を2群に分割した。HeatMap では、資産をドメイン単位で行に配置し、列に severity をとってイベント件数を可視化することで、各ドメインがどの severity 帯でどの程度指摘を受けているかを把握できる。さらに、type ごとに色相を割り当てているため、危険度分布に加えて、全体としてどの type の指摘が多いかを俯瞰できる。一方、TreeMap では subject 単位で件数を集計し、矩形の面積で件数規模を、色の段階で severity 代表値を表すことで、どの subject が多く観測され、それらがどの程度の危険度に位置づくかを同時に把握できる。

その結果、2群で type 構成および severity 分布が異なることを、可視化により短時間で把握できることを確認した。次章では、得られた差分の解釈と運用上の活用可能性、ならびに本手法の適用上の制約について考察する。

4. 考察

4.1 提案表現の有効性

TreeMap は subject ごとの規模（件数）と危険度 severity）を同時に示すため、運用者が重点領域（件数が多い／危険度が高い）を短時間で抽出できる。HeatMap は Asset × severity 分布を可視化し、指摘が偏在する資産や深刻度帯を把握しやすい。両者は「重点領域抽出（TreeMap）」と「分布把握（HeatMap）」を分担し、ASM 出力の要約表現として相補的に機能する。

2群を比較すると、type 構成と severity 分布の差が可視化から確認できた（例：Group A は CertificateWeakness 等が目立ち、severity=7 帯が多い一方、Group B は WebsiteWeakness が支配的で severity=4 帯に集中する）。また TreeMap 上位の subject から、Group A では証明書設定や外部公開構成に関する指摘、Group B では Web 応答ヘッダやソフトウェア起因の指摘が重点領域として抽出された。

4.2 今後の課題

今後は、(1) ノイズ低減のための除外規則の精緻化、(2) subject 粒度を安定させる辞書整備・同義語統合、(3) type 横断比較を補助する定量集計（例：type 別件数や高 severity 件数）の併記、を検討する。加えて、外部観測に起因する誤検出の影響を把握するため、(4) 他の ASM 製品や資産台帳・構成管理データベース、あるいは手動検証結果との突合により誤検出率や適合率等の 精度指標に基づいて ASM 出力の妥当性を評価することも課題である。

5. 結論

本研究では、ASM が出力する検出結果を運用判断へ接続しやすい形で俯瞰可能にすることを目的として、可視化用の集約処理を行い、HeatMap および TreeMap による要約表現を提案した。HeatMap により Asset × severity の分布を可視化し、重点的に確認すべき資産領域の抽出を摘内容の抽出に寄与し得ることを示した。

さらに、IP アドレス条件で分割した2群比較により、type 構成および severity 分布が異なることが可視化から確認できた。このことは、提案表現が実運用での比較・把握（状況認識と初動トリージ）に有用である可能性を支持する。

一方で、外部観測に基づく ASM 出力にはノイズや誤検出が含まれ得るため、本研究の結果は検出真偽の確定を意図するものではなく、分布の偏りや重点領域の抽出に主眼を置く。今後は、重複を含む件数指標の併記、詳細表示への導線整備、および subject 統合の高度化に加え、他のシステムや資産台帳等との突合による誤検出評価を行うことで、意思決定支援効果をさらに高めることが課題である。

- [1] 安藤謙, “組織規模の脆弱性情報からの脅威の把握に関する研究”, 九州大学 工学部電気情報工学科 卒業論文 2026 年 2 月

7.8 笠原 義晃（先端サイバーネットワーク研究部門 助教）

九州大学 Microsoft 365 環境への多要素認証導入に関する中間報告

1. はじめに

九州大学では、本学構成員へ Microsoft Office アプリケーションを提供するため、2007 年から包括契約を利用している。2014 年の契約形態変更により Office 365 ProPlus が全構成員で利用可能となったことから、2015 年中に内部的な検証を実施し、2016 年に Office 365 (A1 ライセンス +ProPlus) の学内提供を開始した [1]。その後、2018 年に認証システムとアカウント管理を再構成することになり、Office 365 環境を新規テナントで再構築した [2]。2020 年度には再度契約形態が変更になり、Microsoft 365 A3 でのサービス提供となった。2026 年現在もその構成で運用している。

本学では、2016 年の導入直後から、Microsoft 365 (旧 Office 365) での多要素認証 (MFA: Multi-Factor Authentication) による利用者アカウントの保護について調査検討し、試行錯誤を続けてきた [3]。本稿では、その過程で得られた知見と、現在進行中の多要素認証必須化について報告する。

なお、本稿では当時の名称である Azure Active Directory ではなく Entra ID で記載している (モジュール名など一部例外あり)。

2. 「ユーザごとの MFA」機能のセルフサービス有効化

我々は、2016 年の Office 365 導入時点から、本サービスには MFA 機能が内蔵されている事を把握していた。当時の環境で利用可能だった MFA 機能は、Microsoft 365 の認証基盤である Entra ID の機能で、「ユーザごとの MFA(Per-User MFA)」と呼ばれている。この機能は教育機関向け無償ライセンスである A1 ライセンスから利用できる。MFA を利用したい場合、利用者からは有効化の操作はできず、管理者が各アカウントについて個別に有効化操作をする必要がある。

Office 365 の導入当初から利用者アカウント保護のため MFA を導入したいという意向はあったが、本学の利用者規模 (約 3 万アカウント) において、365 の管理者が全学で一括、もしくは利用者の希望に応じて個別に MFA を有効化するという運用は現実的ではなかった。そのため、本学で別途運用している全学認証基盤の利用者設定用ページに 365 の MFA 設定を切り替えるスイッチを設置し、MFA 利用を希望する利用者がセルフサービスで設定を変更するという運用を念頭に認証基盤への追加実装を実施した。具体的には、利用者による設定の on/off に応じて、Azure AD PowerShell モジュールを介して 365 側の MFA 設定を enabled/disabled に変更する構成とした。

当初この機能は動作しているように見えたが、その後 365 側の設定値は enabled/disabled の二値ではなく、enforced/enabled/disabled の三値であり、enabled/disabled の切り替えでは正常に動作しない可能性がある事が判明した。具体的には、enabled にすると次のサインインで MFA のための本人確認情報の登録を求められ、登録が完了すると enforced になり、以降のサインインで毎回 MFA による本人確認が求められるという仕組みになっていたが、一度 enforced になったアカウントに対し enabled に上書設定すると enforced には戻らないという 365 側の仕様上の問題があった。当時の事情により、構築済のシステムに対しさらに再設計して修正する事ができなかったため、この機能は公開を中止することになった。

現在「ユーザごとの MFA」はレガシー機能として扱われており、「条件付きアクセス」が利用できる環境 (有償の Entra ID P1 ライセンス以上) での利用は推奨されていない。しかし、Entra ID の無料ライセンスしか利用できない環境ではこの機能を利用したい場合もあると思われる。2026 年現在、

Microsoft Graph API により Per-User MFA の設定に加えて各アカウントごとの MFA 情報の登録状態も取得可能なため、enforced/enabled/disable の設定値を適切に制御可能と思われる。

3. 「条件付きアクセス」を利用したセルフサービス MFA 有効化と設定変更

2020 年から包括契約の内容が変更になり、本学では Microsoft 365 A3 ライセンスを通して Microsoft 365 サービスと Microsoft Office アプリケーションを提供することになった。A3 ライセンスには Entra ID P1 ライセンスが含まれていることから、「条件付きアクセス」機能を利用可能であることが判明した。

「条件付きアクセス」は Entra ID のプレミアム機能で、管理者がポリシーを設定することで、365 へのサインイン時にユーザやグループ・利用したいサービス・接続元 IP アドレス等の様々な条件をトリガーとして、認証をブロックしたり、MFA を追加で要求したりすることができる。また、Entra ID P2 ライセンスでは「リスクベース認証」が利用できる。通常と異なる環境やネットワークからのサインインなど、不正利用のリスクを Entra ID が自動で評価し、リスクレベルを決定する。このリスクレベルを条件とするポリシーが設定可能となるため、不審なサインインの時のみ MFA による本人確認を要求する、といった処理ができる。本学では、A3 への切り替え前に Entra ID P2 を含む A5 ライセンスの評価を実施し、リスクベース認証は必須との判断から A3 ライセンスに加えて Entra ID P2 を別途契約している。

利用者の利便性を考えると、可能であれば「MFA 情報が登録されている利用者がサインインしてきたら MFA を要求する」(=MFA 情報がない利用者には MFA を要求しない) というポリシーを実装したかったが、条件付きアクセスではサインインしてきたアカウントに MFA 情報が登録されているかどうかを条件にできない。このため、「MFA 有効化済」を意味するセキュリティグループに所属する利用者のみ MFA を要求する、というポリシーを用意した。条件付きアクセスポリシーによりサインイン時に MFA を要求する場合、MFA 情報が未登録だと MFA 情報を登録するフローが強制的に開始され、登録が完了するまでサインインできないため、利用者への侵襲は大き目である。

この「MFA 有効化済」グループは何らかの手段で更新する必要がある。当初案は、全アカウントについて MFA 情報が登録されているかどうかを Power Automate の定期実行でスキャンして、利用者が MFA 情報を登録すると自動で MFA が有効化される、というものであった。実際に試作してみると当時の API には機能が不足しており、安定的に全アカウントの MFA 情報登録状況を取得できなかった。このため、利用者が手動で有効化する構成とした。別途サーバなどのインフラを準備しなかったため、MFA の有効化は Microsoft Forms をトリガーとして Power Automate でグループへの追加処理をする構成とした。

利用者本人が MFA 情報を登録する前に第三者によりアカウントに不正アクセスされ、不正な MFA 情報を登録されてしまうことの懸念から、MFA の有効化操作をするまでは MFA 情報を登録するページ(マイアカウント内のセキュリティ情報)自体を条件付きアクセスでブロックする構成とした。不正アクセスに対する安全性は高くなる一方で、利用者が何らかの理由でセキュリティ情報ページにアクセスした際にブロックメッセージが出てしまい、このメッセージは本学側で変更できないため、利用者から問い合わせが発生するという難点があった。

また、MFA 有効化直後は接続元ネットワークを問わず常時 MFA が必要なポリシーとし、かつセッションの保持期間を 7 日とすることで、MFA による本人確認操作を頻繁に実施させ、強制的に慣れてもらう構成とした。そのままのポリシーで長期利用し続けることは利用者への負担が大きいため、別の Forms で MFA の設定を変更可能とした。Forms での選択により、いくつかのセキュリティグループの登録状況を変更することで「常時 MFA」「学外からのみ MFA」「リスク検知時のみ MFA」の選択と、セッション維

持期間「7日」「90日（365の既定値）」から選択可能とした。これらの機能は2021年12月から任意登録期間として提供を開始した。

4. MFA 必須化に向けたポリシー変更

他大学での MFA 導入状況等を鑑み、本学でも 2024 年ごろから全学必須化について検討を進め、2026 年 9 月からの MFA 全アカウント必須化を目標として進めることとなった。

必須化に際して、MFA 登録に関するポリシーを見直し、利用者への侵襲度を下げて心理的障壁を減らす事になった。具体的には、(1) リスクベース認証を活用し、利用者には MFA をなるべく要求せずに、不正アクセスのみをブロックする、(2) MFA の有効化操作を、MFA 登録の後にする、の二点を実施することとなった。

(1) については、MFA 有効化後の標準設定を「リスク検知時のみ MFA」に緩和することにした。Entra ID のリスク検知が期待通り動作する事が前提となるが、事務系を中心にほとんどの利用者は職場や自宅の同じ環境から Microsoft 365 を利用することから、無用な MFA 要求はしない一方で、リスク検知により不正アクセスを未然に防ぐ構成にすることで、利用者への侵襲度を下げた。同時に、セッション長 7 日の設定は廃止することとした。なお、常時 MFA、及び学外からのみ MFA の設定変更フォームは残し、利用者が希望すればより安全なポリシーへ設定可能とした。

(2) については、MFA 情報登録ページの常時ブロックポリシーを、不正アクセスの可能性があるリスク検知時のみブロックするポリシーに変更した。これにより、利用者は任意のタイミングで MFA 情報を登録し、それがうまく完了した所で MFA を有効化する、という順序で作業できる。当初のポリシーでは、MFA を有効化すると次のサインインで MFA を登録完了しなければサインインできなくなるため、操作マニュアルの警告も強い内容となっており、利用者の忌避感が高くなっていたと思われる。

一方で、利用者が MFA 用に登録したスマートフォン等を紛失した際に管理者に依頼せずに回復できるようにするため、利用者マニュアルで MFA 情報を最低二種類登録することを強調した。

2026 年 1 月から MFA 必須化への移行期間とし、利用者には MFA 設定を依頼する通知を実施した。また、本学では 4 月からの新入生に対して 3 月中に入学前学習として必携 PC の設定等を実施させているところ、2026 年度の新入生から準備内容に MFA 設定を追加した。

これらの施策により、2026 年 4 月時程で全アカウントの 25% が MFA 設定を完了した状況である。強制必須化の期限である 9 月以降は、全アカウントについて MFA を有効化した状態にし、MFA 登録が未完了の利用者はサインイン時に MFA 情報の登録を強制する予定である。管理者への問い合わせを減らすためには、強制必須化までにできるだけ登録率を 100% に近づけておく必要があり、登録率が上がらなければ強制必須化を延期せざるを得ない可能性がある。引き続き通知等により設定の必要性について注意喚起していく予定である。

5. 再度のポリシー変更（予定）

1 月からの移行期間に利用者対応をしていた担当者から「問い合わせの多くが電話対応となっており、かつ一件あたり 30 分以上かかるケースがあるため、サポート担当の負担が大きい」という報告があった。スマートフォン等での MFA 操作に慣れている学生などと、そうでない利用者との間に操作内容の理解度にかなり差がある。さらに、スマートフォンを常用していない利用者は PC への認証アプリの導入が必要となる場合がある。しかし、PC 用の認証アプリ自体あまり一般的でなく、また MFA 登録時に画面上の QR コードの取り込み操作や認証コードの直接入力など、利用者にとってなじみのない特殊な操作が必要で混乱することが多いため、電話での説明も困難な状況になりやすい。このような状況で、MFA 情報を

一つ登録するのにも困難な利用者について、二つ登録することを強制するのはさらに難しい。4月現在での登録率の低さを考えると、デバイス紛失等での管理者対応の負荷を下げるより、利用者によるMFA情報登録を容易にする方が優先度が高いと考えられることから、登録は最低一つでよい、と方針変更することになった。

また、MFA情報の登録状況について精査した所、MFA情報登録を有効化より先にした結果、MFA情報を登録しただけで有効化操作を忘れている利用者がかなりいることがわかった。現状の構成では、MFAの有効化操作を忘れると条件付きアクセスで保護されず、MFA登録した意味がない。そこでEntra IDのAPIについて再度調査したところ、Identity and access reports APIで全アカウントのMFA登録状態を一括で取得できる「List userRegistrationDetails」が実装されていることが判明した。このAPIを利用することで、安定的にMFAの登録状況を収集し、MFA有効化済のセキュリティグループに自動的に追加する処理が実現可能となった。このため、利用者による手動でのMFA有効化操作も廃止することにした。これにより、利用者はMFA情報を1つ登録するだけで、自動的にリスクベース認証での保護下に入るようになる。

2026年に入って、Entra IDでパスキーが正式対応となり、MFAとしても利用可能となった。パスキーはスマートフォンのみならずWindowsやmacOSでもネイティブ対応が進んでいることから、動作検証により大きな問題がなければPCでのMFAの選択肢として案内する予定である。これにより、懸案であったPC環境のみでのMFA対応もユーザ体験が改善する事を期待している。

参考文献

- [1] Kasahara, Y., Shimayoshi, T., Obana, M. and Fujimura, N.: Our Experience with Introducing Microsoft Office 365 in Kyushu University, in Proceedings of the 2017 ACM Annual Conference on SIGUCCS, SIGUCCS '17, pp. 109–112, New York, NY, USA (2017), ACM.
- [2] Shimayoshi, T., Kasahara, Y. and Fujimura, N.: Renovation of the Office 365 Environment in Kyushu University: Integration of Account Management and Authentication, in Proceedings of the 2019 ACM SIGUCCS Annual Conference, SIGUCCS '19, pp. 135–139, New York, NY, USA (2019), Association for Computing Machinery.
- [3] Kasahara, Y. and Shimayoshi, T.: Our Design and Implementation of Multi-Factor Authentication Deployment for Microsoft 365 in Kyushu University, in Proceedings of the 2022 ACM SIGUCCS Annual Conference, SIGUCCS '22, pp. 56–61, New York, NY, USA (2022), Association for Computing Machinery.

7.9 美添 一樹 (先端計算科学研究部門 教授)

柔軟な探索空間制御による汎用分子設計フレームワーク

我々が先行研究で開発してきた探索と機械学習に基づく分子設計フレームワーク ChemTS / ChemTSv2 を発展させ、利用者が柔軟に探索空間を定義・切り替えられる新たなフレームワーク ChemTSv3 を開発した [1]。従来の多くの分子生成手法では、SMILES 文字列、分子グラフなどの分子表現や、特定の生成モデル等によって探索対象となる候補分子の空間、すなわち化学空間が暗黙に固定されていた。そのため、実際の創薬・材料探索のように、広い探索から局所的な最適化へ移行したい場合や、途中で評価関数・生成規則・分子表現を変更したい場合には柔軟性に限界があった。本研究では利用者がそれぞれの目的に応じて、探索空間そのものを柔軟に設計することを可能とした点に貢献がある。

ChemTSv3 の主な貢献は、節点の定義 (Node)、枝の定義 (Transition)、報酬 (Reward)、除外設定 (Filter) のモジュールを交換することで、利用者が容易にモンテカルロ木探索 (MCTS) に基づく木探索を制御可能としたことにある。Node は探索木上の状態を表し、SMILES や SELFIES などの文字列表現、分子グラフ、FASTA で表されるタンパク質配列などを同一の枠組みで扱うことができる。Transition はある Node から次の Node への遷移を表す。Reward は目的関数であり、機械学習予測器、シミュレーション、多目的最適化などを組み込むことが可能である。Filter は望ましくない構造の除外や、計算コストの高い評価を避けるための枝刈りに用いられる。これにより ChemTSv3 は、ChemTSv2 が持っていた柔軟な報酬関数・フィルタ設計を継承しつつ、探索空間の定義をユーザが自由に変更可能とした。

特に重要な拡張は、探索途中で探索空間を動的に切り替えられる点である。たとえば、初期段階では SMILES/SELFIES に基づく文字列上の自己回帰モデルで広い化学空間を探索し、有望なヒット化合物が得られた後、その化合物を起点として局所的な探索を行うことができる。これは、創薬で典型的に行われる広域探索からヒット化合物周辺の局所最適化へ進む流れを、探索空間の切り替えとして実現するものと言える。

上記のアプローチの有効性を検証するために、医薬候補分子の設計、タンパク質配列設計、および標準ベンチマーク問題を用いた評価を行った。医薬候補分子の設計では、まず生成モデルを用いて広い候補空間を探索し、得られた有望候補を起点として局所的な探索に切り替えることで、評価値の改善が確認された。また、予測器の訓練データから除外されていた既存の EGFR 阻害剤 Erlotinib および Gefitinib を再発見しており、現実的な探索にも有用であることが示唆された。さらに、GFP のタンパク質配列設計では予測蛍光強度の向上を確認し、PMO benchmark では探索空間を切り替える設定が切り替えを行わない設定を上回る性能を示した。ChemTSv3 の実装は GitHub で公開されている [2]

関連する総説論文 [3] では、低分子設計における生成 AI の発展を、化学における逆問題の解決を支援・加速する技術として整理し、2016 年頃以降に急速に発展した深層学習ベースの分子生成手法を概観している。対象となる手法は、変分オートエンコーダなどの初期の深層生成モデルから、近年注目される大規模言語モデルや拡散モデルまで幅広く、低分子化合物の設計において AI がどのように化学空間探索を支援し得るかを体系的に整理している。

1. S. Fujii, Y. Murakami, T. Yoshizawa, S. Ishida, N. Cho, M. Ohta, T. Honma, K. Yoshizoe, M. Sumita, K. Tsuda, K. Terayama, "ChemTSv3: Generalizing Molecular Design via Flexible Search Space Control.", ChemRxiv preprint (2025).
DOI: 10.26434/chemrxiv-2025-kdvrt
2. ChemTSv3: A unified tree search framework for molecular generation
<https://github.com/molecule-generator-collection/ChemTSv3>
3. M. Sumita, S. Ishida, K. Yoshizoe, R. Tamura, K. Terayama, K. Tsuda, "Molecular Design with Artificial Intelligence: Progress and Perspectives for Small Molecules.", Chemical Reviews, 126, 3007–3054 (2026).
DOI: 10.1021/acs.chemrev.5c00689

7.10 渡部 善隆（先端計算科学研究部門 准教授）

2025年度は、昨年度からの継続研究として、非線形関数方程式に対する精度保証付き数値計算手法の一般化・汎用化に取り組んだ。得られたおもな研究成果は以下の通りである。

- (1) 非線形関数方程式に対し、適切な離散化により得られた近似解の周りで非線形関数方程式を線形化した作用素を与える。この線形化作用素は無限次元 Newton 型不動点問題に基づく非線形問題の解の検証において重要な役割を果たす。2025年度は、Hilbert 空間における線形化作用素の計算機による可逆性検証条件の改善、および逆作用素ノルムの上界・下界を効率的に評価する手法の構築を推進した。具体的には、2 階楕円型作用素における収束オーダーの最適評価および定量的な逆作用素ノルムの誤差限界を与えることに成功し、国際学術誌にて成果を公開した。
- (2) 成果 1 で得られた線形化作用素の可逆性検証と逆作用素ノルム評価を用いて、非線形関数方程式を無限次元 Newton 型作用素による不動点形式に同値変形させ、適切な Hilbert 空間における弱形式を満たす解の存在検証に適用する研究を推進した。不動点問題を計算機により解くためには、計算機による取扱いが可能な有限次元部分空間の設定と無限次元空間から有限次元部分空間への直交射影の定量的誤差評価が必要となる。2025年度は、様々な形状の領域および重調和問題にも対応する区分的 5 次 Hermite 関数、および、Legendre 基底関数基底関数より構成される有限次元部分空間に対する直交射影の効率的誤差評価手法を推進した。
- (3) 成果 1 および成果 2 により得られた理論および基盤技術を、非線形関数方程式の解の存在検証に適用した。2025年度は、Navier-Stokes 方程式から導かれる Elkouh 問題および Bermann 問題の解に対する精度保証付き数値計算に取り組み、いくつかの非自明解の検証に成功するとともに研究成果を公開した。
- (4) 3 次元波動方程式から導かれる特異項を含む無限次元線形微分作用素の逆作用素ノルムの効率的な上界評価に成功し、国際学術誌にて成果を公開した。また、得られた知見を数学的な未解決問題のひとつである波動方程式の爆発解の安定性解析に適用する国際共同研究を推進した。

研究成果は以下の学術論文・国際会議会議録として出版された。

- ▶ Yoshitaka Watanabe and Shuting Cai: A computer-assisted proof for solutions of Elkouh's equation related to Navier-Stokes equations, *Fixed Point Theory and Algorithms for Sciences and Engineering*, vol. 7, (May 2025). <https://doi.org/10.1186/s13663-025-00790-9>
- ▶ Shuting Cai and Yoshitaka Watanabe: Stability analysis of Kolmogorov flow by an eigenvalue excluding method, *Japan Journal of Industrial and Applied Mathematics*, Volume 42, Issue 5 (December 2025), pages 2107-2132. <https://doi.org/10.1007/s13160-025-00709-2>; <https://hdl.handle.net/2324/7402145> (embargoed access)
- ▶ Kenta Kobayashi and Yoshitaka Watanabe: Efficient numerical verification procedure of norm bound for infinite-dimensional differential operator with singular term, *Japan Journal of Industrial and Applied Mathematics*, Volume 43, article number 19, (February 2026). <https://doi.org/10.1007/s13160-026-00776-z>

- ▶ Takehiko Kinoshita, Yoshitaka Watanabe, and Mitsuhiro T. Nakao: Quantitative lower bounds estimates for the norm of resolvent using Galerkin approximation and its applications, Japan Journal of Industrial and Applied Mathematics, Volume 43, article number 28, (March 2026). <https://doi.org/10.1007/s13160-026-00787-w>

計算機援用証明の遂行においては、九州大学情報基盤研究開発センターのスーパーコンピュータシステム玄界 (Genkai) を活用した。

共同研究

2025 年度に行なった共同研究は以下の通りである。(敬称略)

1. 無限次元線形化作用素の可逆性判定とノルム評価の効率化
共同研究者: 中尾 充宏 (早稲田大学) 木下 武彦 (佐賀大学)
2. 非線形波動方程式に対する計算機援用証明
共同研究者: Michael Plum, 長藤 かおり (ドイツ・Karlsruhe 工科大学)
3. Navier-Stokes 方程式に対する計算機援用証明と安定性解析
共同研究者: Shuting Cai (中国・Fujian Jiangxia 大学)
4. Navier-Stokes 方程式の解の検証および解の形状に関する精度保証付き数値計算
共同研究者: 宮路 智行 (京都大学)
5. 特異項を含む線形作用素の逆作用素ノルム評価の高性能化
共同研究者: 小林 健太 (一橋大学)

7.11 南里 豪志（先端計算科学研究部門 准教授）

非ブロッキング集団通信の実用上の効果に関する研究

研究の背景

大規模並列計算機では、計算機の台数に応じた性能向上が要求される。特に、特定の問題サイズのアプリケーションにおいて計算時間を計算機の台数に反比例して減少させることのできる強スケーリング性は、様々な研究分野で必要とされている。この強スケーリング性の実現に向けた技術の一つに、並列計算に伴って発生する通信時間を計算とオーバーラップさせて隠蔽する非ブロッキング通信がある。特に、計算機全体で総和や収集、転置操作などの定型パターンの通信を実施する集団通信は、計算機の規模に応じて所要時間が増加するため、その隠蔽を可能にする非ブロッキング集団通信 (Non-Blocking Collective communications: NBC) の利用によるスケーラビリティの向上が期待されている。

しかし、現在この NBC の使用は限られたアプリケーションにとどまっている。その主な理由として、非ブロッキング集団通信による通信隠蔽効果に関する情報が不十分であることが挙げられる。計算と通信のオーバーラップにはアルゴリズムの修正が必要であるため、通信隠蔽による性能向上が期待できなければ、プログラマによる非ブロッキング集団通信の積極的な利用は期待できない。

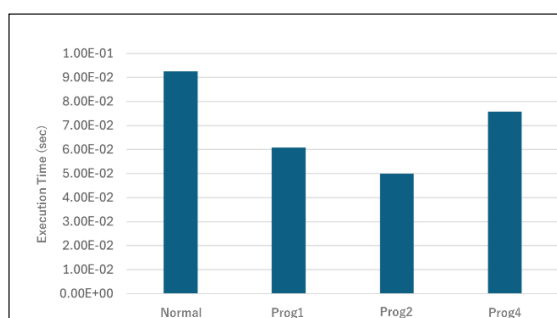
そこで、このような情報を提供するため、JHPCN のプロジェクトとして、日本の各計算機センターの研究者、米国および欧州の通信ライブラリ関係者、および日米のアプリケーション研究者とともに 2024 年度より非ブロッキング集団通信の各実装技術の利用方法の調査、基本性能評価方法の確立、および実アプリケーションにおける効果の検証に取り組んでいる。2025 年度の九州大学グループの主な成果は、プログレススレッドのコア割り当て最適化技術による NBC の実装と通信隠蔽効果の検証、および NBC の隠蔽効果計測用ベンチマークプログラムにおける計算量自動設定技術の実装である。それぞれ、以下に報告する。

プログレススレッドのコア割り当て最適化技術による NBC 実装

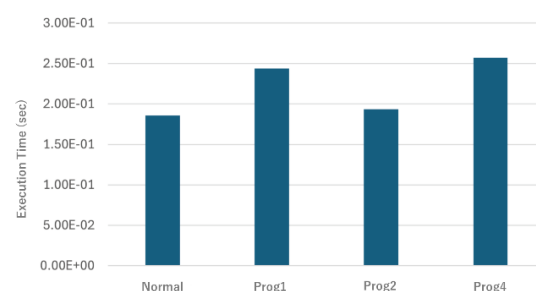
2024 年度のプロジェクトにおいて、NBC 向けの進行方式、特にオフロード機構である SHARP には、PPNs（ノード当たりのプロセス数）に関する制約があることが分かった。そこで今年度は、もう一つの進行方式であるプログレススレッドに着目し、コアごとに複数のプログレススレッドを割り当てるコア割り当て最適化技術を実装して通信隠蔽効果を検証した。この技術は、計算スレッドができるだけ多くのコアを利用可能となる。コアへのプログレススレッドの割り当てには、プログレススレッド内で `pthread_setaffinity_np` 関数を適用することで実現した。また、割り当て先のコアは、NUMA 構成に基づいて NUMA ノードに分散するように選択する。今回は、Flat-MPI での実行において、この割り当て技術による効果を評価した。そのため、以下の 4 つのケースにおける実行時間を比較した。

- Normal: オーバーラップなしで実行
- Prog1: 各プロセスに 1 つの進行スレッドを割り当て、1 コアあたり 1 スレッドで実行
- Prog2: 各プロセスに 1 つの進行スレッドを割り当て、1 コアに 2 スレッドをバインドして実行
- Prog4: 各プロセスに 1 つの進行スレッドを割り当て、1 コアに 4 スレッドをバインドして実行

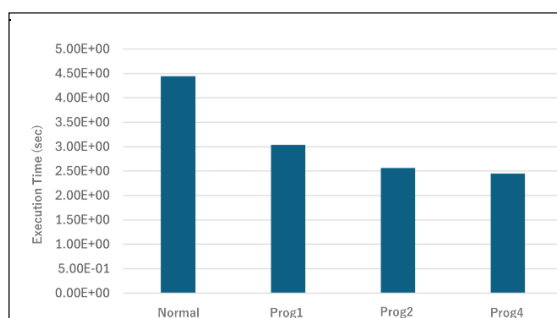
実験は、九州大学のスーパーコンピュータ玄界のノードグループ A において、32 ノードを用い、MPI_Allgather と行列積演算をオーバーラップさせる形で実施した。行列積演算の呼び出し回数は、計算時間の比率が Small-comp の場合には 1/2、Large-comp の場合には 2 となるように調整した。ノード当たりのプロセス数は、計算に利用可能なコア数に応じて決定した。その結果、Normal では 1 ノード当たり 120 プロセス、Prog1、Prog2、Prog4 ではそれぞれ 60、80、96 プロセスで実行した。図 1 および図 2 に、メッセージサイズ 16 バイトにおける Small-comp と Large-comp の結果をそれぞれ示す。図 1 から、Prog2 が他の方式よりも高い性能を示していることが分かる。これは、1 つのコアに 2 本の進行スレッドを割り当てることで、オーバーラップの効果が得られていることを意味する。一方で、Prog4 は、1 つのコア上で 4 本の進行スレッドを切り替える際のコンテキストスイッチのオーバーヘッドのため、Prog2 ほどの効果は得られなかった。図 2 では、通信に対する計算の比率が高くなるにつれて、Normal が最も高速となっている。これは、Normal では各ノードの 120 コアすべてを計算に利用するのに対し、Prog1、Prog2、Prog4 では通信のために一部のコアを消費するため、計算に使えるコア数が少なくなるためである。図 3 および図 4 に、メッセージサイズ 1MB の場合の結果を示す。図 3 から、Prog4 が他の方式よりも高い性能を示していることが分かる。したがって、メッセージサイズが大きくなると、コンテキストスイッチのオーバーヘッドは無視できる程度になる。また、図 4 に示すように、Prog4 は Normal よりも高速であった。この結果より、計算時間が通信時間より長い場合であっても、4 スレッドを一つのコアに割り当てることで通信隠蔽効果が得られることが分かった。



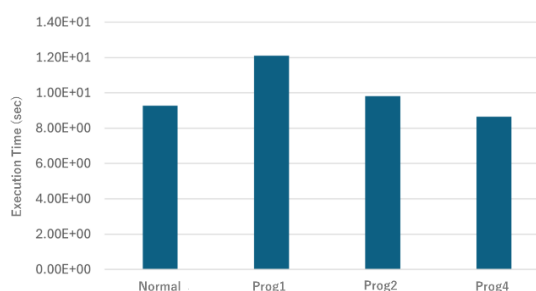
(図 1) 1 コアに複数のプログレススレッドを割り当てることによる通信隠蔽の効果 (16 byte Allgather + 計算比 1/2)



(図 2) 1 コアに複数のプログレススレッドを割り当てることによる通信隠蔽の効果 (16 byte Allgather + 計算比 2)



(図 3) 1 コアに複数のプログレススレッドを割り当てることによる通信隠蔽の効果 (1 MB Allgather + 計算比 1/2)

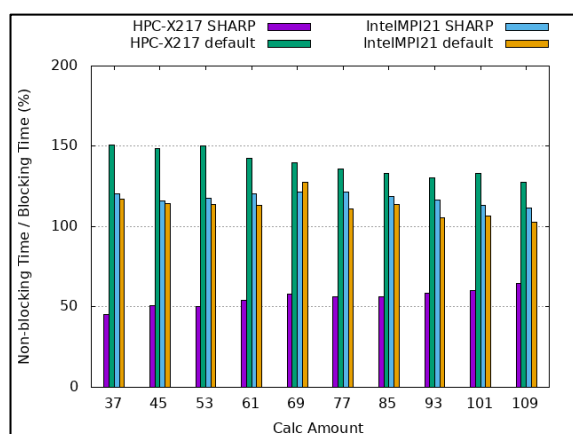


(図 4) 1 コアに複数のプログレススレッドを割り当てることによる通信隠蔽の効果 (1 MB Allgather + 計算比 2)

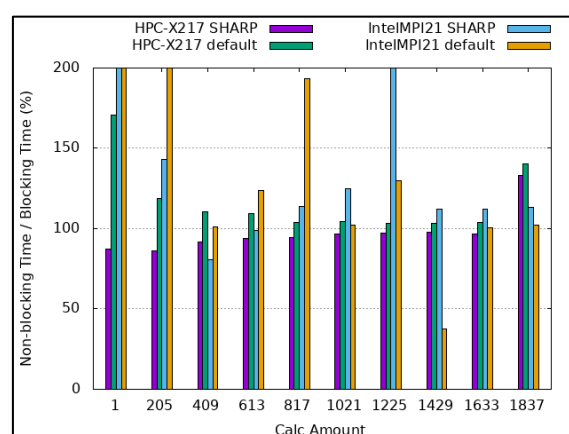
NBCの隠蔽効果計測用ベンチマークプログラムにおける計算量自動設定技術

2024年度において、異なるNBC実装間で、通信隠蔽がアプリケーション性能に与える効果を比較できるベンチマークプログラム NBC-Eff を構築した。今年度は、利用者の負担を軽減するため、対象となるすべてのNBC機構の通信速度に応じて、通信とオーバーラップさせる計算量を自動的に調整する仕組みを追加した。この仕組みでは、候補となる各実装の通信時間を測定し、それらのオーバーラップ効果を評価するのに十分な計算量の範囲を算出する。次に、その範囲を10分割し、各実装について実行時間を測定する。

図5に、1ノード当たり1プロセスの場合におけるIntel MPIとHPC-Xのベンチマーク結果を示す。この場合、SHARPを有効にしたHPC-Xが他よりも高い性能を示している。一方、図6には、1ノード当たり8プロセスの場合の同様の比較を示す。この場合、SHARPは顕著な優位性を示していない。また、計算量が増加するにつれて、Intel MPIはHPC-Xよりも良好な性能を示している。これらの結果は、本ベンチマークが、NBCによるオーバーラップの実際の効果を公平に比較するために利用できることを示している。



(図5) Intel MPIとHPC-XでのMPI_lallreduceの通信隠蔽効果の比較 (1プロセス/ノード)



(図6) Intel MPIとHPC-XでのMPI_lallreduceの通信隠蔽効果の比較 (8プロセス/ノード)

7.12 大島 聡史（先端計算科学研究部門 准教授）

All Core Computing と HPC 向けコード生成 AI 活用に関する研究

高性能計算向け GPU は多くの計算コアを搭載しており、それらをフル活用することで高い演算性能を発揮している。高速なメモリも搭載しているが、その性能を発揮するためにも多数の計算コアが動作する必要がある。スーパーコンピュータ（スパコン）に搭載された GPU ではそのコア数が 1 万を超えるため、そこで動かすプログラムは数万の並列度を有するものでなければ GPU の性能を使い切ることができない。一方で現実の計算にはそこまで高い並列度を持たないものも多い。そこで、小規模の計算を多数同時に行うことで総並列度を高めて計算効率を改善する手法が用いられている。これはバッチ計算と呼ばれており、特に小規模な密行列同士の行列積和演算は深層学習において重要な計算であるため、主要な数値計算ライブラリに専用の関数（バッチ計算関数）が用意されており、ハードウェアもそれに合わせた性能向上を続けている。NVIDIA 社の GPU に搭載されている Tensor Core はその最たるものである。しかし、バッチ計算は単一のプロセスとメモリ空間から全ての計算を発行する設計が基本であり、例えば MPI+OpenMP によるハイブリッド並列化プログラムの GPU 化には活用しにくい場合がある。

本研究では低ランク近似行列計算法の GPU を用いた高速化に継続して取り組んでいる。低ランク近似行列計算法は、密行列（の一部）を使用メモリ量の少ない行列に置き換えて（近似して）保持や計算を行う計算法であり、これを用いるとメモリに乗りきれないような巨大な密行列を扱うことが可能となる。現在は低ランク近似行列の 1 つであるブロック低ランク行列（Block Low-Rank, BLR）に対する QR 分解（BLR-QR）を主な計算対象としている。BLR-QR の内部では行列積和演算など主要な数値計算ライブラリにより提供されている計算を多数行っており、一見するとバッチ計算関数により容易に高速化が行えそうな問題にも見える。しかし我々の計算対象は MPI+OpenMP によるハイブリッド並列化プログラムを元にしており、既存のバッチ計算関数を適用するにはプログラムの大改造が必要となってしまう、実用的ではない。そこで本研究では、複数のプロセスで GPU を共有利用し、その際に各プロセスが利用できる計算資源を制限する機能である Multi-Process Service(MPS) や Multi-Instance GPU(MIG) を用いたマルチプロセス GPU 実行によるプログラムの高速化に取り組んでいる。GPU 資源を多く使用するのが非効率な計算を逐次的に行うよりも、GPU 資源を多く使わない効率的な計算を並列的に行うほうが、総計算時間を短くできるだろうというのが基本的な発想である。2025 年度にはこの計算法を NVIDIA GH200 に適用し、得られた成果や知見を国際会議 [1] やシンポジウムなど [2] で発表した。

従来の GPU プログラム最適化においては、CPU に GPU を邪魔させないことが基本的な最適化戦略の一つであった。一方 GH200 は CPU と GPU が一つにパッケージ化されたプロセッサであり、CPU-GPU 間のデータ通信速度が速く、従来の CPU・GPU 環境と比べてお互いの有するメモリに自由にアクセスできる。さらにマルチプロセス GPU 実行により複数プロセスが GPU を共有・分割利用すると、プロセスあたりの GPU 性能と CPU 性能の差が小さくなり、CPU も使わなければもったいない状況も生じる。AMD も高性能な CPU・GPU 混載プロセッサをリリースしており、同様の状況と言える。さらに現在の GPU は従来の CPU・GPU と比べてより低精度な演算の性能を重視したハードウェアへと進化しており、CPU と GPU の得意とする計算の差異は大きくなりつつある。こうした状況から、計算ノードの性能をフルに発揮するには CPU と GPU を全て活用した計算、All Core Computing（造語）の重要性が高まっていると考え、本研究ではその技術開発を進めている。現在は前述の BLR-QR に対するマルチプロセス GPU 実行をベースに、CPU へも計算を割り当ててさらなる高性能化を目指している。

一方で現在急速にコード生成 AI 技術が進歩しており、HPC 向けの並列計算コードも高いレベルで生成可能となりつつある。しかし大規模なプログラムの十分な最適化を容易に行えるレベルまでは到達していないと考えられており、HPC 向けのコード生成 AI 技術について様々な研究が進められている。低ランク近似行列計算法や All Core Computing で用いているプログラムはやや難しい（典型的ではない）HPC コードであるため、これらのプログラミングにおけるコード生成 AI の有用性調査およびコード生成 AI 技術の評価を行っている。最新の状況を国際会議 [3] で報告した。

発表論文等

- [1] Satoshi Ohshima, Akihiro Ida, Masatoshi Kawai, Takeshi Fukaya, Rio Yokota, “A Study on the Performance and Usability of Managed Memory and Unified Memory for Accelerating Numerical Calculation Program”, IEEE, 2025 IEEE 18th International Symposium on Embedded Multicore/Many-core Systems-on-Chip (MCSoc), 2025.12.31 published, pp.41-48, doi.org/10.1109/MCSoc67473.2025.00017
- [2] 大島聡史, 伊田明弘, 河合直聡, 深谷猛, 横田理央, “BLR-QR における Managed Memory や Unified Memory の有用性と性能について”, 第 17 回 自動チューニング技術の現状と応用に関するシンポジウム (ATTA2025), 工学院大学新宿キャンパス, 2025.12.23
- [3] Satoshi Ohshima, Shunsuke Nakano, Yusuke Endo, “Optimization of GEMM using AMX, and can Code Generative AI generate its optimal code?”, 2026 Conference on Advanced Topics and Auto Tuning in High-Performance Scientific Computing, Cosmology Building, National Taiwan University, 2026.03.20-21

7.13 樋口 祐次（先端計算科学研究部門 准教授）

ソフトマターの分子シミュレーション

(1) ソフトマター周囲の水の動態

ソフトマター周囲の水の動態は、生体材料の生体適合性や生体分子の機能発現との関係が指摘されている。このため、生体材料の分子設計や、生体分子の機能理解のためには、ソフトマター周囲の水分子の動態解明が重要な課題となっている。そこで、双極性イオンを側鎖にもつ高分子周囲の水の動態に関して量子化学計算に基づく分子動力学シミュレーションを用いて調べた。興味深い発見の一つとして、高分子化することにより、単独のモノマーでは見られない、側鎖間による水分子のトラップ現象を明らかにした [1]。

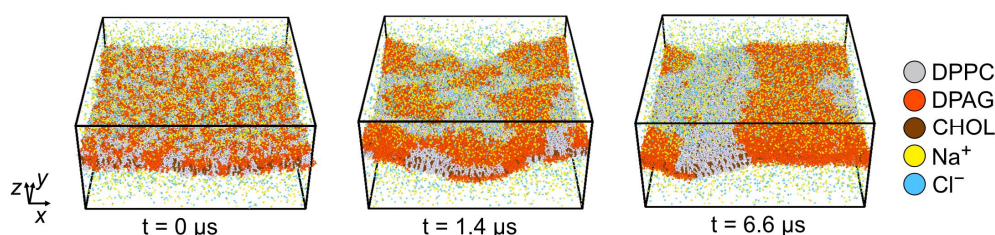
(2) リン脂質二重膜の相分離とイオン分布の相関

リン脂質二重膜は生体膜のモデルとして研究が行われている。生体中では多くのイオンが含まれていること、リン脂質分子も電荷を持つため、互いに相関しながらイオン分布や構造が決定されている。このため、生体膜の構造やその物性理解には、生体膜近傍でのイオン分布解明が求められるが、実験的に直接観測することが困難な対象である。そこで、粗視化分子動力学法を用いて、リン脂質分子の相分離過程におけるイオン分布変化を調べた（図1）。相分離界面において、イオン分布変化はシグモイド関数でフィッティングできることが分かった。荷電脂質分布とイオン分布の相関を定量化し、理論計算とも比較した。荷電脂質膜と電解質溶液の相互作用に関する分子レベルの理解に貢献した [2]。

(3) 高分子材料におけるトポロジー解析

ゴムは製造過程において架橋されるが、実際には架橋点間距離、架橋点の空間分布、架橋効率などは理想的に制御できないため、それら特徴量と力学特性の相関関係の解明が求められている。そこで、架橋により作成された複雑な網目構造に対してグラフ理論に基づくトポロジー解析を行うことで、低ひずみでは独立な網目構造の総数が、高ひずみでは架橋の均一性が重要な要素となることを明らかにした [3]。

1. M. A. Saleh, Y. Higuchi*, S. Shiimoto, T. Anada, M. Hishida, M. Tanaka*, Polym. J. 57, 1127-1139 (2025).
2. Y. Higuchi*, H. Ito, N. Shimokawa, J. Reščič, S. Maset, and K. Bohinc*, J. Mol. Liq. 444, 129168 (2026).
3. Y. Ishikura, E. Uehara, Y. Higuchi*, Soft Matter 22, 2704-2712 (2026).



(図1) 粗視化分子動力学シミュレーションによるリン脂質二重膜(生体膜モデル)の相分離過程。異なる色は異なる種類のリン脂質分子を表し、時間の経過とともに同種の分子が集まって相分離していく様子を示している。

7.14 小出 洋（情報システムセキュリティ研究部門 教授）

AI と Moving Target Defense を活用したサイバーセキュリティ研究

近年のサイバー攻撃は、標的型攻撃の長期潜伏化、IoT 機器の急速な普及、クラウドや分散基盤の複雑化などにより、従来以上に高度化・複雑化している。とりわけ、侵入後に長期間潜伏しながら情報窃取を行う Advanced Persistent Threat (APT) は、重要組織に対して大きな脅威となっている。また、IoT 環境では、計算資源や電力に制約のある機器が多数接続されるため、重い防御機構をそのまま適用することが難しい。さらに、複雑なネットワーク環境に対する脆弱性診断やペネトレーションテストにおいても、網羅性と効率性をどのように両立させるかが重要な課題となっている。このような背景の下、本研究室では、AI を用いた攻撃検知技術、Moving Target Defense (MTD) に代表される能動的防御技術、および自動化された攻撃評価技術を組み合わせ、攻撃を受けにくく、かつ侵入後も被害を抑制できる情報システムの実現を目指して研究を進めている。

第一に、自動ペネトレーションテストに関する研究を進めた。2025 年度の成果として取りまとめた Multi-Objective Policy Optimization for Effective and Cost-Conscious Penetration Testing では、複雑なネットワーク環境に対して有効かつ低コストな侵入試験方策をどのように学習させるかを課題とした [1]。実環境に近いネットワークでは、深い侵入や重要資産への到達を狙うほど、多くの試行回数、時間、通信量、検知リスクを要する傾向がある。一方で、過度に保守的な方策では重大な脆弱性を見逃す可能性がある。そこで本研究では、報酬最大化とコスト最小化を同時に扱う多目的強化学習の枠組みを導入し、Lagrangian に基づく方策最適化によって、攻撃効果と運用コストのバランスを動的に調整する手法を提案した。各種ベンチマーク環境で評価した結果、提案手法は高い侵入性能を維持しつつ、時間コストの削減に有効であることを示した。これは、自動化された脆弱性評価をより実践的なものへと近づける成果である。

第二に、AI を活用した情報窃取攻撃の検知に関する研究を進めた。Detecting Advanced Persistent Threat Exfiltration With Ensemble Deep Learning Tree Models and Novel Detection Metrics では、APT によるデータ持ち出しを、攻撃の初期段階ではなく、実際の通信転送段階で検知することを目的とした [2]。既存研究の多くは侵入の兆候や特定のトンネル通信に着目しているが、侵入の早期検知に失敗した場合には十分に機能しない可能性がある。そこで本研究では、攻撃者が既に内部に侵入している状況を想定し、外向き通信の振る舞いそのものに注目することで、秘匿性の高い情報流出を検知する手法を検討した。具体的には、コマンド・アンド・コントロール通信への偽装や、小分割転送による隠蔽といった APT の特徴を分析し、それらを捉えるための新たな検知指標を導入したうえで、アンサンブル深層学習モデルによる分類器を構築した。その結果、高い検知性能と実用的な計算効率の両立が可能であることを示し、侵入後対策として有効なネットワーク監視方式の実現可能性を示した。この研究は、侵入を完全に防ぐことのみを前提としない防御設計の重要性を示すものである。すなわち、侵入後の挙動監視を継続的に行い、情報流出段階で被害を把握・抑止するという観点から、大学情報基盤や重要インフラにも適用可能な実践的監視技術の基盤を与えている。特に、精度のみならず、安定性や誤検知・見逃しの扱いを意識した評価を進めた点に特徴がある。

第三に、Moving Target Defense に基づく能動的防御技術の研究を進めた。Proposal and Implementation of a Security Enhancement Method using Route Hopping MTD for Mesh Networks では、IoT メッシュネットワークにおいて通信経路を動的に変更するルートホッピング型 MTD の提案と

実装を行った [3]。メッシュネットワークは耐障害性や通信範囲の拡張に優れる一方、異なるノードを経由してパケット転送を行うため、単一の侵害ノードに起因する盗聴、改ざん、DoS などの脅威にさらされやすい。そこで本研究では、AODV を拡張し、経路探索時に複数の候補経路を確保したうえで、転送時に経路を動的に切り替える方式を実装した。これにより、固定的な経路への依存を避け、単一ノードを起点とする中間者攻撃や妨害攻撃への耐性向上を図った。実機（Raspberry Pi）上での実装と評価を通じて、提案方式の実装可能性と、その防御的有效性を確認した。

これら三つの研究は、それぞれ異なる層におけるセキュリティ課題に対応しているが、全体としては相補的な関係にある。自動ペネトレーションテストは、システム運用前または運用中の脆弱性把握と評価の高度化に資する。APT 情報流出検知は、侵入後の被害把握と早期対処に有効である。ルートホッピング型 MTD は、通信経路そのものを動的に変化させることで攻撃成立条件を揺さぶる能動的防御として機能する。すなわち、脆弱性評価、侵入後監視、通信防御という三つの観点から、より抗堪性の高い情報システムを構成するための基盤技術を提示している。

加えて、著者らは研究成果を個別の論文発表にとどめず、教育・実装・運用の接点に接続することを重視している。サイバーセキュリティ分野では、理論上有効な方式であっても、運用上の負荷や導入コストの問題により実装に至らない例が少なくない。そのため、著者らは攻撃モデルの設定、検知・防御方式の提案、プロトタイプ実装、性能評価までを一貫して行うことで、現実の情報基盤への適用可能性を検証している。今後も、情報システムや IoT 環境を対象とした実環境評価を進めるとともに、AI と MTD を含む要素技術を組み合わせた実践的なセキュリティ技術の高度化を継続していく方針である。

研究業績

1. Xiaojuan Cai, Lulu Zhu, Zhuo Li, Hiroshi Koide, "Multi-Objective Policy Optimization for Effective and Cost-Conscious Penetration Testing," In Proceedings of the 21st International Conference on Web Information Systems and Technologies (WEBIST 2025), pp. 374-385, 2025.
2. Xiaojuan Cai, Haibo Zhang, Chuadhry Mujeeb Ahmed, Hiroshi Koide, "Detecting Advanced Persistent Threat Exfiltration With Ensemble Deep Learning Tree Models and Novel Detection Metrics," IEEE Access, 2025.
3. Yuto Ikeda, Hiroshi Koide, "Proposal and Implementation of a Security Enhancement Method using Route Hopping MTD for Mesh Networks," CYBER 2025, pp. 24-30, 2025.