

3-2-4 プライバシー保護連合学習技術 DeepProtect

3-2-4 DeepProtect: Privacy-preserving Federated Learning System

Le Trieu Phong 王立華 金森祥子 小野元 野島良 大久保美也子
LE Trieu Phong, WANG Lihua, KANAMORI Sachiko, ONO Hajime, NOJIMA Ryo, and OHKUBO Miyako

現代のあらゆるシーンで大量のデータを活用して新たな価値を創出する試みが進んでいる。しかし、銀行や病院が所有する個人情報や機密を含むデータは組織を横断した共有に制限がかけられており、データの価値を引き出すことが難しい。セキュリティ基盤研究室は暗号技術を取り入れたプライバシー保護連合学習技術「DeepProtect」を開発した。機密データを共有せずに深層学習を実現するこの技術はデータの安全を担保しつつ複数の組織に分散したデータの真の価値を引き出すことができる。本稿では DeepProtect の技術的特徴を説明するとともに、応用と実証実験についても紹介する。

The use of personal and confidential data, such as that held by banks and hospitals, has prompted concerns about privacy breaches and data leaks, thereby restricting data sharing across organizations. To address this, Security Fundamentals Laboratory has developed a federated learning technology called “DeepProtect” that incorporates advanced encryption methods. This technology facilitates deep learning without the need to share confidential data, thus maintaining privacy and ensuring data security while enabling the extraction of valuable insights from data distributed across various entities. This paper details the technological aspects of “DeepProtect,” its applications, and the outcomes of pilot experiments utilizing the system.

1 まえがき

様々な組織や業態でデータが大量に蓄積される昨今、多様なビジネスシーンで、大量に蓄積されたデータを新たな価値・ビジネスにつなげるために利活用しようという試みが始まっている。一方、蓄積したデータが個人情報であったり機密性の高いデータが含まれている場合、利活用という観点ではハードルが高く、十分に活用しきれていないという現状がある。例えば、振り込め詐欺やマネーロンダリングなどの金融犯罪は巧妙化が著しく、深刻な社会問題となっているが、多くの金融機関では、それぞれが保有する金融取引データに対して、個別にルールベースのモニタリングツールを用いて対策を行っているのが現状である。各々の金融機関で蓄積しているデータを利活用した機械学習技術を用いた不正取引の自動検知システム導入へのニーズが高まってはいるものの、ひとつの金融機関では学習データ量が十分ではなく検知精度が上がりづらい。複数の金融機関でデータを統合すれば十分な学習ができ、検知精度は上がると期待されるが、個人情報保護が課題となり、他の金融機関とのデータ連携が難

しく、複数の金融機関のデータの統合は簡単ではない。

上述のような状況において、セキュリティ基盤研究室はフェデレーテッドラーニング(連合学習)という技術に暗号技術を融合し、パーソナルデータなど機密性の高いデータを互いに開示することなく安全に深層学習を用いて解析することができるプライバシー保護連合学習技術 DeepProtect を開発した。DeepProtect は、複数組織間で連合して深層学習を行う際に、組織外部に送信する情報(深層学習のパラメータ)を統計情報化し、かつ、暗号化することによって個人識別ができない状況で統合し、各組織の学習モデルを更新することを可能とする。現在、セキュリティ基盤研究室は、DeepProtect を活用して金融分野における不正送金の自動検知システムの実現に向けた実証実験を進めている。

今回は、DeepProtect の仕組みや技術的な特徴・機能を紹介するとともに、プライバシー機能の拡張技術やそれらの知見を活かした深層学習アルゴリズムへの応用について紹介し、最後に DeepProtect を用いた実証実験について紹介する。

2 DeepProtect の誕生

データを有効活用するには大量のデータが必要になると考えられている。ただし、プライバシー上の懸念やEUの一般データ保護規則(General Data Protection Regulation; GDPR)*¹や日本の個人情報保護法*²などの規制により、データが他の組織と共有できない可能性があるため、課題が生じる。

プライバシー保護連合学習技術は、複数の組織が持つデータセットでデータプライバシーを保護しながら機械学習モデルを学習できるようにする機械学習手法である。基本的に、機械学習モデルは各組織で個別に学習される。その後、学習されたモデルは中央サーバと安全に共有される。中央サーバは、各組織からの更新モデルを安全に結合し、グローバルモデルを更新する。この仕組みを繰り返すことにより、生データを他の組織と共有せず安全に機械学習ができるようになる。

暗号は、データを秘密鍵無しでは読み取り不可能な形式に変換し、プライバシーを保護するために広く採用されている技術である。プライバシー保護連合学習技術 DeepProtect [1][2] は、暗号を使用し、中央サーバに情報を漏えいしないようにする。

3 理論的な進展と今後の展望

3.1 DeepProtect の仕組み及び性能向上・拡張

機械学習を実現するためのアプローチは複数存在する。中でもニューラルネットワークを用いてレイヤー化された構造の中で学習するアプローチとして深層学習がある。DeepProtect はこのアプローチに類する手法を取っている。

また、深層学習では、よりよいパラメータ(「重み」とも呼ぶ)を見つけるために「確率的勾配降下法」と呼ばれる手法が使われる。確率的勾配降下法は、最適化アルゴリズムであり、モデルのパラメータを、トレーニングデータのランダムに選択された部分集合に関する損失関数の勾配を計算することで更新する。損失関数は、モデルの予測がトレーニングデータの実際の目標値にどれだけ適合しているかを定量化し、モデルのパフォーマンスを評価する尺度として機能する。しかし、損失関数の勾配情報から元々のデータの情報がどの程度漏れるのか、については明らかにされていない。

Phong らは、[1][2]や後続の研究[3]等で、損失関数の勾配によりデータの情報が漏えいする可能性があることを示している。通常、学習するデータセットはいくつかのグループに分けて学習を実施する。この分けられたグループをバッチと呼び、各グループのサイズ

をバッチサイズと呼ぶ。文献[1][2]で観察されるように、学習のバッチサイズが1の場合、ニューラルネットワークに関わる損失関数 J は、次の数式を満たす：

$$\frac{\delta J}{\delta w_{ik}^{(1)}} = \xi_i \cdot x_k$$

ここで、 $w_{ik}^{(1)}$ は、ニューラルネットワークのレイヤー1の入力 x_k とレイヤー2の隠れノード i を接続する重みパラメータである。また、 ξ_i はバイアス項に関連付けられた勾配を表す。したがって、データ x_k は損失関数 J の勾配から漏れる可能性がある。

学習にあるバッチサイズが1より大きい場合、勾配には上記の方式の右辺のような形式を持つ複数の情報の合計が含まれるため、情報が乱され、勾配からの情報漏えいを防ぐことができる。実際、文献[4]で証明されているように、データの機密性はサブセットと問題に基づいている可能性があり、バッチサイズ値が大きいほどデータのプライバシーと機密性に優れているという観察も裏付けられる。さらに、文献[4]で提案されているように、ローカルトレーニングを通じてバッチサイズを効率的に大幅に増やすことが可能である。

さらにセキュリティとプライバシー保護を強化するために、暗号技術を組み込むことは有用な手段のひとつである。我々は機械学習の学習プロセスの処理と親和性の高い(加法)準同型暗号を導入した。(加法)準同型暗号方式によりデータを暗号化したまま学習を行うことが可能となる。これらの措置は、データの機密性を保護し、プライバシーのリスクを軽減するために実施される。また、プライバシーを保護する連合学習システムを設計及び評価するときは、セキュリティとプライバシーと実用性のバランスを考慮する必要がある。

DeepProtectシステム概要を図1に示す。DeepProtectには、 N 組織(例：銀行)の連合学習参加者と中央サーバがあり、その動作は以下のとおりである。連合学習参加者はデータセットをローカルに所有し、サーバや他の参加者などの外部者にデータセットを送信しないが、学習モデルのコピーをローカルで実行し、多くの通信ラウンドでサーバと対話する。ニューラルネットワークの初期のランダムな重み W_{global} は、例えば学習参加者Aにより初期化され、最初に準同型暗号 $\text{Enc}(W_{global})$ をサーバに送信する。暗号化 Enc の目的は、サーバに対して重みの機密性を保護することである。どの通信ラウンドでも、学習参加者はサーバから最新の暗号化された $\text{Enc}(W_{global})$ を取得し、復号を

*1 <https://gdpr-info.eu/>

*2 <https://www.ppc.go.jp/index.html>

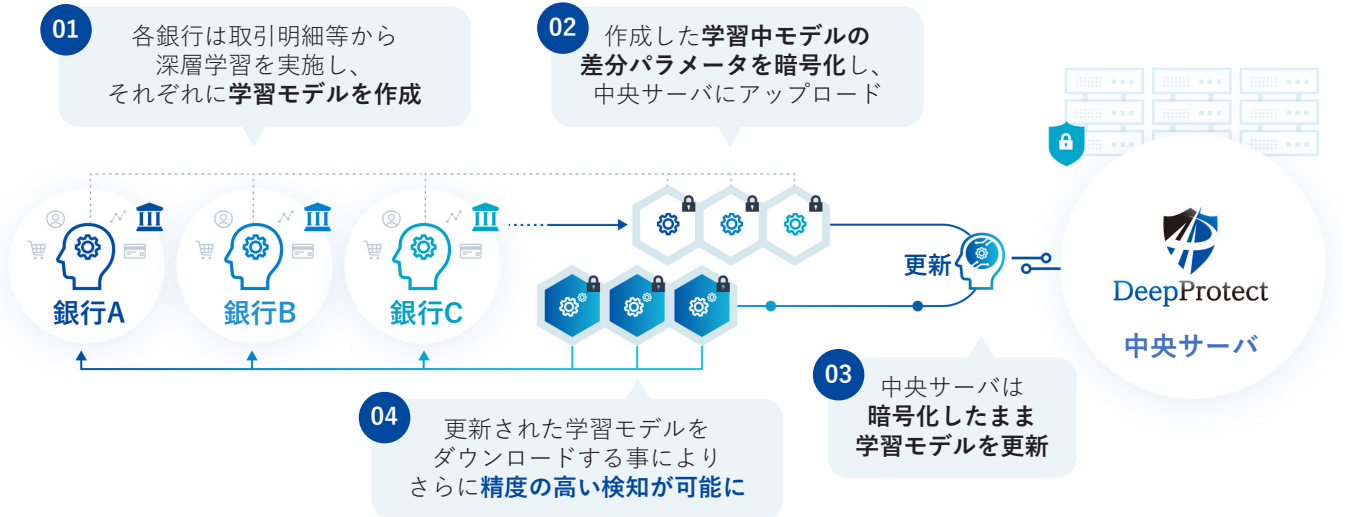


図1 プライバシー保護連合学習技術 DeepProtect

実行して W_{global} を取得する。そして、 W_{global} によるニューラルネットワークと学習参加者 i でのローカルデータバッチの各実行後に得られる勾配ベクトル $G(i)$ に学習率 α が乗算される。次に、各学習参加者からの暗号化された $\text{Enc}(-\alpha \cdot G(i))$ がサーバに送信される。サーバでは、暗号化 Enc の加法準同型特性を利用して、暗号化された重みパラメータを再帰的に更新する。特に、サーバは次の計算を行う：

$$\text{Enc}(W_{global}) + \text{Enc}(-\alpha \cdot G(i)).$$

したがって、重み W_{global} は、 $W_{global} - \alpha \cdot G(i)$ となる。表記上

$$W_{global} \leftarrow W_{global} - \alpha \cdot G(i)$$

として更新され、これは機械学習によく使われている確率的勾配降下法 (stochastic gradient descent, SGD) と同一である。

上記の暗号化 Enc は加法準同型であることのみが必要であり、公開鍵ベースと秘密鍵ベースの両方を使用できる。前者のインスタンス化については、文献 [2] で詳しく説明されている。後者は、文献 [5] に次のようにワンタイムパッドタイプで構築できる： $\text{Enc}_K(W) = W + \text{PRF}_K(\text{iter})$ 。ここで、 K は学習参加者間で共有される秘密鍵であり、サーバには知らせない。 PRF は安全な擬似ランダム関数、 iter はすべての関係者に知られている固有の通信ラウンドインデックスである。準同型特性がない場合でも秘密鍵暗号を使用する方法は、文献 [4] に記載され、さらに不正な学習参加者や学習参加者とサーバの不誠実な共謀など、比較的複雑なシナリオで重みを保護する方法についても説明されている。

3.2 差分プライバシーを取り入れたプライバシー保護機能の強化

差分プライバシー [6] と呼ばれている技術は、外部の観察者が特定の個人の情報がデータセット内に存在するかどうかを判断できないことを保証する。差分プライバシーは、入力データ、計算結果、またはアルゴリズム自体に適用できる。これらのアプローチは、特定のアプリケーションに応じて、各学習参加者によってローカルで実行される。DeepProtect に学習参加者間で通信する情報に差分プライバシーを追加するアプローチを検討している [7]。

まず、Laplace 分布や Gaussian 分布に従う乱数 (Laplace ノイズ、Gaussian ノイズ) を使用することで、DeepProtect に差分プライバシーを追加することができる。例えば、学習参加者間で共有される重みは、様々な手法を使用し差分プライバシーにすることができる。文献 [8]-[10] に示すように、各学習参加者で重みパラメータを更新するときに、ローカル勾配情報に Laplace や Gaussian ノイズを追加する。もう 1 つのアプローチは、出力摂動を使用することである。さらに、ラベルなしの公開データが利用可能な場合、各学習参加者は [11] で説明されている手法をローカルで適用できる。差分プライバシーの合成定理により、複数の差分プライベート重みが続いて送信される場合に使用できる。差分プライバシーは、特定のデータ項目に対する共有重みの依存性を軽減し、それによって各ローカルデータセットの保護を強化する。その一方、トレーニングプロセスにノイズが導入されるため、学習精度とプライバシーの間にトレードオフが生じる可能性がある。

上記の差分プライバシー技術のほか、ドロップアウト技術を使うアプローチもある。ドロップアウトは、

ニューラルネットワークのトレーニング中に使用される手法で、ニューラルノードとその接続がランダムにドロップされ、各ノードのトレーニングデータへの依存度を下げるのに役立つ。これは、文献 [12] で観察されているように、差分プライバシーと同様の効果をもたらす可能性がある。

3.3 学習方法の多様化

一般的に深層学習により検出された結果について、なぜその結果が抽出されたのかの要因を明示すること（説明可能性）は難しい。一方、決定木を用いるアプローチではある程度の因果関係を示すことが可能であるとされている。そこで、AIの説明可能性と効率性の観点から、勾配ブースティング決定木 (GBDT: Gradient Boosting Decision Trees) に基づく効率的な連合学習プロトコル (eFL-Boost) [13] の採用も検討している [5]。

eFL-Boost は、決定木の構築を次の二段階に分割する。

①木構造の決定：ビルダーと呼ばれる役割に選ばれた1組織によって行われる。グローバルモデルの更新の際に生じるバイアスを抑制するために、毎回 N 組織から異なるビルダーによって行われる。ビルダーは、自身のデータセット及び更新前のグローバルモデルを用いて GBDT における一般的な決定木構築アルゴリズムを実行することによって、最新のローカルモデル（すなわち、決定木 T ）の木構造 $\mathcal{T}$ のみを構築し、他の参加組織と共有する（図2）。

②葉の重みの計算：モデルが更新される度に全参加組織が連合して共同で行う。

参加組織 $d \in \{1, \dots, N\}$ は、自身のデータセットを用いて、 $\mathcal{T}$ の各葉に対して、損失関数 L の勾配 $g_i = \partial_{\hat{y}_i} L(y_i, \hat{y}_i)$ 及び $h_i = \partial_{\hat{y}_i}^2 L(y_i, \hat{y}_i)$ の合計、 $G_d = \sum_{i \in D} g_i$ と $H_d = \sum_{i \in D} h_i$ を計算する。次に、加法準同型暗号を用いてこれらの結果を暗号化し、暗号文 $\text{Enc}(G_d)$ 及び $\text{Enc}(H_d)$ をサーバに送信する。

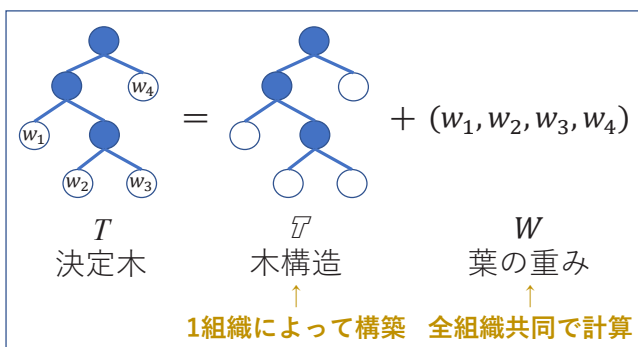


図2 決定木モデルは木構造及び葉の重みで構成

サーバは、加法準同型特性を利用して、暗号化された勾配値の集約を行い、その結果

$$\text{Enc}(G) = \sum_{d=1}^N \text{Enc}(G_d) \text{ 及び}$$

$$\text{Enc}(H) = \sum_{d=1}^N \text{Enc}(H_d)$$

を全参加組織に送る。各参加組織は、 $\text{Enc}(G)$ と $\text{Enc}(H)$ から G と H を復号し、葉の重み $W = -\frac{G}{H+\lambda}$ を計算する。そして、葉の重み W を木構造 $\mathcal{T}$ に追加することで完成した決定木モデル T を、グローバルモデルに追加することで、グローバルモデルは更新される。

eFL-Boost [13] は先行研究の特殊な GBDT ベース連合学習方式 FL-XGBoost [14] を発展させたスキームであり、予測精度の向上、通信コストや他組織に開示する情報の低減を実現した。

さらに、分析対象の状況は逐次変化することから、学習モデルも更新されていくことが望ましいと考えられる。このように継続的に学習モデルを更新していく学習方法として継続学習がある。逐次状況に対してもっとも適切な検出を行うことを目指し、継続学習の機能を追加する検討も行っている。

中央サーバ不要な継続学習プロトコルの確立：全ての連合学習参加者が納得できる中立的立場の中央サーバを確保することは困難であり、この問題のために連合学習に参加できない組織が発生し得る。この問題を解決するために、Blockchain 技術を使った、中央サーバ不要の連合学習プロトコルの研究に取り組んだ。決定木の連合学習方法である FL-XGBoost と eFL-Boost に Blockchain 技術を取り入れることで、中央サーバが存在せずとも、連合学習参加者だけでグローバルモデルを継続的に更新できるプロトコルを提案した [15]。この提案プロトコルでは分類問題において優れた予測精度を維持しつつ中央サーバをプロトコルから排除することができた。

忘却防止のための工夫：継続学習における課題の1つに、継続学習を行っていると新しいデータに対する予測精度だけが長く古いデータに対する予測精度が落ちてしまう「忘却」と呼ばれるものがある。この問題を克服するために、eFL-Boost に古いデータも再利用（リプレイ）する機能を追加した学習アルゴリズムを研究中である [16][17]。このアルゴリズムを DeepProtect に導入すれば、セキュリティ強化機能と高い継続学習能力が両立できると期待される。

決定木における差分プライバシーアプローチ：トレーニングされた決定木モデルから学習で利用したデータの漏えいを防ぐため、我々は Li らが提案したデータセットの k -匿名と差分プライバシーの関係 [18] を活用し、決定木学習において枝刈りの匿名化への効果に着目し、属するデータ数が少なく識別リスクが高

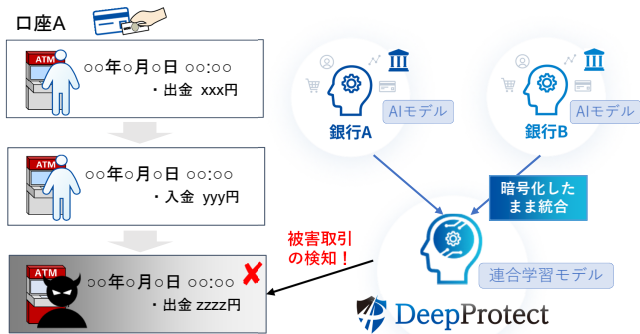


図3 被害取引の検知

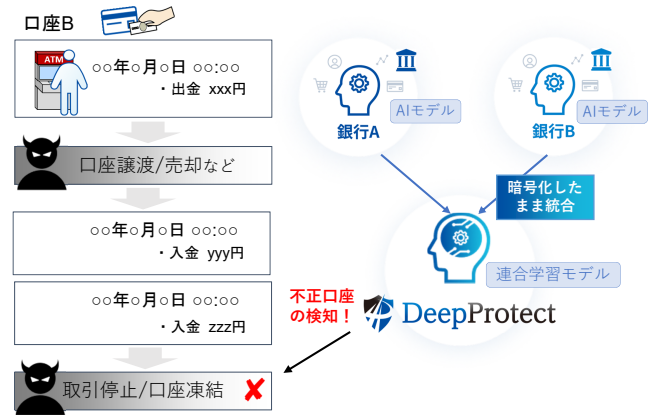


図4 不正口座の検知

くなる葉を剪定することで、ノイズ無しで差分プライバシーを満たす決定木の構築方法を提案した [19][20]。

4 社会実装に向けた活動 (CREST, AIP 加速, etc)

近年、金融犯罪手法は複雑化・巧妙化しており、早急な対策が求められている。中でも、振り込め詐欺等の特殊詐欺は深刻な社会問題となっている。対策の一つとして、AIを用いた不正取引の自動検知技術などが注目を集めている。精度の高い検知結果を得るためには、大量のデータによる学習が不可欠となる。しかし、単独の金融機関では十分な量の学習データを用意することはなかなか容易なことではない。複数の金融機関がデータを持ち寄れば、十分な量の学習データが準備できるが、個人情報を含む金融取引データを各金融機関外に持ち出すことはできないため、容易に実現することはできない。

DeepProtect では、暗号化したデータ（勾配情報）を用いて学習を行うことから、複数の金融機関による学習への障壁が大幅に軽減される。

セキュリティ基盤研究室、神戸大学及びエルテスは、データを外部に開示することなく機密性を保ったまま機械学習を行う独自開発のプライバシー保護連合学習技術 DeepProtect 等を活用し、複数の金融機関（千葉銀行、三菱 UFJ 銀行、中国銀行、三井住友信託銀行、伊予銀行）と連携して不正送金等を自動検知するシステムの実現を目指し、実証実験に取り組んだ。（報道発表:「プライバシー保護連合学習技術を活用した不正送金検知の実証実験を実施～被害取引の検知精度向上や不正口座の早期検知を確認～」国立研究開発法人情報通信研究機構、国立大学法人神戸大学、株式会社エルテス（2022年3月10日））

各銀行の検知目的に応じて、被害取引の検知を目的とする「被害検知グループ」(図3)と、不正口座の検知を目的とする「加害検知グループ」(図4)とで、それぞ

れ実証実験を行った。

被害検知グループの実証実験には2行の銀行が参加し、「単独組織のデータのみを用いた通常の機械学習モデル（個別学習モデル）による検知精度」と「2行のデータを用いた DeepProtect モデル（連合学習モデル）による検知精度」の比較を行った。図5の例に示すように、連合学習モデルにより検知精度が向上し、1日当たりの被害取引のアラート件数を600件としたときの検知精度は、当初目標としていた80%以上を達成した。また、個別学習モデルでは検知できなかった不正取引が検知された事例も確認した。なお、検知率は、検知漏れの少なさを表す指標である再現率で示している。

加害検知グループの実証実験には4行の銀行が参加し、「通常取引に使われていた口座が、あるタイミングから犯罪に悪用され、取引停止・凍結されたもの」を不正口座と定義し、不正口座を検知することとした。個別学習モデルとハイブリッドモデル（個別学習モデルと連合学習モデルの組合せ）で比較し、モデルの性能を検知率に加えて、不正口座に対して凍結時点より何週前に検知できるかという検知タイミングで評価した。その結果、図6の例に示すように、個別学習モデルと比較してハイブリッドモデルが高い性能を示し、検知率80%以上を達成した。さらに、実データでの不正口座凍結よりも20～50週程度の早期検知が可能であることが確認された。

本実証実験は、2019年度から JST CREST「イノベーション創発に資する人工知能基盤技術の創出と統合化」の加速フェーズ研究課題として採択された研究課題「プライバシー保護データ解析技術の社会実装」（課題番号 JPMJCR19F6）の下で実施した。

2022年度からは JST AIP 加速課題として採択された「秘匿計算による安全な組織間データ連携技術の社会実装」（課題番号 JPMJCR22U5）に取り組んでおり、銀行における不正送金業務への実運用に向け、更なる検知性能の向上やシステム実装に取り組んでいる。

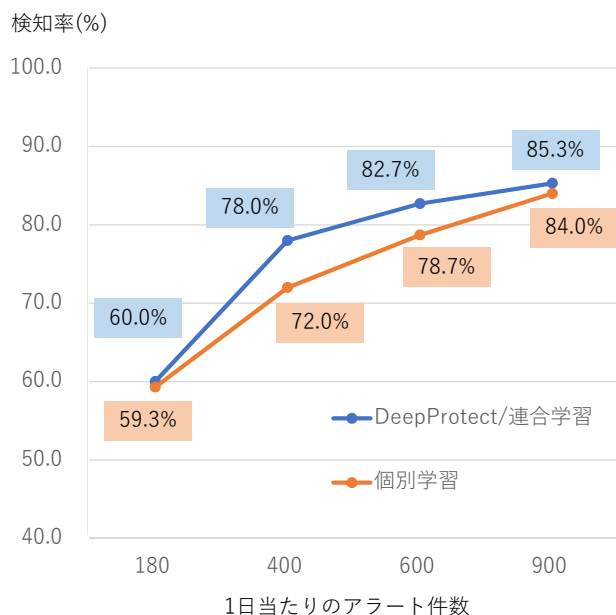


図5 被害検知グループの検知率（一例）

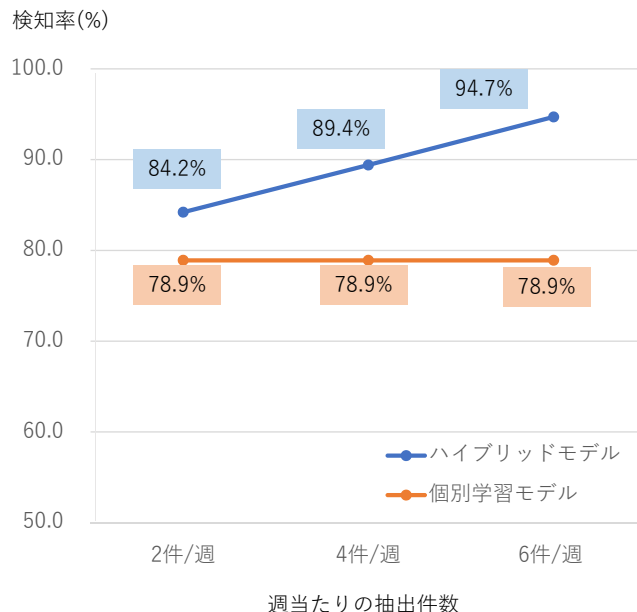


図6 加害検知グループの検知率（一例）

また、高度委託共同研究の下、神戸大学及びEAGLYS株式会社と連携し、大規模な顧客データを持つ4銀行の協力を受けた実証実験も行っている。上記実証実験により抽出された課題として、各銀行のデータフォーマットが統一されていない点が指摘されていた。各銀行のデータから解析できるデータを精査し、これらのデータフォーマットを統一し、不正利用検知の精度の向上を目指した実証実験となっている。

その他、銀行以外の企業からの業務委託契約により提供された取引データを用いた実証実験を実施した。機械学習に使用する共通特徴量を選定・統合、プライバシー保護連合学習技術DeepProtectを活用した連合学習モデルを作成し、個別学習モデルでは不正取引と判定できなかったデータが、連合学習で判定できた解析結果などを得ることができた。

5 むすび

DeepProtectは、データを暗号化したまま学習モデルの更新が行えるため、データのプライバシーを保護しつつ複数の組織による連合学習を実現することができる。これまでの、学習対象のデータが機微な情報を含むことから他組織と協力して機械学習を行うことが困難であった利用用途やサービスについてもDeepProtectを用いることにより、複数組織による連合学習が可能となり、データの利活用の促進とプライバシー保護を実現することができる。

今後は、金融分野のデータを用いた実証実験を実施してきた知見を活かし、更なる精度の追求及び利便性の向上を目指すとともに、他分野への適用拡大を目指している。

【参考文献】

- 1 Le Trieu Phong, Yoshinori Aono, Takuya Hayashi, Lihua Wang, and Shiho Moriai, "Privacy-Preserving Deep Learning: Revisited and Enhanced," ATIS 2017, pp.100–110.
- 2 Le Trieu Phong, Yoshinori Aono, Takuya Hayashi, Lihua Wang, and Shiho Moriai, "Privacy-Preserving Deep Learning via Additively Homomorphic Encryption," IEEE Trans. Inf. Forensics Secur. 13(5), pp.1333–1345, 2018.
- 3 Ligeng Zhu, Zhijian Liu, and Song Han, "Deep Leakage from Gradients," NeurIPS 2019, pp.14747–14756.
- 4 Le Trieu Phong and Tran Thi Phuong, "Privacy-Preserving Deep Learning via Weight Transmission," IEEE Trans. Inf. Forensics Secur. 14(11), pp.3003–3015, 2019.
- 5 Sachiko Kanamori, Taeko Abe, Takuma Ito, Keita Emura, Lihua Wang, Shuntaro Yamamoto, Le Trieu Phong, Kaiken Abe, Sangwook Kim, Ryo Nojima, Seiichi Ozawa, and Shiho Moriai, "Privacy-Preserving Federated Learning for Detecting Fraudulent Financial Transactions in Japanese Banks," J. Inf. Process. 30, pp.789–795, 2022.
- 6 Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam D. Smith, "Calibrating Noise to Sensitivity in Private Data Analysis," TCC 2006, pp.265–284.
- 7 Le Trieu Phong, Tran Thi Phuong, Lihua Wang, and Seiichi Ozawa, "Frameworks for Privacy-Preserving Federated Learning," IEICE Trans. Inf. Syst. 107(1), pp.2–12, 2024.
- 8 Tran Thi Phuong and Le Trieu Phong, "Distributed differentially-private learning with communication efficiency," J. Syst. Archit. 128, 102555, 2022.
- 9 Le Trieu Phong and Tran Thi Phuong, "Differentially private stochastic gradient descent via compression and memorization," J. Syst. Archit. 135, 102819, 2023.
- 10 Le Trieu Phong and Tran Thi Phuong, "Differentially-Private Distributed Machine Learning with Partial Worker Attendance: A Flexible and Efficient Approach," The 12th Conference on Information Technology and Its Applications, CITA 2023, Lecture Notes in Networks and Systems, vol.734. Springer, Cham. https://doi.org/10.1007/978-3-031-36886-8_2
- 11 Nicolas Papernot, Martín Abadi, Úlfar Erlingsson, Ian J. Goodfellow, and Kunal Talwar, "Semi-supervised Knowledge Transfer for Deep Learning from Private Training Data," ICLR 2017.
- 12 Prateek Jain, Vivek Kulkarni, Abhradeep Thakurta, and Oliver Williams, "To Drop or Not to Drop: Robustness, Consistency and Differential Privacy Properties of Dropout," CoRR abs/1503.02031, 2015.
- 13 Fuki Yamamoto, Seiichi Ozawa, and Lihua Wang, "eFL-Boost: Efficient Federated Learning for Gradient Boosting Decision Trees". IEEE Access

- 10, pp.43954–43963, 2022.
- 14 Fuki Yamamoto, Lihua Wang, and Seiichi Ozawa, “New Approaches to Federated XGBoost Learning for Privacy-Preserving Data Analysis”, ICONIP2020-27th International Conference on Neural Information Processing, LNCS 12533, pp.558–569, Springer, 2020.
- 15 Septiviana Savitri Asrori, Lihua Wang, and Seiichi Ozawa, “Permissioned Blockchain-based XGBoost for Multi Banks Fraud Detection,” ICONIP (3) 2022, vol.LNCS 13625, pp.683–692, Springer, 2023.
- 16 三浦 啓吾, 井上 広明, 金 相旭, 王 立華, 小澤 誠一, “動的サンプリングを使用した勾配ブースティング決定木の連合追加学習,” 第30回インテリジェント・システム・シンポジウム (FAN2022) 予稿集, pp.235–239, Sept. 2022.
- 17 北野 優斗, 王 立華, 小澤 誠一, “継続学習型連合学習モデルにおける効率的なリプレイデータの選択,” IEICE Technical Report, ISEC2023-95 (2024-03), pp.135–141, March 2024.
- 18 Ninghui Li, Wahbeh H. Qardaji, and Dong Su, “On sampling, anonymization, and differential privacy or, k-anonymization meets differential privacy,” AsiaCCS 2012: pp.32–33, 2012.
- 19 Ryo Nojima and Lihua Wang, “Differential Private (Random) Decision Tree without Adding Noise,” ICONIP (9) 2023, vol.CCIS 1963, pp.162–174, Springer, 2024.
- 20 Ryosuke Wakabayashi, Lihua Wang, Ryo Nojima, and Atsushi Waseda, “Security Evaluation of Decision Tree Meets Data Anonymization,” ICISSP 2024, pp.853–860, 2024.



野島 良 (のじま りょう)

サイバーセキュリティ研究所
セキュリティ基盤研究室
上席研究員
立命館大学 情報理工学部 教授
博士 (工学)
暗号、耐量子計算機暗号、秘密計算、機械学習



大久保 美也子 (おおくぼ みやこ)

サイバーセキュリティ研究所
セキュリティ基盤研究室
主任研究員
博士 (工学)
暗号理論、暗号プロトコル
【受賞歴】
2018 年 電子情報通信学会 SCIS 2018 イノベーション論文賞
2015 年 市村清新技術財団市村学術賞 功績賞
2000 年 電子情報通信学会 SCIS 2000 論文賞



Le Trieu Phong (れ ちゅう ふおん)

サイバーセキュリティ研究所
セキュリティ基盤研究室
主任研究員
博士 (学術)
暗号、プライバシー保護機械学習
【受賞歴】
2023 年 IEEE Signal Processing Society,
Best Paper Award 受賞



王 立華 (おう りつか)

サイバーセキュリティ研究所
セキュリティ基盤研究室
主任研究員
博士 (工学)
暗号、プライバシー保護機械学習
【受賞歴】
2023 年 IEEE Signal Processing Society,
Best Paper Award 受賞



金森 祥子 (かなもり さちこ)

サイバーセキュリティ研究所
セキュリティ基盤研究室
主任研究技術員
プライバシー、ユーザブルセキュリティ
【受賞歴】
2024 年 2024 年度山下記念研究賞
2023 年 CSS2023 優秀論文賞 (兼・UWS 優秀論文賞)



小野 元 (おの はじめ)

サイバーセキュリティ研究所
セキュリティ基盤研究室
研究員
博士 (統計科学)
プライバシー、統計、AI セキュリティ