



Title	Fast Secure Federated Learning against Semi-honest and Dishonest Adversaries
Author(s)	増田, 大輝
Citation	大阪大学, 2024, 博士論文
Version Type	VoR
URL	https://doi.org/10.18910/98692
rights	
Note	

The University of Osaka Institutional Knowledge Archive : OUKA

<https://ir.library.osaka-u.ac.jp/>

The University of Osaka

論文内容の要旨

氏 名 (増 田 大 輝)	
論文題名	Fast Secure Federated Learning against Semi-honest and Dishonest Adversaries (セミオネストおよびディスオネストな攻撃者に対する安全な連合学習の高速化)
論文内容の要旨 様々な分野で広く利用されている機械学習システムは、スマートフォンや Internet of Things (IoT) デバイスなど複数のエッジデバイス（ユーザー）から収集された教師データを活用することに関心が高まっている。複数のユーザーによって収集された教師データは膨大で多様であるため、高品質な機械学習サービスを提供する可能性がある。しかし、教師データのプライバシー漏洩がこれらのサービスの普及を妨げている。 この問題に対処するため、連合学習が注目されている。連合学習は、ユーザーの教師データを共有するのではなく、ユーザーがローカルで学習した機械学習モデルを共有することで教師データのプライバシーを保護する。連合学習では、各ユーザーはサーバーからモデルを受信して自身の教師データで学習し、学習したモデル（ローカルモデル）を返送する。サーバーはユーザーのローカルモデルを集約（加算）し、集約結果を使用してモデルを更新する。 連合学習の可能性は認識されているが、セミオネストな攻撃者とディスオネストな攻撃者に対して脆弱である。セミオネストな攻撃者とは、ローカルモデルからユーザーの教師データを推測するプライバシー攻撃を行うセミオネストなサーバーやセミオネストなユーザーである。ディスオネストな攻撃者とは、セミオネストな攻撃者に加え、悪意のある情報（悪意のあるローカルモデルなど）を送信してモデルの精度を低下させるセキュリティ攻撃（ビザンチン攻撃）を行うディスオネストなユーザー（ビザンチンユーザー）も含まれる。セミオネストな攻撃者に対する攻撃を軽減するために、ユーザーのローカルモデルを隠蔽しながら集約を行う secure aggregation と呼ばれる方法が提案されている。また、ディスオネストな攻撃者に対する攻撃を軽減するために、secure aggregationを実行しながら、悪意のある情報を検出して除去する Byzantine-resilient secure aggregation と呼ばれる方法が提案されている。しかし、secure aggregation と Byzantine-resilient secure aggregation は計算コストと通信コストの観点で効率が悪い。 本論文の目標は、secure aggregationとByzantine-resilient secure aggregationを行う連合学習のパフォーマンスを向上させることで連合学習の適用範囲を広げることである。既存の secure aggregation と Byzantine-resilient secure aggregation の分析を通じて、モデルサイズのデータ転送回数が通信のボトルネックであり、単一ノードでのモデルサイズのデータに関わる集中的な計算が計算のボトルネックであると特定した。モデルサイズのデータの転送回数を減らし、さらに単一ノードでのモデルサイズのデータに関わる集中的な計算をアルゴリズムレベルで最適化することでこれらのボトルネックを解消した。	

論文審査の結果の要旨及び担当者

氏 名 (増 田 大 輝)			
	(職)	氏 名	
論文審査担当者	主 査	准教授	小泉 佑揮
	副 査	教授	村田 正幸
	副 査	教授	渡辺 尚
	副 査	教授	山口 弘純
	副 査	教授	下西 英之

論文審査の結果の要旨

連合学習は、複数のユーザーが協調してモデルをトレーニングする手法である。各ユーザーがローカルに自身のデータでモデルをトレーニングし、トレーニングされたモデルをサーバーで集約することで、データそのものを共有せずに機械学習モデルを構築できる。連合学習には、セミオネストな攻撃者に共有されるモデルから学習データの情報が漏洩するプライバシー上の脅威に加えて、ディスオネストな攻撃者によってモデルの性能が劣化するセキュリティ上の脅威がある。

これらの脅威に対しては、セキュアアグリゲーションを用いる。これは、各ユーザーがトレーニングしたモデルをモデルと同サイズの乱数ベクトルを用いたOne-Time Padにより暗号化し、それを集約する手法である。集約後に自動的に乱数ベクトルが消え、モデルが復号される。ユーザーが処理中に離脱した場合に備え、その乱数モデルを消すための冗長ベクトルの共有が必要である。本博士論文では、この冗長ベクトルの処理がセキュアアグリゲーションの性能を決めることに注目し、これら脅威に対する安全性を保ちつつ処理の高速化を両立する方法を提案している。提案と貢献は以下の二点である。

博士論文の前半では、セミオネストな攻撃者のみを想定し、セキュアアグリゲーションの冗長化ベクトルの処理の高速化を達成している。まず、既存手法であるSecAggとLightSecAggは、冗長データの共有における計算と通信コストのバランスに課題があることを指摘している。SecAggは、冗長ベクトルの生成情報のみを通信で共有し、サーバーで冗長ベクトルを復元するため、通信コストが低い一方で、サーバにおける計算コストが高い。LightSecAggは、冗長ベクトル自体を共有するため、サーバーにおける冗長ベクトル復元をする必要がない一方で、通信コストが高い。筆者が設計したBalancedSecAggは、冗長ベクトルの生成情報のみを共有しつつ、冗長ベクトルの復元をユーザーが分散して処理することで、計算コストをユーザー間で分散し、処理時間を短縮した。

博士論文の後半では、セミオネストな攻撃者に加えて、ディスオネストな攻撃者が存在する場合の高速なセキュアアグリゲーション手法を設計している。ディスオネストな攻撃者が冗長ベクトルを改竄する攻撃に対しては、冗長ベクトルの検証情報も共に共有する手法が有効である。既存手法であるBREAは、冗長ベクトルと検証情報を共にユーザー間で共有しており、ユーザー間の通信コストが高いのに対して、本論文で提案したBREA-SVは、検証情報の処理をサーバーに委譲することで、通信コストを低下し、処理の時間を短縮している。

いずれの手法も、セキュリティ解析による安全性の検証に加えて、シミュレーションにより既存手法と比較して高速であることを示している。

本論文で取り扱う連合学習は、大量のデータをそのプライバシーを保護しながら集約するために有効な技術である。筆者は、連合学習において、プライバシー上の脅威とセキュリティ上の脅威に対する耐性を保ちつつ、セキュアアグリゲーションの高速化を達成しており、連合学習とセキュアアグリゲーションの実用化にむけた1つの解を提示している。EUのGeneral Data Protection Regulation (GDPR) などに代表されるユーザーのデータに関するプライバシー保護の機運により、データのサイロ化が進み大量の良質な学習データを確保することが難しい現状を鑑みると、本論文の成果の重要性が高まると予想される。以上のような理由から、本論文は博士（情報科学）の学位論文として価値のあるものと認める。