

基于生物特征识别的隐私保护可验证联邦学习

周 浩 戴 华 杨 庚 黄喻先 王周生

(南京邮电大学计算机学院、软件学院、网络空间安全学院 南京 210023)

摘 要 本文提出了一种基于生物特征的隐私保护可验证联邦学习方法(Biometric-based Privacy-Preserving Verifiable Federated Learning, BPPVFL),以解决联邦学习(Federated Learning, FL)中隐私保护和验证效率的双重挑战。传统的FL方法,通常不考虑隐私保护,使其易受数据泄露的威胁,而更安全的方法如基于零知识证明或同态哈希的FL方法,带来了显著的计算和通信开销。BPPVFL提出了基于生物特征的身份验证和数据完整性验证机制,实现了针对参与者敏感数据的隐私保护和高效的身份和数据验证,减少了客户端和服务端端的验证开销。此外,该方法使用针对生物特征数据的自适应噪声机制,在隐私保护和模型准确性之间取得了平衡。从理论和实验两方面证明了在BPPVFL中客户端的验证通信开销与客户端数量 N 和梯度维度 d 无关,从而实现了在大模型和高维度数据情况下的高效验证,即使在梯度维度 d 增加时,客户端的验证通信开销仍保持不变。在三个真实世界的生物特征数据集(SigD、BIDMC和TBME)上进行的大量实验表明,与隐私保护方法NbAFL相比BPPVFL的准确率最高提升了81%,同时与可验证方法VerifyNet相比BPPVFL的客户端验证通信开销最多减少了85%,服务器端通信开销与梯度维度独立且减少了90%以上。该方法在不同数据维度和隐私预算的环境下表现出出色的可扩展性和高效性。理论分析和实验结果表明,BPPVFL能够有效防止身份伪造和数据篡改,同时确保对敏感生物特征信息的强大隐私保护。BPPVFL为隐私保护联邦学习,尤其是生物特征应用,提供了一种有前景的解决方案。通过在隐私、验证开销和模型性能之间取得平衡,BPPVFL为安全联邦学习提供了一种实用且高效的方法。该工作作为未来针对敏感和高维数据环境的隐私保护机器学习方法研究奠定了基础。

关键词 联邦学习;隐私保护;可验证;边缘计算;生物特征识别

中图法分类号 TP309 **DOI号** 10.11897/SP.J.1016.2025.01848

Privacy-Preserving Verifiable Federated Learning Based on Biometric Recognition

ZHOU Hao DAI Hua YANG Geng HUANG Yu-Xian WANG Zhou-Sheng

(School of Computer Science, Nanjing University of Posts and Telecommunications, Nanjing 210023)

Abstract This paper proposes a novel framework called Biometric-based Privacy-Preserving Verifiable Federated Learning (BPPVFL), which addresses two critical and long-standing challenges in Federated Learning (FL): ensuring strong privacy protection for participants and achieving efficient, scalable verification of clients and data. As FL has become an increasingly important paradigm for distributed machine learning, particularly in privacy-sensitive domains such as healthcare, finance, and biometric authentication, its security and efficiency requirements have grown more demanding. Despite its decentralized nature, conventional FL often fails to provide adequate safeguards for user data, making it vulnerable to privacy breaches, model

收稿日期:2024-09-19;在线发布日期:2025-05-21。本课题得到国家自然科学基金(No. 62372244)、南京邮电大学引进人才科研启动基金项目(No. NY224058)和医院管理方向重点实验室开放课题(No. 2024LYKC003)资助。周 浩,博士,副教授,中国计算机学会(CCF)会员,主要研究领域为联邦学习、隐私保护、可验证和差分隐私。Email:haozhou@njupt.edu.cn。戴 华(通信作者),博士,教授,中国计算机学会(CCF)会员,主要研究领域为数据管理与安全、数据库安全、行为识别。E-mail:daihua@njupt.edu.cn。杨 庚,博士,教授,中国计算机学会(CCF)高级会员,主要研究领域为计算机网络与通信、信息与网络安全、分布与并行计算、无线传感器网络与安全。黄喻先,博士研究生,主要研究领域为隐私保护、拜占庭攻击和联邦学习。王周生,博士,主要研究领域为联邦学习、生物验证、边缘计算。

poisoning, and unauthorized participation. To overcome these limitations, BPPVFL introduces a privacy-preserving mechanism that integrates biometric-based identity authentication with verifiable model update mechanisms. This dual strategy ensures that only authenticated participants verified using immutable biometric features can take part in the collaborative training process, and that any data sent by clients can be efficiently verified for integrity. In contrast to cryptographic methods such as zero-knowledge proofs or homomorphic hashing, which significantly increase computational and communication complexity, BPPVFL is designed to maintain lightweight overhead while ensuring robust security. This makes it a practical and scalable solution for real-world FL deployments. A key innovation of BPPVFL is its use of an adaptive noise injection mechanism specifically designed for biometric data. This technique introduces calibrated noise during the training process to preserve privacy, without compromising the integrity or utility of the data. Unlike conventional differential privacy methods that apply fixed noise levels, the adaptive approach in BPPVFL adjusts noise dynamically based on the characteristics of the input, ensuring a better trade-off between privacy preservation and model performance. As a result, BPPVFL achieves high model accuracy while still meeting strong privacy guarantees for sensitive biometric information. From a theoretical standpoint, BPPVFL offers a significant advantage: the client-side verification communication overhead is provably independent of both the number of clients N and the gradient dimensionality d . This decoupling is particularly important for modern machine learning models, which often involve millions of parameters. While traditional verification methods become increasingly inefficient as model complexity grows, BPPVFL maintains constant communication costs on the client side, ensuring scalability even in high-dimensional learning settings. To validate the effectiveness and efficiency of the proposed framework, extensive experiments were conducted on three real-world biometric datasets: SigD, BIDMC, and TBME. The results show that, compared to the state-of-the-art privacy-preserving FL approach NbAFL, BPPVFL achieves up to 81% improvement in model accuracy. Furthermore, in comparison to VerifyNet, a leading verifiable FL method, BPPVFL reduces client-side verification communication overhead by as much as 85%. It also reduces server-side communication overhead by over 90%, while maintaining independence from the gradient dimensionality. Additionally, BPPVFL demonstrates strong adaptability and scalability across a wide range of conditions, including varying privacy budgets and data dimensions. It effectively prevents malicious behaviors such as identity spoofing and model update tampering, while maintaining high performance in distributed learning environments. In conclusion, BPPVFL provides a practical, secure, and efficient solution for privacy-preserving and verifiable federated learning, particularly in applications involving biometric data. By achieving a balanced trade-off between privacy, communication overhead, and learning performance, BPPVFL sets a new benchmark for secure FL systems. This work lays a solid foundation for future research into privacy-aware machine learning techniques tailored for sensitive, high-dimensional data environments, and contributes to the broader goal of building trustworthy AI systems in distributed settings.

Keywords federated learning; privacy protection; verifiable; edge computing; biometric identification

1 引言

随着生成式人工智能的迅猛发展,各领域的模型日益依赖于个人数据输入,以实现更高的性能和

更个性化的用户体验^[1]。尤其是在物联网(Internet of Things, IoT)设备的广泛应用背景下,如智能穿戴设备和传感器,个人数据的持续收集和使用为模型的训练和更新提供了丰富的资源^[2]。然而,这种趋势也引发了严重的安全和隐私问题。敏感生物特

征数据(如指纹、虹膜、语音、心电图、光电容积图等)在各种应用场景中因其独特性和不可更改性而特别容易成为攻击目标。在身份验证系统、医疗健康监测、移动支付、智慧城市等领域,这些生物特征数据通常需要在设备与服务器之间频繁传输和共享,从而暴露在潜在的攻击风险之中。攻击者在这些场景中利用中间人攻击,通过截获和篡改通信数据,冒充合法参与者或者操纵系统的行为。例如,在智能穿戴设备的健康监测应用中,攻击者可以通过拦截并篡改传输的心电图或光电容积图数据,伪造用户的健康状况或身份,进一步获取敏感的健康信息或执行未经授权的操作。在移动支付和金融交易领域,指纹或面部识别等生物特征的泄露可能导致账户被盗用,严重危害用户的财产安全。即便是在智慧城市中,虹膜或步态识别等生物特征的广泛应用也使得个人身份信息暴露在更大的风险之下,特别是当这些信息被不法分子用于跟踪、监视或定位特定目标时,潜在的威胁更为严重。对手可以通过篡改发送给服务器的梯度信息,或者操纵发送给客户端的聚合模型结果,以推测或重构客户端的隐私信息,从而实施更加隐蔽和复杂的攻击手段^[3]。因此,如何在数据的收集、存储和传输过程中,有效地认证参与者身份并确保交互过程的安全性,防止攻击者获取和操纵个人隐私信息,成为当前亟待解决的重大问题^[4]。当前亟需开发出能够提供强有力生物特征隐私保护的解决方案,不仅要在身份认证阶段保护用户数据不被恶意篡改和伪造,还需要在数据传输和聚合阶段提供足够的安全性,以抵御日益复杂的攻击形式。

目前的研究主要集中在利用生物特征(如指纹、虹膜、语音等)进行身份验证,以应对客户端身份认证的挑战^[5]。这些基于生物特征的身份验证方法在移动环境中展现了一定的有效性,通过提取和匹配用户的生物特征来识别合法参与者。然而,这些方法在应对日益复杂的攻击手段方面存在显著局限。现有的生物特征认证方案大多专注于验证用户的身份,而忽略了数据在传输和处理过程中的完整性和真实性,这使得系统难以防止攻击者通过截获和篡改数据进行伪造和重放攻击^[6]。一旦攻击者成功劫持了通信链路,便可以利用篡改后的生物特征数据进行未经授权的操作,甚至通过重放合法用户的旧数据来蒙混过关,威胁系统的安全性。此外,现有的生物特征身份验证方法通常缺乏对生物特征隐私的有效保护,尤其是在数据的收集、存储和共享过程

中。一方面,传统的生物特征识别方法如指纹和虹膜识别通常依赖于集中式的数据存储和管理,这为潜在的攻击者提供了集中的攻击目标。数据一旦被攻破,所有存储的生物特征信息可能被大规模泄露,从而带来严重的安全和隐私风险。另一方面,在边缘计算和分布式学习环境下,虽然数据不直接上传至中央服务器,但在参与者设备和服务器之间频繁交换的生物特征数据依然存在被窃取和滥用的风险。尤其是在复杂环境下,诸如基于心电图和光电容积图的生物特征认证机制同样面临数据完整性和隐私保护的不足^[7-8]。这些机制在应对复杂攻击场景时,难以有效保护生物特征数据免遭重构和分析攻击。

在此背景下,加强生物特征隐私保护和数据完整性验证的重要性愈发凸显。现有研究表明,通过集成多种技术,如差分隐私、同态加密、安全多方计算等,能够在一定程度上保护数据隐私和安全性。然而,这些方法往往依赖于复杂的计算和高昂的通信开销,对于资源受限的边缘设备或需要实时响应的应用场景而言,难以提供切实可行的解决方案。此外,传统的基于密钥的认证和完整性验证机制在面对密钥被复制、篡改和伪造的风险时,难以确保安全性,尤其是在攻击者具备强大计算能力和多种攻击手段的情况下^[9-16]。因此,当前的生物特征认证方法面临着重大的挑战和不足,急需引入新的机制来增强生物特征数据的隐私保护和数据完整性验证,以更好地应对多样化的攻击和安全威胁。这些新的方法应能够在保证系统安全性的同时,兼顾计算和通信开销,特别是在联邦学习等新兴分布式学习框架中,实现更高效和安全的生物特征认证和数据保护。

联邦学习作为一种新兴的机器学习范式,致力于在客户端和服务端之间进行协作训练的过程中保护用户数据隐私,从而避免服务器直接接触到敏感的原始用户数据^[17-18]。与传统的集中式学习方法不同,联邦学习将模型训练任务分布到各个客户端设备上,在本地处理数据,并通过共享模型更新而不是原始数据来完成训练过程,这一特点使得联邦学习在数据隐私保护方面具有天然的优势^[19]。然而,尽管联邦学习的框架设计初衷是为了保护用户的隐私数据,其在身份验证和数据完整性验证方面依然面临重大挑战^[9-10]。现有的基于密钥的身份验证和数据完整性验证机制,通常依赖于客户端和服务端之间的共享密钥来确保通信的安全性和数据的完整

性。然而,这些机制存在密钥容易被复制、篡改和伪造的风险。一旦攻击者成功获取或伪造密钥,就能够冒充合法参与者进行恶意操作,从而危害系统的安全性和数据的可信度^[11-12]。此外,密钥管理本身也是一个复杂的过程,尤其是在分布式环境中,如何安全、有效地分发和更新密钥是一个难以解决的问题^[13-14]。这种基于密钥的验证方法在联邦学习的背景下尤为不适用,因为联邦学习通常涉及大量的边缘设备,这些设备的计算资源和存储能力有限,难以负担高开销的加密和解密操作。过高的通信和计算开销不仅增加了系统的负担,还可能降低设备的响应速度,影响用户体验^[16]。另一方面,尽管一些研究尝试利用安全多方计算来保护联邦学习过程中的数据隐私,这些方法同样面临攻击者针对中央服务器的脆弱性问题^[15]。在安全多方计算的框架下,多个参与者可以在不泄露各自数据的前提下协同计算某个函数的值,这种方法能够在一定程度上提升数据的隐私性。然而,一旦中央服务器被攻破或受到恶意篡改,这种保护机制便不复存在。攻击者可以通过操纵参与者提交的模型更新或者篡改聚合结果,直接影响联邦学习的结果,甚至反推出参与者的私密数据。因此,这种依赖于中央服务器的架构在面对强大的攻击者时依然存在较大的安全隐患^[20-22]。更为关键的是,联邦学习中对参与者身份的认证和对传输数据的完整性验证缺乏更加灵活和鲁棒的解决方案^[23-24]。目前的研究主要集中在如何提升联邦学习中数据的隐私保护水平,但对身份验证的关注相对较少。攻击者可以通过伪造身份,参与到联邦学习的训练过程中,甚至可能通过提交恶意的模型更新来影响全局模型的性能,或者获取其他合法参与者的敏感信息。这种模型篡改攻击和推理攻击对系统的安全性和数据隐私构成了严重威胁。此外,由于联邦学习框架中数据更新频繁,参与者设备和中央服务器之间的交互复杂且通信量大,现有的验证机制难以在不增加额外计算和通信开销的前提下提供足够的安全性和验证能力。因此,联邦学习亟需一种高效且安全的身份验证和数据完整性验证机制,这种机制应能够适应分布式学习的特点,充分考虑边缘设备的资源限制,同时提供对抗各种攻击的能力。为了在不增加系统开销的情况下有效验证参与者的身份并确保数据的完整性,必须设计出新的解决方案,例如基于生物特征的轻量级认证和验证方法。这些方法可以利用用户设备本身的生物特征信息,结合联邦学习的隐私保护优势,提供一种更加安全且高效的认证机制,抵御中间人攻击、模型篡改攻击以及其他恶意攻击。这样一来,联邦学习不仅能在保护数据隐私方面继续发挥其优势,还能进一步增强系统的安全性

和鲁棒性,为各种实际应用场景中的生物特征识别和验证提供更可靠的解决方案。

本文针对联邦学习在生物特征识别场景中的隐私保护与安全验证问题,提出了一种基于生物特征隐私保护的验证联邦学习框架(BPPVFL)。本文的主要工作包括:(1)设计了一种自适应噪声机制,通过动态调整噪声强度与类型,平衡隐私保护与模型性能,并有效保护用户数据隐私;(2)提出了签名验证与聚合结果完整性验证机制,结合用户私钥与梯度签名,防止恶意用户的身份伪造与数据篡改,增强联邦学习框架的安全性和鲁棒性;(3)设计了基于生物特征的数据完整性验证机制,确保聚合梯度的安全与完整性,有效识别并隔离恶意参与方。论文的主要贡献包括:

(1)提出一种BPPVFL方法,能够有效提供对生物特征的隐私保护,防止对手通过重构生物特征数据进行攻击。

(2)设计一种基于参与者生物特征的身份验证机制,防止对手通过冒名顶替进行攻击,与经典的隐私保护方案NbAFL^[11]相比,该机制对于参与者身份的验证准确率最高提高了81%。

(3)提出基于生物特征的数据完整性验证机制,证明了客户端的验证通信开销与客户端数量 N 和梯度维度 d 无关,从而实现了在大模型和高维度数据情况下的高效验证,即使在梯度维度 d 增加时,客户端的验证通信开销仍保持不变。与经典的^[25]相比,BPPVFL的验证通信开销减少了90%。

2 相关工作

本节主要关注FL中基于生物特征的身份验证和差分隐私保护两方面面临的挑战并分析本文提出BPPVFL的研究动机。

2.1 基于生物特征的身份验证

生物特征认证系统,如指纹、虹膜扫描和语音识别,由于其难以复制的特性,被广泛应用于各种安全敏感场景中,提供了一道强有力的防线以抵御未经授权的访问^[26-27]。尽管这些系统因其独特的生物特征难以被复制而受到青睐,但它们在隐私保护方面依然存在严重的漏洞和挑战。具体而言,这些系统的核心问题在于生物特征数据的唯一性和不可更改性。一旦面部或指纹等生物特征数据被泄露或被黑客窃取,这些数据就会永久失效,难以通过传统的密码重置或更换方法来修复^[28]。例如,面部和指纹识别系统仍然

容易遭受诸如呈现攻击等的破坏。在这种攻击中,攻击者使用受害者面部的照片或指纹模具等复制品来欺骗认证系统,从而实现未经授权的访问^[29]。更为严重的是,一旦这些生物特征数据被泄露或被不法分子获取,用户无法像更改密码那样简单地更改其生物特征数据。这种不可更改性进一步放大了生物特征数据的隐私风险和安全威胁。针对这些挑战,研究人员开始探索生理信号如光体积描记(Physical Property Gradient, PPG)作为更安全的替代方案,因为这些信号难以被外界观察或复制^[30]。PPG 信号在与传统生物特征认证方式结合时,可用于活体检测,提供额外的一层安全保障^[31]。此外,最新的研究表明,PPG 信号本身也可以作为一种唯一标识符用于个体认证,展现出其作为一种安全且可靠的生物特征验证方法的巨大潜力^[32-33]。PPG 通过可穿戴传感器测量血容量变化,其持续性和实时性使得其更难以被复制和欺诈性使用,从而增强了认证过程的安全性^[34]。

然而,尽管 PPG 等新兴生物特征信号在隐私保护方面展现了巨大的潜力,目前大多数生物特征解决方案仍依赖于集中式的数据集和训练方法。这种集中化的方法不仅增加了数据泄露的风险,还使得用户的生理信号可能会被不法分子或第三方滥用。这种对中心化数据集的依赖,使得用户的生物特征数据在传输和存储过程中面临被拦截、滥用或篡改的风险,进一步放大了隐私和安全方面的威胁。为了应对这些问题,研究人员正在探索将生物特征认证技术与联邦学习相结合,以进一步增强用户隐私的保护^[35-36]。联邦学习通过避免将原始数据直接上传到中心服务器,提供了一定程度的隐私保护。然而,这种方法也并非万无一失。在开放的联邦学习环境中,由于存在大量的异质参与者,这些系统依然容易遭受身份伪造攻击和中间人攻击。此外,攻击者还可以通过分析模型更新的信息,反向推断出用户的生物特征数据,从而造成严重的隐私泄露。尽管联邦学习在保护数据隐私方面展现了前景,但针对生物特征数据的专门隐私保护方法依然匮乏,需要进一步的研究与开发。

综上所述,当前生物特征认证系统在隐私保护方面的不足是一个亟待解决的问题。在现有方法的基础上,发展更为先进的、针对生物特征的隐私保护技术,确保这些不可更改且高度敏感的数据能够在安全可靠的环境下使用,是未来研究的重要方向。

2.2 基于生物特征的身份验证

在基于联邦学习的差分隐私保护领域,已有一些研究对模型参数在上传和广播阶段的全局敏感度

进行了深入分析,并引入差分隐私噪声以保护梯度信息,证明了模型的收敛性^[37]。此外,还有研究提出了基于混合差分隐私和自适应压缩的边缘计算联邦学习框架,用于处理工业数据,验证了其在防止推理攻击方面的有效性^[38]。另一类研究指出梯度变化对训练数据的敏感性是衡量信息泄漏风险的关键指标,并基于此提出了一种通过扰动梯度来匹配信息泄漏风险的防御方法,实现了高效的隐私保护^[39]。还有研究提出了一种针对模型反演攻击和成员推理攻击的差分隐私防御方法,通过对置信度得分向量的修改和标准化来模糊原始数据,同时确保分类精度不受损^[40]。

这些研究尽管在差分隐私保护方面做出了贡献,但在应对生物特征数据的隐私保护时仍存在一定的局限性。生物特征数据由于其独特性和高度敏感性,通常不符合独立同分布假设,特别是当参与者仅拥有单一类别的数据时。这类数据极易受到差分隐私保护机制的影响。例如,尽管一些研究通过隐私预算分配方法来实现差分隐私保护,将隐私预算划分为两部分 $\epsilon = \epsilon_l + \epsilon_r$,用于主题词分布的随机抽样和对特定主题中词频加入拉普拉斯噪声^[41],但在面对生物特征数据时,这种方法可能无法有效防止信息泄露,因为隐私预算分配并未考虑数据的非独立同分布特性。同样,有些研究通过对抗学习和隐私预算分配来将差分隐私保护嵌入特定层和学习过程中,以提升未标记目标数据的分类精度^[42]。然而,当应用于生物特征数据时,这些方法往往缺乏足够的精细度和适应性。均匀分配隐私预算的策略未能充分考虑参与者之间数据分布的差异,特别是在单一类别数据的情况下,可能导致隐私保护效果不足。

综上所述,现有研究^[37-38,42]虽然对联邦学习中的差分隐私保护方法做出了贡献,但在针对生物特征数据的隐私保护方案中,特别是在参与者数据非独立同分布且仅包含单类数据的情况下,仍然存在局限性。目前大多数研究中均采用将隐私预算均匀分配到每轮参数上的方式,这种策略未能有效防止攻击者通过多次窃取模型参数来推断数据间的关联性。特别是对于生物特征数据,由于其独特的敏感性和非独立同分布特性,现有的隐私预算分配策略可能引入过量噪声,影响模型性能。因此,迫切需要针对生物特征数据和参与者仅有单一类别数据的情况,开发更为精细化和有效的隐私预算分配方法,以增强隐私保护效果。

3 预备知识

本节介绍了联邦学习、差分隐私、数字签名和同态哈希等隐私保护和可验证的机制和方法。

3.1 联邦学习

在我们设计的联邦学习系统中,一个应用服务器和 N 个客户端共同工作,每个客户端 $U_{i:1 \leq i \leq N}$ 拥有数据集 $D_{i:1 \leq i \leq N}$ 。FL系统的主要目标是在保护客户端隐私的同时,尽可能地最小化人工智能模型的损失函数。我们将每个在线边缘客户端表示为 $U_i \in \{U_1, U_2, \dots, U_M\}$ 。

我们的目标是在保护客户端交互的生物特征隐私的同时,找到一个全局最小损失函数 GF 。因此,我们的FL方案通过联邦学习减少 U_i 的局部损失函数 F_i ,并在服务器上聚合所有上传的参数来近似实现这一目标。假设 U_i 的权重为 p_i ,优化问题可以表示为^[25]

$$p = \arg \min_p \sum_{i=1}^M p_i w_i \quad (1)$$

$$p_i = \arg \min_p w_i F_i(p, D_i) \quad (2)$$

在这个交互过程中,客户端通过差分隐私、同态加密和安全多方计算保护其私有信息。最终,服务器获得全局的FL模型。

3.2 差分隐私

在当今的AI时代,越来越多的人使用云服务器来支持更高效的机器学习服务^[43]。然而,将客户端数据存储和训练在云服务器上可能会引发隐私问题,特别是对于高度隐私的数据,如生物特征信息。差分隐私(Differential Privacy, DP)是一种隐私保护框架,可以为分布式数据收集方案提供强大的隐私保护标准^[44]。

在 (ϵ, δ) -DP中, $\epsilon > 0$ 表示数据库中相邻数据集 D_i 和 D'_i 的所有输出的可区分边界, δ 表示在添加隐私保护机制后,对于两个相邻数据集 D_i 和 D'_i ,无法被 e^ϵ 限定的概率。对于任何 δ ,较大的 ϵ 的隐私保护方案对于相邻数据集有更显著的可区分性,因此具有更高的隐私泄露风险。DP的定义如下:

定义 1. $((\epsilon, \delta)$ -DP)^[44]. 一个随机机制 $\mathcal{M}$: $\mathcal{X} \rightarrow \mathcal{R}$,其定义域为 $\mathcal{X}$,值域为 $\mathcal{R}$,如果对于所有可测集 $\mathcal{S} \rightarrow \mathcal{R}$ 和任何两个相邻数据库 D_i 和 $D'_i \in \mathcal{X}$,

$$Pr[\mathcal{M}(D_i) \in \mathcal{S}] \leq e^\epsilon Pr[\mathcal{M}(D'_i) \in \mathcal{S}] + \delta \quad (3)$$

则满足 (ϵ, δ) -DP。

$\oplus$ 表示对应哈希值上的操作^[45]。

文献[44]中定义的高斯机制可以应用于FL框架中,以确保 (ϵ, δ) -DP值。差分隐私的一个重要概念是顺序可组合性^[9]。对于一系列随机算法 $f_1, f_2, f_3, \dots, f_n$,如果 $\forall f_{i:1 \leq i \leq n}$ 满足 ϵ_i -DP,则整个过程满足 $\sum_{i=1}^n \epsilon_i$ -DP。DP的另一个重要概念是敏感度^[9]。具体来说,对于任意函数 f ,其敏感度 ∇f 可以定义如下:

定义 2. 敏感度^[44]. 对于任意函数 f ,以及 $\forall x, x' \in \text{Domain}(f)$,敏感度 $\Delta_2^U f$ 定义为

$$\Delta_2^U f = \max_{x, x'} \|f(x) - f(x')\|_2 \quad (4)$$

其中 x 和 x' 是某个客户端的向量, $\|\cdot\|_2$ 表示向量的 $\mathcal{L}_2$ 范数^[44]。

定义 3. 高斯机制^[44]. 对于 $\forall \epsilon \in (0, 1)$ 和 $c^2 > 2 \ln(1.25/\sigma)$,具有参数 $\sigma \leq c \Delta_2^U f$ 的高斯机制满足 (ϵ, σ) -DP。

在DP机制中,噪声水平与FL方案的收敛速度之间的关系仍然是一个重要的研究问题。

3.3 数字签名

数字签名(Digital Signature, DS)使用公钥基础设施建立了一种可验证的机制,将边缘客户端与其发送的消息链接起来。以下是DS的算法概述:

$(sk, pk) \leftarrow \text{DS.Ini}(\lambda)$: 该算法接受安全参数 λ 作为输入,并生成密钥 (sk, pk) ^[25]。 $z \leftarrow \text{DS.Sign}(sk, x)$: 该算法接受一个密钥 sk 和模型参数 x 作为输入,输出一个签名 z 。 $0, 1 \leftarrow \text{DS.Verify}(pk, x, z)$: 该算法接受一个公钥 pk ,模型参数 x 和签名 z 作为输入,并返回表示失败或成功的验证结果0或1。

3.4 同态哈希

同态哈希(Homomorphic Hashing, HH)是一种密码学技术,旨在确保在数据传输和处理过程中数据的完整性和安全性。同态哈希函数的定义如下:

定义 4. 一个哈希函数 $\text{Hash}(\cdot)$ 被称为同态哈希,如果对任意两个消息 m_1 和 m_2 及其对应的哈希值 $\text{Hash}(m_1)$ 和 $\text{Hash}(m_2)$,存在一个运算符 $\oplus$,使得:

$$\text{Hash}(m_1 \cdot m_2) = \text{Hash}(m_1) \oplus \text{Hash}(m_2) \quad (5)$$

其中 $\cdot$ 表示消息上的某种操作(如加法、乘法等),

这种特性允许在不暴露原始数据的情况下验证

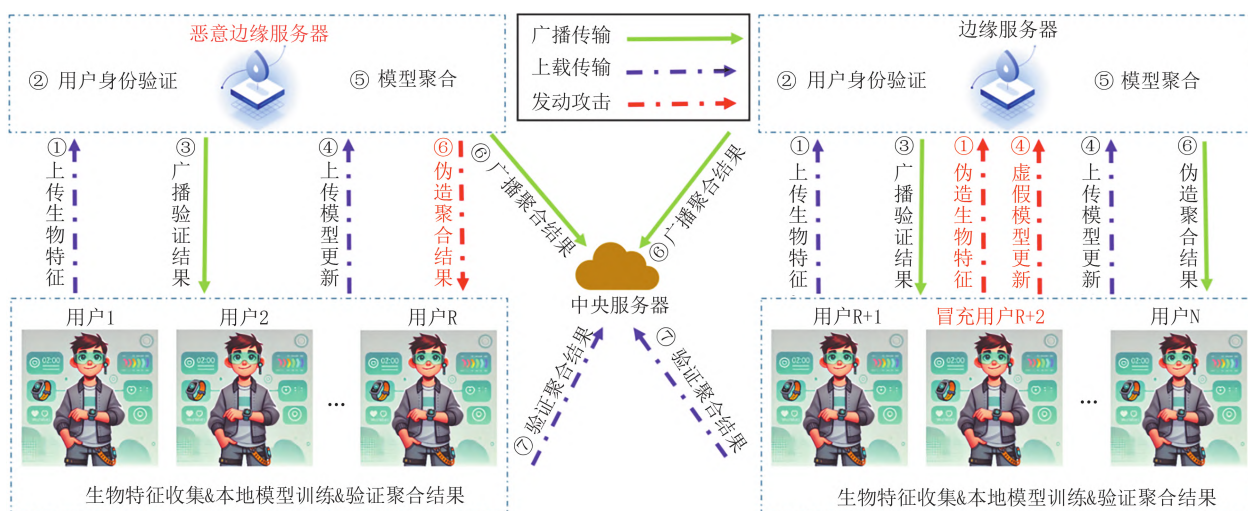


图1 BPPVFL框架

数据的完整性。在联邦学习的场景中,同态哈希可以用于验证客户端上传的梯度是否未被篡改,同时不需要了解具体的梯度值。这在确保系统的安全性和可信性方面具有重要意义。

4 问题定义和威胁模型

4.1 问题定义

我们的问题目标是在实现基于生物特征的身份识别的同时,保护参与者的生物特征隐私,并确保聚合结果的完整性不受篡改。在基于生物特征验证的联邦学习系统中,存在以下三个主要问题:

(1)基于生物特征的身份识别:参与者使用可穿戴生物传感器(例如 Apple Watch、Google Pixel Watch 和 Fitbit Charge)或其他智能设备,提取个人的生物特征数据。为了确保每个参与者的身份认证过程可靠,系统需要设计一种基于生物特征的身份验证机制,能够准确地识别合法用户,并防止恶意用户伪造身份参与联邦学习过程。

(2)参与者生物特征的隐私保护:生物特征数据(例如心率、指纹、面部特征等)属于高度敏感的个人隐私信息。为了保护参与者的隐私,系统需要设计有效的隐私保护机制,确保参与者的生物特征在身份认证和模型训练过程中不被泄露。具体而言,这要求在不泄露原始生物特征数据的情况下,完成身份认证和模型训练的所有必要操作。

(3)聚合结果的完整性保护:在联邦学习过程中,用户在本地设备上训练模型并生成模型梯度更新,这些更新被发送到边缘服务器进行聚合。为了

确保整个系统的安全性和可靠性,必须保证这些聚合结果的完整性不受篡改。任何试图在传输过程中篡改、伪造或操控模型更新的行为都应被检测和防止。

基于上述问题定义,本研究旨在设计一种新的联邦学习方法 BPPVFL,实现对参与者身份的生物特征验证,同时保护参与者生物特征隐私,并确保模型聚合结果的完整性,BPPVFL 框架如图 1 所示。通过使用差分隐私、同态哈希和数字签名等技术,我们的方法能够在保证系统安全性的前提下,提供有效的隐私保护和完整性验证。

4.2 威胁模型

在本方法中,定义合法用户为通过初始注册阶段并经过生物特征认证模型 M_{bio} 验证成功的参与者。合法用户拥有经系统注册并验证成功的合法生物特征数据 Bio_i 。相应地,合法生物特征 Bio_i 是指在初始注册时系统认证成功的、用户真实且唯一的生物特征模板,如指纹、虹膜或心率模式等。与之对应,未授权用户(非法用户)则被定义为未经过注册认证或认证失败的参与者。这些未授权用户无法获得对应的合法私钥 $sk_i = \text{Hash}(Bio_i)$,因此在理论上无法生成与合法用户生物特征数据相匹配的有效签名。

在实现基于生物特征的身份识别、保护参与者生物特征隐私以及保证聚合结果完整性这三个目标的过程中,我们需要应对多种潜在的威胁。本节定义了涉及的威胁模型,包括来自中央服务器、边缘服务器、用户和恶意对手的威胁。

在基于生物特征的联邦学习系统中,身份伪造

是一个主要威胁。恶意对手可能会伪造合法用户的身份,利用虚假生物特征数据 $\widehat{Bi}o_u^{fake}$ 参与联邦学习过程。为确保系统的安全性,系统需要验证每个参与者的身份是否是真实的。假设用户 R 提供的生物特征数据为 $\widehat{Bi}o_u$, 如果存在伪造用户 R' , 满足:

$$\text{Hash}(\widehat{Bi}o_u) = \text{Hash}(\widehat{Bi}o_u^{fake}) \quad (6)$$

则身份验证机制会被攻破。因此,需要设计一个稳健的生物特征哈希和验证算法,使得恶意对手无法生成有效的伪造生物特征数据 $\widehat{Bi}o_u^{fake}$ 。

中央服务器和边缘服务器在参与聚合和验证过程中,可能对用户的生物特征数据和模型梯度数据感兴趣,试图推断用户的隐私信息。为防止隐私泄露,我们使用差分隐私机制来保护用户的生物特征数据 $\widehat{Bi}o_u$ 。对于每个用户 R 和其对应的生物特征数据 $\widehat{Bi}o_u$, 需要满足 (ϵ, δ) -DP 的约束条件:

$$\Pr[\mathcal{M}(\widehat{Bi}o_u) \in \mathcal{S}] \leq e^\epsilon \Pr[\mathcal{M}(\widehat{Bi}o'_u) \in \mathcal{S}] + \delta \quad (7)$$

其中 $\mathcal{M}$ 表示用于保护隐私的随机机制, $\widehat{Bi}o_u$ 和 $\widehat{Bi}o'_u$ 是任意两个相邻数据集。 ϵ 越小,隐私保护越强; δ 是一个很小的概率,用来表示隐私泄露的风险。

在联邦学习过程中,边缘服务器收集用户上传的模型梯度更新 w_i , 并将其聚合为 $\tilde{w}$ 后发送至中央服务器。恶意对手可能会尝试篡改这些聚合结果,影响最终的模型性能。为了保证聚合结果的完整性,系统必须能够验证从边缘服务器到中央服务器的聚合结果的准确性。假设聚合结果为 G_{global} , 通过使用数字签名和同态哈希,我们需要确保:

$$\text{Verify}(pk, \tilde{w}, z) = 1 \quad (8)$$

其中, Verify 是验证算法, pk 是公钥, $\tilde{w}$ 是聚合结果, z 是数字签名。如果恶意对手试图篡改结果 $\tilde{w}$, 则验证将失败, $\text{Verify}(pk, \tilde{w}, z) \neq 1$ 。

非法用户可能会向模型提交经过精心设计的虚假数据,以使联邦学习过程输出偏离预期。为了防止这种威胁,本方法要求在聚合过程之前对所有用户上传的梯度进行验证。对于每个上传的梯度 g_u , 需要满足以下验证条件:

$$\text{Sign}(sk, g_u) \neq \text{Sign}(sk, g_u^{fake}) \quad (9)$$

其中, g_u 是真实用户的梯度更新, g_u^{fake} 是伪造用户上传的虚假梯度, Sign 是签名算法, sk 是私钥。如果签名验证失败,则可以检测到伪造数据。

综上所述,本威胁模型覆盖了联邦学习过程中

身份伪造、生物特征隐私泄露、聚合结果篡改等多种潜在威胁。我们的方法通过生物特征验证、差分隐私保护、同态哈希和数字签名等技术手段来应对这些威胁,确保系统的安全性和完整性。

5 BPPVFL 方法

5.1 BPPVFL 方法的核心思路

在基于生物特征的隐私保护可验证联邦学习方法(BPPVFL)中,我们针对威胁模型中的四类潜在威胁(身份伪造、隐私泄露、聚合结果篡改、恶意数据投毒),提出了三项核心策略,以确保系统的安全性、隐私性和完整性。

传统密码学方案在身份验证过程中可能会受到身份伪造攻击的威胁,尤其是在分布式环境下,攻击者可能伪造合法用户的身份来参与联邦学习。为解决这一问题,BPPVFL 利用生物特征的难以复制性和抗伪造性来加强身份验证。生物特征数据(如指纹、虹膜、心率模式等)具有唯一性和不可伪造性,因此使用生物特征进行身份认证,可以显著减少伪造用户的攻击风险。

在 BPPVFL 方法中,每个用户 U_i 使用其可穿戴设备(如智能手表、智能手机等)提取其生物特征数据 Bio_i , 并通过生物特征哈希函数 $\text{Hash}(Bio_i)$ 将其转换为唯一的生物特征标识符。这些标识符用于在联邦学习过程中的身份验证,只有具有合法生物特征的用户才能被认证为合法参与者:

$$sk_i = \text{Hash}(Bio_i) \quad (10)$$

其中, sk_i 是用户 U_i 的数字签名密钥。通过这种方式,可以有效抵御身份伪造攻击。

在联邦学习过程中,生物特征数据的隐私保护至关重要。由于生物特征数据的敏感性,任何数据泄露都会导致严重的隐私问题。因此,BPPVFL 通过差分隐私机制为生物特征数据加入自适应噪声,以实现对用户隐私的保护。

具体而言,每个客户端在上传生物特征数据之前,BPPVFL 首先根据不同类型的生物特征数据的敏感程度和隐私需求动态计算隐私预算 (ϵ, δ) 。例如,对于指纹或虹膜等高度敏感的生物特征,系统自动选择较低的隐私预算(即较小的 ϵ 和 δ),以提高隐私保护水平,而对于心率等敏感度较低的生物特征数据,则适当提高隐私预算以保障模型准确性。确定隐私预算后,根据差分隐私机制中的自适应方法,

动态调整高斯噪声的标准差 σ :

$$\sigma = f(\epsilon, \delta, S_{Bio}) \quad (11)$$

其中, $f(\cdot)$ 是自适应函数, 用于综合考虑隐私预算 ϵ , δ 和生物特征数据的敏感度 S_{Bio} (例如数据方差、信息熵或特征唯一性程度等), 以动态确定合适的噪声强度。敏感度较高的特征数据会导致较大的 σ , 即更多的噪声加入, 确保用户隐私得到充分保护, 而敏感度较低的数据则选择较小的噪声, 以维持模型性能。最终得到扰动后的数据:

$$\widehat{Bio}_i = Bio_i + \beta_i, \quad \beta_i \sim N(0, \sigma^2) \quad (12)$$

其中, β_i 是服从高斯分布的噪声项, σ 是噪声的标准差, 确保添加的噪声量满足 (ϵ, δ) -差分隐私条件。通过这种方式, 任何潜在的攻击者即使获取了扰动后的生物特征数据 $\widehat{Bio}_i$, 也无法准确重构出用户的原始生物特征数据 Bio_i , 从而保护了用户的隐私。通过自适应差分隐私机制, BPPVFL 方法实现了针对不同生物特征数据的自适应差分隐私保护, 既有效保护了用户的隐私安全, 又保证了模型训练的准确性。

在联邦学习的聚合过程中, 恶意对手可能会尝试篡改模型参数的聚合结果, 或通过恶意数据投毒来操控模型性能。为了解决这些问题, BPPVFL 方法通过结合数字签名和同态哈希技术来确保聚合结果的完整性和抗投毒攻击能力。每个用户在计算本地模型梯度 g_i 后, 使用其生物特征哈希生成的私钥 sk_i 对其进行数字签名, 以确保上传的梯度确实来自合法用户:

$$z_i = \text{Sign}(sk_i, \widehat{g}_i) \quad (13)$$

其中, $\widehat{g}_i = g_i + \gamma_i$, $\gamma_i \sim \mathcal{N}(0, \sigma_g^2)$ 。服务器在接收到梯度更新 g_i 及其签名 z_i 后, 可以使用相应的公钥 pk_i 验证梯度的合法性。如果验证失败, 则拒绝该梯度更新, 从而有效抵御恶意用户的攻击。为了保证聚合结果的完整性, 方法采用同态哈希技术对聚合操作进行加密验证。具体而言, 同态哈希函数 $H(\cdot)$ 满足:

$$H(g_i + g_j) = H(g_i) \oplus H(g_j) \quad (14)$$

在聚合过程中, 边缘服务器和每个合法参与者分别独立地计算聚合梯度的同态哈希值, 以实现多方验证, 防止单一来源的哈希值被攻击者伪造。具体而言, 每个参与者首先计算本地梯度 $\widehat{g}_i$ 的同态哈希值:

$$H(\widehat{g}_i) = \text{HomHash}(\widehat{g}_i) \quad (15)$$

边缘服务器在接收到所有参与者上传的梯度后, 使用同态哈希函数的同态特性进行梯度聚合并计算聚合哈希值:

$$H(\tilde{w}) = \bigoplus_{\widehat{g}_i \in G} H(\widehat{g}_i) \quad (16)$$

同时, 每个参与者本地也独立计算本地视角的聚合哈希值 (即参与者视角的期望聚合值):

$$H(\tilde{w}^V) = \bigoplus_{U_i \in G} H(\widehat{g}_i) \quad (17)$$

然后, 所有参与者独立地将本地计算得聚合哈希值 $H(\tilde{w}^V)$ 直接上传至中央服务器。中央服务器最终收到边缘服务器上传的聚合梯度 $\tilde{w}$ 及其对应的哈希值 $H(\tilde{w})$, 同时也收到各个参与者上传的独立哈希值 $H(\tilde{w}^V)$ 。中央服务器根据以下规则验证聚合结果的完整性:

$$\text{VerifyA}\left(pk, \tilde{w}, \{H(\tilde{w}^V)_i\}_{i=1}^n, H(\tilde{w})\right) = \begin{cases} \text{True, 如果 } H(\tilde{w}) = H(\tilde{w}^V)_i, \forall i \\ \text{False, 其他} \end{cases}$$

该多方哈希值对比策略确保任何攻击者即使伪造梯度更新和对应的哈希值, 也必须同时成功伪造所有参与者独立计算得到的聚合哈希值, 这是实际中极难实现的。攻击者成功的概率近似为

$$\Pr\left[\text{VerifyA}\left(pk, \tilde{w}^{fake}, \{H(\tilde{w}^V)_i^{fake}\}_{i=1}^n, H(\tilde{w})^{fake}\right) = \text{True}\right] \approx 0 \quad (18)$$

通过以上机制, 系统能够有效防止攻击者同时篡改聚合结果和哈希值, 而不被检测到。因此, BPPVFL 方法通过多方独立计算与验证机制, 确保了聚合结果的完整性和安全性。

5.2 BPPVFL 方法的整体架构

基于生物特征的隐私保护可验证联邦学习方法 (BPPVFL) 旨在通过一种安全、高效的架构, 实现基于生物特征的身份验证、生物特征隐私保护和联邦学习模型的完整性保证。该架构主要包括三个关键组成部分: 参与者、边缘服务器和中央服务器, 各个组成部分通过安全的交互过程协同工作, 保证整个系统的安全性和隐私性。

每个参与者 (用户) 使用可穿戴生物传感器 (如智能手表或其他智能设备) 来收集生物特征数据, 并在本地设备上执行联邦学习任务。具体而言, 参与者在本地设备上进行以下操作: 参与者首先提取个人的

生物特征数据 Bio_i (例如心率、步态、指纹等), 并对其进行处理, 以消除噪声和不相关的特征。然后, 将预处理后的生物特征数据 $\widehat{Bio}_i$ 通过差分隐私机制添加自适应噪声, 得到保护隐私的生物特征数据:

$$\widehat{Bio}_i = Bio_i + \beta_i, \quad \beta_i \sim N(0, \sigma^2) \quad (19)$$

参与者使用本地数据集 D_i 进行本地模型训练, 计算本地模型梯度 g_i :

$$g_i = \nabla L_{local}(M_i, D_i) \quad (20)$$

其中, M_i 是参与者 i 的本地模型, L_{local} 是本地损失函数。参与者生成用于验证身份的生物特征哈希值, 并使用该哈希值生成的私钥对本地梯度 g_i 进行加噪后得到 $\widehat{g}_i$, 并对其进行签名, 同时计算该梯度的同态哈希值 $H(\widehat{g}_i)$:

$$z_i = \text{Sign}(sk_i, \widehat{g}_i), \quad H(\widehat{g}_i) = \text{HomHash}(\widehat{g}_i) \quad (21)$$

其中 $\widehat{g}_i = g_i + \gamma_i$, $\gamma_i \sim \mathcal{N}(0, \sigma_g^2)$ 。

参与者将生物特征标识符 $\widehat{Bio}_i$ 、本地模型梯度 $\widehat{g}_i$ 、数字签名 z_i 和同态哈希值 $H(\widehat{g}_i)$ 一并发送到最近的边缘服务器。边缘服务器作为高性能计算设备 (如无人机、小型基站等), 在参与者与中央服务器之间充当中间节点, 负责验证参与者身份并聚合本地模型更新。边缘服务器的操作流程如下:

边缘服务器接收参与者发送的扰动后的生物特征标识符 $\widehat{Bio}_i$ 后, 利用预设的生物特征验证模型 M_{bio}^* 进行身份验证。具体而言, 模型 M_{bio} 是在系统初始注册阶段利用合法用户的真实生物特征数据 (如指纹、虹膜或心率模式) 进行训练而构建的二分类器, 能够区分合法用户和非法用户的生物特征数据。身份验证过程如下:

首先, M_{bio} 接收上传的扰动后的生物特征数据 $\widehat{Bio}_i$, 输出该生物特征与合法用户注册模板的匹配相似度得分 $S(\widehat{Bio}_i)$:

$$S(\widehat{Bio}_i) = M_{bio}(\widehat{Bio}_i)$$

然后, 将该相似度得分与预先定义的认证阈值 τ 进行比较, 其中阈值 τ 是通过交叉验证方法确定的, 目的是平衡假接受率 (FAR) 和假拒绝率 (FRR), 确保认证系统的鲁棒性与安全性:

$$\text{VerifyI}(M_{bio}, \widehat{Bio}_i) = \begin{cases} \text{True}, & \text{若 } S(\widehat{Bio}_i) \geq \tau \\ \text{False}, & \text{若 } S(\widehat{Bio}_i) < \tau \end{cases}$$

若参与者的生物特征数据 $\widehat{Bio}_i$ 与注册模板的相似度超过阈值 τ , 则该参与者通过身份验证, 被视为

合法参与者, 能够继续参与后续的联邦学习过程; 否则, 系统将拒绝该参与者的请求。通过以上明确的机制定义, 模型 M_{bio} 实现了可靠且明确的生物特征身份验证, 确保只有合法用户能够参与联邦学习任务, 从而增强了系统对身份伪造攻击的鲁棒性。

边缘服务器对所有合法的本地梯度 $\widehat{g}_i$ 进行聚合, 计算得到聚合梯度 $\tilde{w}$, 并将结果发送到中央服务器:

$$\tilde{w} = \frac{1}{|G|} \sum_{g_i \in G} \widehat{g}_i \quad (22)$$

其中, G 表示所有通过验证的参与者的梯度集合。

中央服务器在整个架构中起到管理和协调的作用, 负责启动和监督联邦学习任务, 以及进行最终的全局模型聚合。中央服务器的操作流程如下: 中央服务器从多个边缘服务器接收聚合梯度 $\tilde{w}$, 并对每个聚合结果进行验证, 以确保其完整性和正确性。验证方法包括检查数字签名和同态哈希值, 确认聚合过程是否没有受到篡改。

$$\text{VerifyA}(pk, \tilde{w}, z) \rightarrow \text{True/False} \quad (23)$$

其中, pk 是由各个参与方的局部公共密钥 pk_i 通过安全多方计算得到 $pk = f(pk_1, pk_2, \dots, pk_n)$, $f(\cdot)$ 表示一种预定义的安全聚合函数, 以确保所有参与方的公共密钥贡献到全局密钥中, z 是由各个参与方的局部参数 z_i 根据安全聚合协议进行聚合得到 $z = g(z_1, z_2, \dots, z_n)$, g 为哈希函数或其他安全聚合操作, 确保中央服务器能够验证全局梯度的完整性和一致性。在聚合过程中, 为了保证 pk 和 z 的安全性与可靠性, 中央服务器不会直接访问或存储各个 pk_i 和 z_i 的原始值, 而是通过同态加密或秘密共享机制完成聚合过程。这样, 参与方的局部信息在聚合过程中仍保持私密性, 中央服务器只能得到全局聚合结果 pk 和 z 。如果所有聚合结果均通过验证, 中央服务器将这些结果进一步聚合, 生成全局模型参数 w , 并将其广播给所有边缘服务器和参与者。

$$w = \frac{1}{|E|} \sum_{\tilde{w}_E \in E} \tilde{w}_E \quad (24)$$

其中, E 表示所有边缘服务器的集合, $\tilde{w}_E$ 是边缘服务器 E 的聚合结果。中央服务器将聚合后的全局模型参数 w 发送给所有边缘服务器, 边缘服务器再将其广播给参与者。参与者在本地接收并验证这些参数, 并根据验证结果决定是否更新本地模型。

整个架构设计确保了各个组件之间的安全交互与有效通信。参与者通过边缘服务器与中央服务器

进行交互,确保身份验证和数据传输的安全性;中央服务器通过边缘服务器与参与者保持通信,发布全局模型结果,并接收反馈。在每一轮联邦学习的迭代过程中,系统通过生物特征验证、差分隐私保护、数字签名和同态哈希等技术手段,有效抵御威胁模型中的各种攻击,保障联邦学习过程的安全性和可信性。BPPVFL 的整体架构设计为分布式环境下的联邦学习提供了一个安全的、可验证的和隐私保护的解决方案。

算法 1. BPPVFL

输入:参与者本地数据集 D_i ,生物特征数据 Bio_i ;生物特征验证模型 M_{bio} ;差分隐私参数 ϵ, δ 和噪声方差 σ^2 ;数字签名密钥对 (sk_i, pk_i) 。

输出:联邦学习全局模型参数 w 。

1. 初始化验证模型 M_{bio} 和联邦模型 M_{global} 。
2. FOR 每个参与者 U_i
3. 提取并预处理生物特征数据 Bio_i 。
4. 添加自适应噪声 $g_i = \nabla L_{local}(M_i, D_i)$ 。
5. 计算哈希值 $Hash(Bio_i)$,生成签名密钥 sk_i 。
6. 计算本地梯度 $\hat{g}_i = \nabla L_{local}(M_i, D_i)$ 。
7. 生成扰动梯度 $\hat{g}_i = g_i + \gamma_i, \gamma_i \sim \mathcal{N}(0, \sigma_g^2)$ 。
8. 使用 sk_i 对 $\hat{g}_i$ 签名, $z_i = Sign(sk_i, \hat{g}_i)$ 。
9. 计算同态哈希值 $H(\hat{g}_i) = HomHash(\hat{g}_i)$ 。
10. 将 $\widehat{Bio_i}, \hat{g}_i, z_i, H(\hat{g}_i)$ 发送到边缘服务器。
11. ENDFOR
12. FOR 每个边缘服务器 E_j
13. FOR 每个接收到的参与者 U_i 的消息
14. 使用验证模型 M_{bio} 对 $\widehat{Bio_i}$ 进行身份验证。
15. IF 身份验证成功
16. 使用 pk_i 验证签名 $VerifyS(pk_i, \hat{g}_i, z_i)$ 。
17. IF 签名验证成功
18. 将 $\hat{g}_i$ 纳入聚合集合 G 。
19. ENDFOR
20. ENDFOR
21. ENDFOR
22. 计算聚合梯度 $\tilde{w} = \frac{1}{|G|} \sum_{g_i \in G} \hat{g}_i$ 。
23. 计算 $\hat{g}_i$ 的同态哈希值 $H(\tilde{w}) = \bigoplus_{H(\hat{g}_i) \in G} H(\hat{g}_i)$ 。
24. 将 $\tilde{w}$ 和 $H(\tilde{w})$ 发送到中央服务器。
25. ENDFOR
26. 接收聚合梯度 $\tilde{w}_E$ 和同态哈希值 $H(\tilde{w}_E)$ 。
27. FOR 每个边缘服务器的结果 $(\tilde{w}_E, H(\tilde{w}_E))$
28. 使用 $VerifyA(pk, \tilde{w}_E, H(\tilde{w}_E))$ 验证 $\tilde{w}_E$ 。
29. IF 验证成功

30. 将 $\tilde{w}_E$ 纳入全局聚合集合 W 。
31. ENDFOR
32. ENDFOR
33. 计算全局模型参数 $w = \frac{1}{|W|} \sum_{\tilde{w}_E \in W} \tilde{w}_E$ 。
34. 将全局模型参数 w 发布至所有边缘服务器。
35. FOR 每个参与者 U_i
36. 通过边缘服务器接收全局模型参数 w 。
37. 验证全局模型参数的完整性。
38. IF 验证成功
39. 使用 w 更新本地模型。
40. ENDFOR
41. ENDFOR

5.3 BPPVFL 方法的描述

本节详细描述基于生物特征的隐私保护可验证联邦学习方法(BPPVFL)的具体方法如算法 1 所示,包括身份验证、生物特征隐私保护、聚合结果的完整性验证以及抵御恶意数据投毒攻击的能力。

BPPVFL 的首要步骤是对参与者进行基于生物特征的身份验证。具体流程如下:参与者 U_i 使用其本地设备(如智能手表、智能手机)提取生物特征数据 Bio_i ,然后利用生物特征哈希函数计算其哈希值,生成数字签名密钥 sk_i :

$$sk_i = Hash(Bio_i) \quad (25)$$

为了保护生物特征数据的隐私,参与者在上传生物特征数据之前,通过添加自适应高斯噪声 β_i 来进行扰动处理,得到扰动后的生物特征数据 $\widehat{Bio_i}$:

$$\widehat{Bio_i} = Bio_i + \beta_i, \quad \beta_i \sim \mathcal{N}(0, \sigma^2) \quad (26)$$

其中, β_i 是一个符合高斯分布的噪声项,确保扰动后的数据满足 (ϵ, δ) -差分隐私要求。参与者将扰动后的生物特征数据 $\widehat{Bio_i}$ 发送至边缘服务器,边缘服务器使用预设的生物特征验证模型 M_{bio} 验证参与者的身份:

$$VerifyI(M_{bio}, \widehat{Bio_i}) = \begin{cases} \text{True, 验证成功} \\ \text{False, 验证失败} \end{cases} \quad (27)$$

如果身份验证成功,参与者被认为是合法用户,能够继续参与联邦学习过程;否则,其请求将被拒绝。

在身份验证成功后,参与者需要在本地进行模型训练,同时保护生物特征和模型梯度的隐私。每个参与者在本地数据集 D_i 上训练模型 M_i ,计算损失函数 $L_{local}(M_i, D_i)$ 的梯度 g_i :

$$g_i = \nabla L_{local}(M_i, D_i) \quad (28)$$

为进一步保护隐私,参与者对梯度 g_i 添加噪声

γ_i , 然后使用数字签名算法对扰动后的梯度进行签名, 生成数字签名 z_i :

$$\hat{g}_i = g_i + \gamma_i, \quad \gamma_i \sim N(0, \sigma_g^2) \quad (29)$$

$$z_i = \text{Sign}(sk_i, \hat{g}_i) \quad (30)$$

其中, γ_i 是服从高斯分布的噪声项, sk_i 是由生物特征哈希值生成的私钥。

参与者还计算梯度的同态哈希值 $H(\hat{g}_i)$, 确保在聚合过程中可以验证其完整性:

$$H(\hat{g}_i) = \text{HomHash}(\hat{g}_i) \quad (31)$$

参与者将扰动后的生物特征数据 $\widehat{Bio}_i$ 、扰动后的梯度 $\hat{g}_i$ 、数字签名 z_i 和同态哈希值 $H(\hat{g}_i)$ 一并发送到边缘服务器, 并广播给其他参与者其数字签名 z_i 和同态哈希值 $H(\hat{g}_i)$ 。

边缘服务器接收到来自多个参与者的上传信息后, 执行以下操作: 边缘服务器首先使用验证模型 M_{bio} 对每个参与者的生物特征数据 $\widehat{Bio}_i$ 进行身份验证, 然后使用参与者的公钥 pk_i 验证上传梯度 $\hat{g}_i$ 的数字签名 z_i , 确保数据的来源合法且未被篡改:

$$\text{VerifyI}(M_{bio}, \widehat{Bio}_i) \rightarrow \text{True/False} \quad (32)$$

$$\text{VerifyS}(pk_i, \hat{g}_i, z_i) \rightarrow \text{True/False} \quad (33)$$

对于通过身份验证和签名验证的参与者, 边缘服务器将其梯度 $\hat{g}_i$ 纳入聚合集合中并广播合法参与者集合 Φ , 使用加权平均的方法计算聚合梯度 $\tilde{w}$, 并生成其同态哈希值 $H(\tilde{w})$:

$$\tilde{w} = \frac{1}{|G|} \sum_{g_i \in G} \hat{g}_i, \quad H(\tilde{w}) = \bigoplus_{H(\hat{g}_i) \in G} H(\hat{g}_i) \quad (34)$$

边缘服务器将聚合梯度 $\tilde{w}$ 及其同态哈希值 $H(\tilde{w})$ 发送至中央服务器。同时, 合法参与者用同样方式本地计算聚合结果的同态哈希值 $H(\tilde{w}^v) = \bigoplus_{U_i \in \Phi} H(\hat{g}_i)$ 发送至中央服务器。

中央服务器接收到来自多个边缘服务器的聚合结果后和边缘客户端的验证信息后, 执行以下操作: 中央服务器使用各边缘服务器和边缘客户端提供的同态哈希值 $H(\tilde{w})$ 、 $H(\tilde{w}^v)$ 对聚合结果进行验证, 确保传输过程中未发生篡改。如果验证通过, 则继续下一步操作:

$$\text{VerifyA}(pk, \tilde{w}, H(\tilde{w}^v)) \rightarrow \text{True/False} \quad (35)$$

中央服务器对来自所有边缘服务器的聚合结果进行进一步的全局聚合, 生成全局模型参数 w :

$$w = \frac{1}{|E|} \sum_{\tilde{w}_E \in E} \tilde{w}_E \quad (36)$$

中央服务器将全局模型参数 w 发送至各边缘服务器, 再由边缘服务器转发给参与者。参与者在接收到全局模型后, 通过验证其签名和完整性来确认模型的合法性, 然后更新本地模型。

BPPVFL 方法通过生物特征的抗伪造性、差分隐私的自适应噪声保护、数字签名和同态哈希的完整性验证, 实现了一个安全、隐私保护和可验证的联邦学习系统, 有效应对了身份伪造、隐私泄露、聚合结果篡改和恶意数据投毒等威胁。

6 BPPVFL 方法的鲁棒性分析

本节对 BPPVFL 在抗身份伪造、模型篡改和投毒攻击三个方面进行了安全性分析和证明, 并证明了 BPPVFL 具有维度独立和参与客户端数量独立的验证开销。

6.1 BPPVFL 方法的安全性分析

首先, 我们证明基于生物特征的隐私保护可验证联邦学习方法 (BPPVFL) 能够实现基于生物特征的身份识别并具备抗伪造能力。具体而言, 我们将分析方法在抵御身份伪造攻击方面的安全性, 证明攻击者无法伪造合法用户的生物特征数据通过身份验证。

定理 1. 在 BPPVFL 方法中, 服务器可以实现基于生物特征的身份识别, 并抵御身份伪造攻击。

证明. 在 BPPVFL 方法中, 参与者的身份是通过其生物特征数据 Bio_i 进行验证的。生物特征数据具有唯一性和难以复制性, 适用于身份认证。以下为 BPPVFL 方法中身份识别与抗伪造的安全性证明:

对于每个合法用户 U_i , 其生物特征数据为 Bio_i , 扰动后的生物特征数据为 $\widehat{Bio}_i = Bio_i + \beta_i$, 其中 $\beta_i \sim \mathcal{N}(0, \sigma^2)$ 是用于隐私保护的高斯噪声。每个合法用户 U_i 使用其生物特征数据 Bio_i 计算哈希值 $\text{Hash}(Bio_i)$, 生成数字签名密钥 $sk_i = \text{Hash}(Bio_i)$ 。这个密钥用来对扰动后的梯度 $\hat{g}_i$ 进行数字签名:

$$z_i = \text{Sign}(sk_i, \hat{g}_i)$$

边缘服务器使用预设的生物特征验证模型 M_{bio} 对每个上传的扰动生物特征数据 $\widehat{Bio}_i$ 进行验证。身份验证过程由以下判断条件确定:

$$\text{VerifyI}(M_{bio}, \widehat{Bio_i}) = \text{True}$$

要证明 BPPVFL 能够抗击身份伪造攻击, 我们需要分析以下几点: 生物特征数据(如指纹、虹膜)具有高度的唯一性, 任何伪造者都难以在不获取原始生物特征的情况下生成与合法用户相同的生物特征数据。因此, 伪造者无法找到一个伪造的生物特征数据 Bio_i^{fake} 使得 $\text{Hash}(Bio_i) = \text{Hash}(Bio_i^{fake})$ 。由于哈希函数 $\text{Hash}(\cdot)$ 是抗碰撞哈希函数, 满足对于任意两个不同的生物特征数据 Bio_i 和 Bio_i^{fake} , 有

$$\Pr[\text{Hash}(Bio_i) = \text{Hash}(Bio_i^{fake})] \approx 0.$$

参与者在上传其生物特征数据之前, 会根据系统设定的隐私参数自适应地添加噪声。由于每次添加的噪声 β_i 都是随机的, 服从正态分布 $\mathcal{N}(0, \sigma^2)$, 因此伪造者难以预测到合法用户的扰动生物特征 $\widehat{Bio_i}$ 。即使伪造者获得了扰动后的生物特征数据 $\widehat{Bio_i}$, 由于噪声的随机性和不可预测性, 伪造者也难以重构出原始的生物特征数据 Bio_i 。

设攻击者 A 试图通过伪造生物特征数据 $\widehat{Bio_i^{fake}}$ 来冒充用户 U_i , 以便通过身份验证。由于哈希函数的抗碰撞性和扰动数据的不可预测性, 伪造者 A 的成功概率可以表示为

$$\Pr[\text{VerifyI}(M_{bio}, \widehat{Bio_i^{fake}}) = \text{True}] \leq \Pr[\text{Hash}(Bio_i) = \text{Hash}(Bio_i^{fake})] + \Pr[\text{Reconst}(Bio_i | \widehat{Bio_i})].$$

其中, $\Pr[\text{Reconst}(Bio_i | \widehat{Bio_i})]$ 表示通过 $\widehat{Bio_i}$ 重建 Bio_i 的概率。由于 $\Pr[\text{Hash}(Bio_i) = \text{Hash}(Bio_i^{fake})] \approx 0$ 和 $\Pr[\text{Reconst}(Bio_i | \widehat{Bio_i})] \approx 0$, 因此:

$$\Pr[\text{VerifyI}(M_{bio}, \widehat{Bio_i^{fake}}) = \text{True}] \approx 0.$$

由于伪造者无法在现实中找到两个不同的生物特征数据使得它们的哈希值相同, 且难以重构合法用户的原始生物特征数据, BPPVFL 方法能够有效实现基于生物特征的身份识别并抗击伪造攻击。因此, 定理得证。

证毕。

定理 2. 在 BPPVFL 方法中, 中央服务器可以通过多方独立哈希值验证机制有效抵抗聚合结果篡改攻击, 并保证聚合结果的完整性。

证明. 在 BPPVFL 方法中, 梯度聚合的完整性通过同态哈希与多方验证策略保障。每个参与者

首先使用私钥生成的签名保证了梯度上传的合法性; 随后, 每个参与者和边缘服务器分别独立计算聚合梯度的哈希值, 并将各自独立计算的聚合哈希值发送给中央服务器进行交叉验证。具体而言, 边缘服务器计算的聚合哈希值为

$$H(\tilde{w}) = \bigoplus_{\widehat{g_i} \in G} H(\widehat{g_i}).$$

每个参与者本地也独立计算聚合哈希值为

$$H(\tilde{w}^V) = \bigoplus_{U_i \in G} H(\widehat{g_i}).$$

中央服务器进行验证时, 需满足

$$H(\tilde{w}) = H(\tilde{w}^V), \quad \forall i \in G.$$

若攻击者篡改聚合结果 $\tilde{w}^{fake}$, 则为了通过验证, 其需要同时成功篡改所有独立计算的哈希值 $H(\tilde{w}^V)_i$ 。然而, 由于这些哈希值是由不同参与者和边缘服务器独立计算并直接上传至中央服务器, 攻击者不可能同时成功篡改所有独立源头, 故篡改聚合结果成功的概率几乎为零:

$$\Pr[\text{VerifyA}(pk, \tilde{w}^{fake}, \{H(\tilde{w}^V)_i\}_{i=1}^n) = \text{True}] \approx 0.$$

因此, 定理 2 得证。

证毕。

定理 3. 在 BPPVFL 方法中, 服务器通过基于生物特征生成的公钥 pk_i 对参与者提交的梯度签名 z_i 和同态哈希值 $H(\widehat{g_i})$ 进行验证, 从而有效确保提交数据的来源合法性。任何未授权的攻击者(未通过生物特征认证的用户)无法成功生成有效签名, 从而无法提交有效的梯度更新。

证明. 在 BPPVFL 中, 参与者的身份认证和数字签名机制均依赖于合法用户的生物特征数据。具体而言, 系统在初始阶段对每个参与者的生物特征数据 Bio_i 进行验证, 并以此为基础生成唯一的私钥 sk_i 和对应的公钥 pk_i :

$$sk_i = \text{Hash}(Bio_i), \quad pk_i = \text{PublicKeyGen}(sk_i)$$

合法用户在每次梯度上传前, 使用其合法私钥 sk_i 对梯度 $\widehat{g_i}$ 进行数字签名 $z_i = \text{Sign}(sk_i, \widehat{g_i})$ 。服务器端使用对应的公钥 pk_i 验证签名是否有效:

$$\text{VerifyS}(pk_i, \widehat{g_i}, z_i) = \text{True}.$$

然而, 对于未授权用户(非法攻击者), 其不具备经过系统认证的合法生物特征数据, 因此无法获得相应的合法私钥 sk_i 。即使攻击者试图伪造生物特征数据 Bio_i^{fake} 以生成私钥 sk_i^{fake} , 也极难满足哈希碰撞条件:

$$\text{Hash}(Bio_i^{fake}) \neq \text{Hash}(Bio_i).$$

因此,攻击者无法成功伪造出与合法用户完全相同的私钥,也无法生成与合法公钥 pk_i 对应的有效签名 z_i^{fake} 。攻击者生成签名被服务器验证成功的概率极低:

$$\Pr[\text{VerifyS}(pk_i, g_i^{fake}, z_i^{fake}) = \text{True}] \approx 0.$$

未授权用户无法成功通过基于生物特征的签名验证过程。因此,上式的验证机制成立,能有效识别并拒绝未授权用户提交的非法梯度更新。

同时,系统通过进一步使用同态哈希技术对聚合结果进行验证,即使攻击者通过其他手段获得合法签名,也无法在不被检测的情况下篡改聚合结果,从而进一步提高了系统对非法数据投毒攻击的鲁棒性。

证毕。

接下来我们证明基于生物特征的隐私保护可验证联邦学习方法(BPPVFL)能够实现对参与者生物特征的隐私保护。具体来说,我们证明BPPVFL方法通过差分隐私机制,可以有效保护参与者的生物特征数据不被泄露。

定理4. 在BPPVFL方法中,参与者的生物特征数据满足 (ϵ, δ) -差分隐私。

证明. 在BPPVFL方法中,为了实现对参与者生物特征数据的隐私保护,每个参与者在上传其生物特征数据之前,会采用差分隐私机制来添加自适应噪声。以下是证明BPPVFL方法能够实现生物特征隐私保护的具体过程:

对于每个参与者 U_i ,其原始生物特征数据为 Bio_i 。为了保护其隐私,参与者在上传生物特征数据之前,通过添加自适应高斯噪声 β_i 来对其进行扰动处理,得到扰动后的生物特征数据 $\widehat{Bio}_i$:

$$\widehat{Bio}_i = Bio_i + \beta_i, \quad \beta_i \sim \mathcal{N}(0, \sigma^2).$$

其中, β_i 是服从高斯分布 $\mathcal{N}(0, \sigma^2)$ 的噪声项。噪声的方差 σ^2 是根据差分隐私参数 ϵ 和 δ 自适应调整的。

差分隐私(Differential Privacy, DP)是一种强有力的隐私保护机制,旨在保证对任何单个输入的输出分布的影响有限。在BPPVFL方法中,生物特征数据的扰动满足 (ϵ, δ) -差分隐私的定义:

对于任意两个相邻的生物特征数据 Bio_i 和 Bio'_i 以及所有可能的输出集 $\mathcal{S}$,扰动机制 $\mathcal{M}$ 满足:

$$\Pr[\mathcal{M}(Bio_i) \in \mathcal{S}] \leq e^\epsilon \cdot \Pr[\mathcal{M}(Bio'_i) \in \mathcal{S}] + \delta.$$

其中, $\epsilon > 0$ 表示隐私预算, δ 表示隐私泄露的概率。

在BPPVFL中使用的高斯机制(Gaussian Mechanism)已被证明能够确保 (ϵ, δ) -差分隐私。当加入的高斯噪声的标准差 σ 满足以下条件时:

$$\sigma \geq \frac{\Delta_2 f \cdot \sqrt{2 \ln(1.25/\delta)}}{\epsilon}.$$

其中, $\Delta_2 f$ 是函数 f 的 $\mathcal{L}_2$ 敏感度,表示两个相邻输入之间的最大变化范围。因此,对于扰动后的生物特征数据 $\widehat{Bio}_i = Bio_i + \beta_i$,只要 β_i 的标准差 σ 满足上述条件,该机制即满足 (ϵ, δ) -差分隐私要求。假设攻击者A试图通过获取扰动后的生物特征数据 $\widehat{Bio}_i$ 来恢复原始生物特征数据 Bio_i 。由于高斯噪声的随机性和正态分布特性,攻击者需要知道具体的噪声值 β_i ,而这种恢复的难度相当于从一个标准差为 σ 的高斯分布中精确估计出噪声 β_i ,其成功概率趋近于零。具体来说,攻击者恢复原始生物特征数据 Bio_i 的成功概率为

$$\Pr[\text{Reconst}(Bio_i | \widehat{Bio}_i)] \leq \frac{1}{\sqrt{2\pi\sigma^2}} \cdot e^{-\frac{(\beta_i - \mu)^2}{2\sigma^2}} \approx 0.$$

其中, $\mu = 0$,因为高斯噪声 $\beta_i \sim \mathcal{N}(0, \sigma^2)$ 的均值为0。

在BPPVFL方法中,每个参与者的生物特征数据在上传前均经过差分隐私保护处理。通过添加符合高斯分布的噪声,BPPVFL方法确保了攻击者难以推断或重构原始生物特征数据。结合差分隐私的理论保证和高斯机制的数学性质,BPPVFL方法能够有效实现对参与者生物特征的隐私保护。

由于高斯机制满足 (ϵ, δ) -差分隐私的要求,且攻击者难以从扰动后的数据中恢复原始生物特征数据,BPPVFL方法能够实现对参与者生物特征的隐私保护。因此,定理4得证。

证毕。

综上所述,BPPVFL方法通过结合生物特征识别、差分隐私、数字签名和同态哈希,成功地构建了一个安全、隐私保护和可验证的联邦学习系统。该系统能够有效抵御身份伪造、模型篡改、投毒攻击和隐私泄露等多种安全威胁,确保联邦学习过程的安全性、可信性和隐私性。

6.2 BPPVFL方法的验证开销分析

本节分别从客户端和服务端分析了BPPVFL的验证开销。

定理5. 在BPPVFL方法中,每个客户端的验证计算开销和验证通信开销分别为 $O(h) + O(s)$

和 $O(h) + O(s)$ 。

证明. 在 BPPVFL 方法中, 每个客户端在接收到来自服务器的全局模型更新后, 需要验证更新的真实性和完整性。客户端首先对服务器发送的全局模型更新及其附带的数字签名进行验证, 以确保更新的来源可信且未被篡改。验证一个签名的计算复杂度为 $O(s)$, 其中 s 是签名的长度。

客户端计算接收到的全局模型的同态哈希值, 并将其与服务器发送的哈希值进行比较, 以确保模型更新没有被篡改。计算同态哈希的复杂度为 $O(h)$, 其中 h 是哈希值的长度。每个客户端只需要对从服务器接收的一次全局模型更新进行签名验证和同态哈希验证。这两个验证步骤是单次操作, 与客户端数量 N 无关。因此, 对于每个客户端来说, 验证计算的总开销为

$$E_{\text{verify, client}} = O(s) + O(h).$$

由于每个客户端的验证操作独立于其他客户端, 它们不相互影响, 因此该验证开销与客户端数量 N 无关。可得, 每个客户端的验证开销为 $O(h) + O(s)$, 与客户端数量 N 无关。

在 BPPVFL 中, 哈希值的长度是预定义的, 与输入数据的大小无关。因此, 固定长度的哈希值在验证过程中对通信开销的影响是一个常量, 和数据的规模没有直接关系。客户端在验证服务器发送的固定长度哈希值传输开销为固定的 H 字节。因此, 客户端的验证通信开销为

$$C_{\text{client}} = O(H) + O(S).$$

其中, $O(H)$ 是固定长度哈希值的大小, $O(S)$ 是数字签名的大小。由于 H 和 $O(S)$ 都独立于 d , 客户端验证通信开销与梯度维度 d 无关。

证毕。

定理 6. 在 BPPVFL 中, 服务器验证计算开销和验证通信开销分别为 $N \times (O(h) + O(s))$ 和 $N \times (O(H_c) + O(S_c))$ 。

证明. 在 BPPVFL 方法中, 服务器接收所有客户端的局部模型更新, 并需要验证这些更新的合法性和完整性, 以防止恶意客户端上传伪造或被篡改的模型更新。

对于每个客户端, 服务器需要计算其提交的模型更新的同态哈希值, 并将其与客户端提交的哈希值进行比较。如果计算复杂度为 $O(h)$, 那么对一个客户端进行哈希验证的开销为 $O(h)$ 。

为了进一步确保模型更新的真实性, 服务器可

以选择对每个客户端提交的模型更新进行数字签名验证。验证一个签名的计算复杂度为 $O(s)$ 。因此, 对一个客户端进行签名验证的开销为 $O(s)$ 。服务器需要对所有 N 个客户端的模型更新进行验证。由于这 N 次验证相互独立且顺序无关, 因此服务器总的验证开销为

$$E_{\text{verify, server}} = N \times (O(h) + O(s)).$$

由于开销中包含 (N) 的乘法项, 可以看出验证开销与客户端数量 (N) 成正比。

服务器的验证开销为 $N \times (O(h) + O(s))$, 与客户端数量 (N) 成正比。当客户端数量增加时, 服务器的验证开销会显著增加。服务器需要对每个客户端提交的本地模型更新进行验证, 这涉及固定长度哈希值传输开销 (H_c) 字节 (常量)。因此, 服务器的验证通信开销为

$$C_{\text{server}} = N \times (O(H_c) + O(S_c)).$$

其中, $O(M_c)$ 是每个客户端提交的本地模型更新大小, $O(H_c)$ 是固定长度的哈希值大小 (常量), $O(S_c)$ 是数字签名大小。

使用固定长度哈希值优化了验证通信开销的表达式, 使得其计算更加明确。无论客户端和服务端, 哈希值部分的通信开销均是一个常量, 因此在分析通信开销时可以将其视为一个固定的成本。更大的通信开销主要源于签名验证的大小, 而服务器端的开销还会随客户端数量的增加而显著增加。

证毕。

以上定理详细描述了客户端和服务端在 BPPVFL 中验证开销的计算复杂度。定理 5 表明客户端的验证通信开销与客户端数量无关和梯度维度无关。定理 6 表明服务器的验证开销与客户端数量呈线性关系。这些定理为优化联邦学习中的验证机制提供了理论基础。

7 实验结果与分析

本节通过将 BPPVFL 与几种当前最先进的隐私保护和可验证联邦学习方案 (包括 VERIFL^[25]、VerifyNet^[46]、FedAvg^[47] 和 NbAFL^[11]) 进行对比, 评估其在收敛性、隐私保护水平和验证开销方面的性能。为了确保实验结果的广泛性和实际应用性, 我们选择了三个真实世界的生物识别数据集来进行实验, 分别是 SigD^[48]、BIDMC^[49] 和 TBME^[50] 数据集。

7.1 实验设置

SigD数据集派生自MIMIC-III临床数据库^[51]和MIMIC-III波形数据库匹配子集^[52],包含来自232名患者的PPG(光电容积图)波形数据。该数据集用于评估联邦学习方法在大规模、多样化数据中的表现。BIDMC数据集是MIMIC-II数据库的一个子集,包含来自53名重症监护病房(ICU)患者的生理数据,采样率为125 Hz。数据主要包括PPG和呼吸信号^[53]。该数据集测试了方法在重症监护场景中的身份验证性能。TBME数据集包含42例原始PPG信号,每个信号在手术和全身麻醉期间记录8分钟^[50]。通过文献中的心搏分离算法^[54],将连续信号分割成心跳序列,以用于实验中的数据样本。这些数据集的详细统计信息见表1。每个数据集被划分为训练样本和测试样本,设置了每个客户端的训练和测试样本数量,确保实验环境与实际应用场景一致。

表1 身份验证数据集

特征	SigD	BIDMC	TBME
训练样本	105 822	29 694	24 227
测试样本	10 582	2969	2422
客户训练样本数	320	320	320
客户测试样本数	100	100	100

实验中使用的生物识别验证模型如表2所示。模型包括卷积层、批归一化层、池化层、双向LSTM层和全连接层等组件,能够有效提取和学习生物特征数据的时序和空间特征,以实现高效的身份验证。实验在台式机上进行,实验设备的详细配置如表3所示。实验环境配备了64位Ubuntu 16.09 LTS操作系统,Intel i5-13400 CPU,NVIDIA 3090 GPU和32 GB内存,确保了在大规模数据训练和模型验证过程中的高效计算和可靠性。为了确保数据的隐私和安全性,BPPVFL方法采用了以下加密方案(见表4):同态加密采用CKKS(Cheon-Kim-Kim-Song)方案^[45],用于保护模型参数的传输和计算过程中的隐私。数字签名使用NIST P-256椭圆曲线,为模型梯度和聚合结果提供安全的身份认证和完整性验证。通过结合差分隐私、同态加密和数字签名技术,我们确保在整个联邦学习过程中,参与者的生物特征数据和模型梯度数据的隐私性和安全性得到充分保护。

在三个真实世界的生物识别数据集上,我们评估了BPPVFL的以下指标:

(1) 收敛性:通过训练轮次评估全局模型的准

表2 生物识别验证模型

组件	描述
输入维度	(Batch, Features, Time Steps)
卷积层	nn. Conv1d(1, 32, 5, padding=2)
批归一化	nn. BatchNorm1d(32)
池化层	nn. MaxPool1d(2)
卷积层	nn. Conv1d(32, 64, 5, padding=2)
批归一化	nn. BatchNorm1d(64)
池化层	nn. MaxPool1d(2)
LSTM	nn. LSTM(64, 64, bidirectional=True)
全连接层	nn. Linear(128, num_classes)
激活函数	ReLU
输出维度	(Batch, num_classes)

表3 设备环境

设备	台式机
操作系统	64位 Ubuntu 16.09 LTS
CPU	Intel i5-13400
GPU	NVIDIA 3090
内存	32 GB

表4 加密设置

加密机制	加密算法
同态加密	CKKS
数字签名	NIST P-256 曲线

确性和稳定性。

(2) 通信开销:测量在联邦学习过程中,参与者和服务器之间传输的数据量。

(3) 计算开销:评估参与者设备和服务器在不同加密和验证设置下的计算时间和资源消耗。

这些指标的对比实验能够全面展示BPPVFL在隐私保护和安全性保证的前提下,与现有方法相比的优势和性能表现。

7.2 隐私和收敛性评估

为了评估BPPVFL的隐私保护性能,我们选取了经典的联邦学习算法FedAvg(不考虑隐私保护机制)和近期提出的隐私保护联邦学习算法NbAFL作为基准进行对比分析。其中,FedAvg作为无隐私保护的联邦学习基准,反映了理想情况下模型的性能上限,而NbAFL则体现了当前典型的隐私保护方法的效果。我们评估了BPPVFL、FedAvg和NbAFL三种算法在三个真实世界生物识别数据集(SigD、BIDMC和TBME)上的收敛性和隐私保护性能。通过调整迭代次数($T=1, 2, \dots, 20$)和固定隐私预算($\epsilon=2$),我们分析了各算法在不同数据集上的预测准确率和损失函数的变化情况。

如图2所示,我们比较了提出的BPPVFL方法与基准模型FedAvg、NbAFL三种算法在不同迭代次数下的预测准确率,以评估算法的收敛性。结果表明,BPPVFL在所有数据集上均表现出较好的收敛性,特别是在迭代次数较高时,其准确率显著提升,明显优于现有的隐私保护方法NbAFL。例如,在SigD数据集上,BPPVFL的准确率从初始的0.8029逐渐提高到0.8572,相比于NbAFL(准确率为0.3929提高到0.4762),准确率提高了约38.1%;尽管FedAvg在最终收敛准确率(0.9203)上仍然表现最佳,但BPPVFL在实现隐私保护的前提下显著缩小了与FedAvg的差距(提升至0.8572,较NbAFL提升幅度显著)。在BIDMC数据集上,BPPVFL的准确率从0.7464提高到0.8569,而

NbAFL仅从0.3334提高到0.4708,BPPVFL相较于NbAFL准确率提升高达约38.6%;FedAvg准确率达到0.9119,BPPVFL相比之下保持了较为接近的性能,同时实现了隐私保护。在TBME数据集上,BPPVFL的准确率从初始的0.6385提高到0.8639,相比之下,NbAFL最终准确率仅为0.4633,BPPVFL相较于NbAFL的准确率提高了约40.1%。虽然FedAvg仍具备最高准确率(0.9061),但BPPVFL成功在隐私保护的前提下表现出接近的性能。总体而言,BPPVFL在实现隐私保护的前提下显著超过了现有隐私保护方法NbAFL的收敛速度和准确率,接近无隐私保护的FedAvg性能水平,这验证了BPPVFL的有效性和优势。

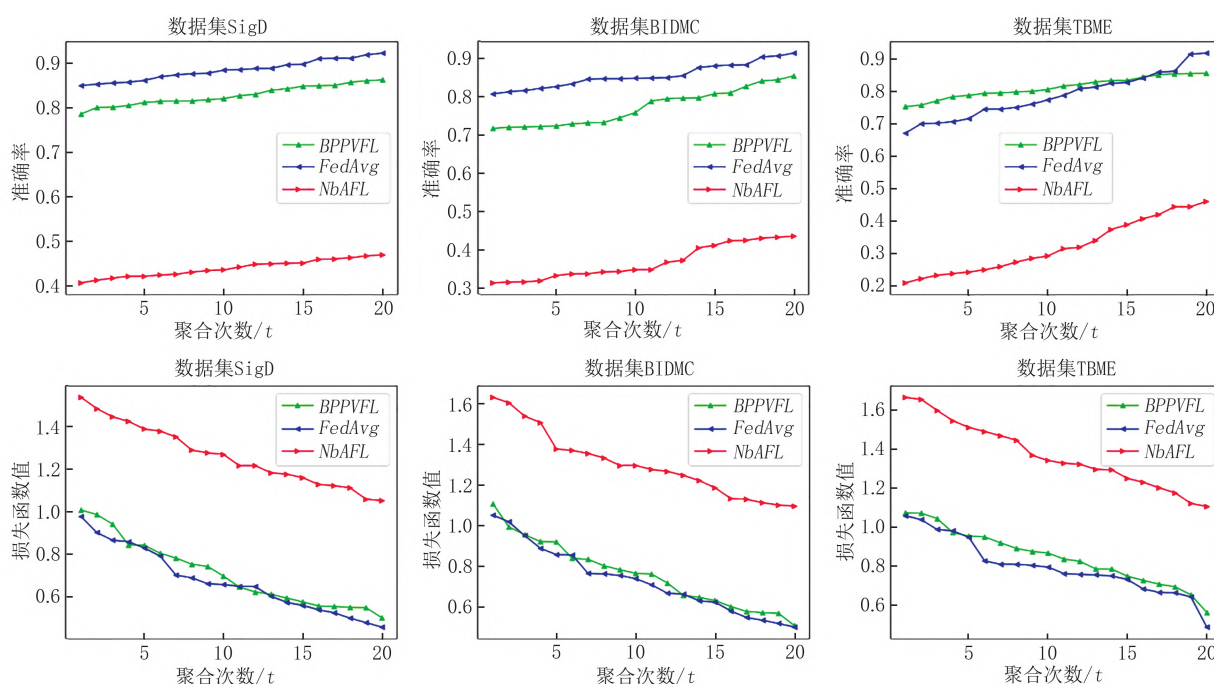


图2 在SigD、BIDMC和TBME数据集上,在不同迭代次数下FedAvg、NbAFL和BPPVFL三个对比方法的平均准确率曲线和损失函数值

如图3所示,我们进一步比较了BPPVFL与基准模型FedAvg(无隐私保护机制)和NbAFL(隐私保护基准)在不同隐私预算(ϵ)下的预测准确率,数据集包括SigD、BIDMC和TBME。实验表明,BPPVFL的预测准确率随隐私预算增大而显著提升,并且始终优于现有隐私保护基准NbAFL。在SigD数据集上,当隐私预算从0.2增加到1.0时,BPPVFL的准确率从0.9241提升到0.9672,相较于NbAFL(从0.6219提升到0.9089),表现出最高约30.2%的相对提升;虽然FedAvg(准确率

0.9718)因无隐私保护保持稳定且略高,但BPPVFL在隐私保护下实现了近似的性能水平。在BIDMC数据集上,BPPVFL的准确率从0.9281提高到0.9699,而NbAFL仅从0.8442提高到0.9321,BPPVFL的准确率比NbAFL最多高出约10.1%,显示出更为有效的隐私保护性能与准确率平衡。相比FedAvg的稳定表现(0.9776),BPPVFL展现出更好的鲁棒性和平衡能力。在TBME数据集上,BPPVFL从0.9225提高到0.9684,而NbAFL在较低隐私预算(0.2)仅为

0.2716,即使最高也未超过0.9,BPPVFL相对提升高达256.7%(低隐私预算条件下);尽管FedAvg(0.9715)仍然最高,但BPPVFL在保证隐私的同时展现出非常接近的性能。综上所述,BPPVFL通过自适应噪声机制,在不同隐私预算

下均表现出比现有隐私保护方法NbAFL显著更好的准确率和稳定性,且性能接近于无隐私保护的FedAvg。这证明了BPPVFL在隐私保护和模型性能间实现了有效权衡,具备明显的实际应用优势。

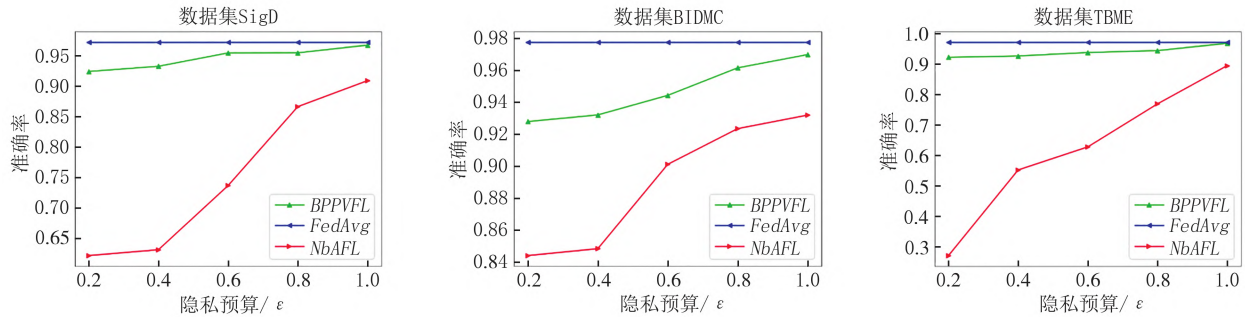


图3 在SigD、BIDMC和TBME数据集上,在不同隐私预算下FedAvg、NbAFL和BPPVFL三种对比方法的平均准确率曲线

BPPVFL在三个数据集上展现了优异的收敛性和隐私保护性能,这得益于其为生物特征数据量身定制的自适应噪声设计,使得在不同的隐私预算下能有效平衡隐私保护强度和模型性能。相较之下,FedAvg没有考虑隐私保护策略,因此表现稳定但未能体现隐私保护的意义。NbAFL虽然在高隐私预算下有所提升,但其在低隐私预算下的表现不佳,难以在强隐私保护需求下实现出色的模型表现。总体而言,BPPVFL能够更好地适应生物识别应用场景,是隐私保护联邦学习的有效解决方案。

7.3 验证开销评估

如图4所示,我们对提出的BPPVFL方法与基准方法VerifyNet(零知识证明)和VeriFL

(传统同态哈希)在不同维度的向量数据下的验证通信开销。实验结果显示,BPPVFL的验证通信开销在各维度下均稳定保持为0.095 KB,而VerifyNet则随着维度增加从 4.88×10^2 KB上升到 5.88×10^2 KB,BPPVFL相较于VerifyNet实现了数千倍的通信开销减少;而相比VeriFL(从35.1 KB到45.7 KB),BPPVFL的通信开销也降低了约99%以上,体现出显著的性能优势。BPPVFL由于采用了基于生物特征的轻量级签名验证机制,验证通信开销与梯度维度无关,在服务器端显著减少了计算与通信负担。相较于VerifyNet和VeriFL,BPPVFL在实际应用中的通信效率优势明显,更适用于资源受限的分布式生物识别场景

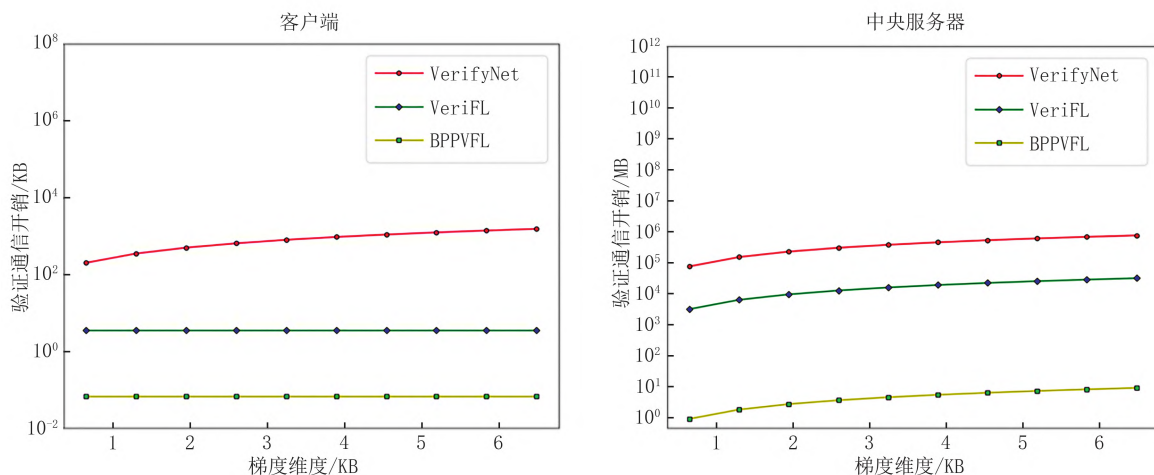


图4 在不同维度向量下,VeriFL、VerifyNet和BPPVFL在服务器和客户端上验证开销

8 结 论

本文提出的BPPVFL方法旨在解决联邦学习中的隐私保护和高效验证问题。通过设计生物特征身份验证机制和数据验证机制,BPPVFL在保持模型性能和收敛性的同时,有效降低了验证开销并增强了隐私保护能力。我们在三种真实世界的生物识别数据集(SigD、BIDMC和TBME)上进行了广泛的实验,并将BPPVFL与当前先进的联邦学习方案进行了对比分析。结果表明,BPPVFL在收敛性、隐私保护和验证开销方面均展现出显著优势。

在收敛性分析中,BPPVFL通过自适应噪声设计充分利用每轮迭代的信息,在所有数据集上均表现出良好的收敛性和模型准确率,准确率最高提升了81%。实验结果还表明,BPPVFL在隐私预算变化下的损失函数表现稳定,能够有效平衡隐私保护和模型性能,这得益于其基于生物特征的差分隐私噪声机制。在验证开销方面,BPPVFL通过生物特征身份验证和CKKS加密,显著降低了客户端和服务端端的验证通信和计算开销,客户端通信开销最多减少了85%,服务器端通信开销减少了90%以上。理论和实验证明了BPPVFL在客户端的验证通信开销维度独立。

未来工作中,我们计划进一步优化方法的隐私保护策略和验证算法,并在更大规模和更多样化的数据集上进行测试,以进一步验证其通用性和鲁棒性。BPPVFL的设计理念为未来隐私保护联邦学习的研究和应用提供了新的方向和可能性。

参 考 文 献

- [1] Wu Q, Dong C, Guo F, et al. Privacy-preserving federated learning for power transformer fault diagnosis with unbalanced data. *IEEE Transactions on Industrial Informatics*, 2024, 20(4): 5383-5394
- [2] Qian W, Shen Q, Wu P, et al. Research progress on privacy protection technology in big data computing environments. *Chinese Journal of Computers*, 2022, 45(4): 669-701 (in Chinese)
(钱文君, 沈晴霓, 吴鹏飞, 等. 大数据计算环境下的隐私保护技术研究进展. *计算机学报*, 2022, 45(4): 669-701)
- [3] Gao Y, Chen L, Han J, et al. Iot privacy-preserving datamining with dynamic incentive mechanism. *IEEE Internet of Things Journal*, 2024, 11(1): 777-790
- [4] Rot P, Grm K, Peer P, et al. Privacy prober: Assessment and detection of soft-biometric privacy-enhancing techniques. *IEEE Transactions on Dependable and Secure Computing*, 2024, 21(4): 2869-2887
- [5] Adhikary S, Karmakar A. A fast rlwe-based IPFE library and its application to privacy-preserving biometric authentication. *IEEE Transactions on Emerging Topics in Computing*, 2024, 12(1): 344-356
- [6] Roig D. O, González-Soler L. J, Rathgeb C, et al. Privacy-preserving multi-biometric indexing based on frequent binary patterns. *IEEE Transactions on Information Forensics and Security*, 2024, 19: 4835-4850
- [7] Coelho K. K, Tristão E. T, Nogueira M, et al. Multimodal biometric authentication method by federated learning. *Biomedical Signal Processing and Control*, 2023, 85: 105022
- [8] Zhang L, Li A, Chen S, et al. A secure, flexible, and ppg-based biometric scheme for healthy iot using homomorphic random forest. *IEEE Internet of Things Journal*, 2024, 11(1): 612-622
- [9] Dwork C, Roth A. The algorithmic foundations of differential privacy. *Foundations and Trends in Theoretical Computer Science*, 2014, 9(3-4): 211-407
- [10] Blum A, Dwork C, McSherry F, et al. Practical privacy: the sulq framework//*Proceedings of the ACM SIGACT-SIGMOD-SIGART, Maryland, USA, 2005*, 128-138
- [11] Wei K, Li J, Ding M, et al. Federated learning with differential privacy: Algorithms and performance analysis. *IEEE Transactions on Information Forensics and Security*, 2020, 15: 3454-3469
- [12] Zhou H, Yang G, Dai H, et al. PFLF: privacy-preserving federated learning framework for edge computing. *IEEE Transactions on Information Forensics and Security*, 2022, 17: 1905-1918
- [13] Li X, Yan H, Cheng Z, et al. Protecting regression models with personalized local differential privacy. *IEEE Transactions on Dependable and Secure Computing*, 2023, 20(2): 960-974
- [14] Sun X, Zhang P, Liu J. K, et al. Private machine learning classification based on fully homomorphic encryption. *IEEE Transactions on Emerging Topics in Computing*, 2020, 8(2): 352-364
- [15] Zhang Q, Xin C, Wu H. GALA: greedy computation for linear algebra in privacy-preserved neural networks//*Proceedings of the 28th Annual Network and Distributed System Security Symposium, Virtually, 2021*, 21-25
- [16] Phong L. T, Aono Y, Hayashi T, et al. Privacy-preserving deep learning via additively homomorphic encryption. *IEEE Transactions on Information Forensics and Security*, 2018, 13(5): 1333-1345
- [17] Lin X, Wu J, Li J, et al. Friend-as-learner: Socially-driven trustworthy and efficient wireless federated edge learning. *IEEE Transactions on Mobile Computing*, 2023, 22(1): 269-283
- [18] Yang Q, Liu Y, Chen T, et al. Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 2019, 10(2): 12:1-12:19
- [19] Yang Q. Ai and data privacy protection: The solution of federated learning. *Information Security Research*, 2019,

- 5(11): 961-965. (in Chinese)
(杨强. AI与数据隐私保护: 联邦学习的破解之道. 信息安全研究, 2019, 5(11): 961-965)
- [20] Mohassel P, Zhang Y. Secureml: A system for scalable privacy-preserving machine learning//Proceedings of the IEEE Symposium on Security and Privacy, San Jose, USA, 2017, 19-38
- [21] Mohassel P, Rindal P. Abys: A mixed protocol framework for machine learning//Proceedings of the ACM Special Interest Group on Security, Audit, and Control, Toronto, Canada, 2018, 35-52
- [22] Xu R, Baracaldo N, Zhou Y, et al. Hybridalpha: An efficient approach for privacy-preserving federated learning//Proceedings of the ACM Conference on Computer and Communications Security, London, UK, 2019, 13-23
- [23] Zhou C, Sun Y, Wang D, et al. A survey of federated learning research. Journal of Network and Information Security, 2021, 7(05): 77-92 (in Chinese)
(周传鑫, 孙奕, 汪德刚, 等. 联邦学习研究综述. 网络与信息安全学报, 2021, 7(05): 77-92)
- [24] Zhu J, Zhang Q, Gao S, et al. A privacy-preserving trusted federated learning model based on blockchain. Chinese Journal of Computers, 2021, 44(12): 2464-2484 (in Chinese)
(朱建明, 张沁楠, 高胜, 等. 基于区块链的隐私保护可信联邦学习模型. 计算机学报, 2021, 44(12): 2464-2484)
- [25] Guo X, Liu Z, Li J, et al. Verifl: Communication-efficient and fast verifiable aggregation for federated learning. IEEE Transactions on Information Forensics and Security, 2021, 16: 1736-1751
- [26] Yang C.-Z, Ma J, Wang S, et al. Preventing deep fake attacks on speaker authentication by dynamic lip movement analysis. IEEE Transactions on Information Forensics and Security, 2020, 16: 1841-1854
- [27] Liu-Jimenez J, Sanchez-Reillo R, Fernandez-Saavedra B. Iris biometrics for embedded systems. IEEE Transactions on Very Large Scale Integration (VLSI) Systems, 2009, 19(2): 274-282
- [28] Mahfouz A, Mahmoud T. M. Eldin A. S. A survey on behavioral biometric authentication on smartphones. Journal of Information Security and Applications, 2017, 37: 28-37
- [29] Abdullakutty F, Elyan E, Johnston P. A review of state-of-the-art in face presentation attack detection: From early development to advanced deep learning and multi-modal fusion methods. Information Fusion, 2021, 75: 55-69
- [30] Yap H. Y., Choo Y.-H., Mohd Yusoh Z. I., et al. An evaluation of transfer learning models in eeg-based authentication. Brain Informatics, 2023, 10(1): 19
- [31] Wang H, Su L, Zeng H, et al. Anti-spoofing study on palm biometric features. Expert Systems with Applications, 2023, 218: 119546
- [32] Li L, Chen C, Pan L, et al. Hiding your signals: A security analysis of ppg-based biometric authentication//Proceedings of the European Symposium on Research in Computer Security, The Hague, The Netherlands, 2023, 183-202
- [33] Zhao T, Wang Y, Liu J, et al. True heart: Continuous authentication on wrist-worn wearables using ppg-based biometrics//Proceedings of the IEEE Conference on Computer Communications, Toronto, Canada, 2020, 30-39
- [34] Li L, Chen C, Pan L, et al. Siga: rppg-based authentication for virtual reality head-mounted display//Proceedings of the 26th International Symposium on Research in Attacks, Intrusions and Defenses, Hong Kong, China, 2023, 686-699
- [35] Guo J, Mu H, Liu X, et al. Federated learning for biometric recognition: a survey. Artificial Intelligence Review, 2024, 57(8): 208
- [36] Gupta H, Rajput T. K., Vyas R, et al. Biometric iris identifier recognition with privacy preserving phenomenon: A federated learning approach//Proceedings of the International Conference on Neural Information Processing, Online, 2022, 493-504
- [37] Wei K, Li J, Ding M, et al. Performance analysis and optimization in privacy-preserving federated learning. arXiv preprint arXiv: 2003.00229, 2020.
- [38] Jiang B, Li J, Wang H, et al. Privacy-preserving federated learning for industrial edge computing via hybrid differential privacy and adaptive compression. IEEE Transactions on Industrial Informatics, 2021, 19(2): 1136-1144
- [39] Wang J, Guo S, Xie X, et al. Protect privacy from gradient leakage attack in federated learning//Proceedings of the IEEE Conference on Computer Communications, London, UK, 2022, 580-589
- [40] Ye D, Shen S, Zhu T, et al. One parameter defense—defending against data inference attacks via differential privacy. IEEE Transactions on Information Forensics and Security, 2022, 17: 1466-1480
- [41] Zhao F, Ren X, Yang S, et al. Latent dirichlet allocation model training with differential privacy. IEEE Transactions on Information Forensics and Security, 2020, 16: 1290-1305
- [42] Wang Q, Li Z, Zou Q, et al. Deep domain adaptation with differential privacy. IEEE Transactions on Information Forensics and Security, 2020, 15: 3093-3106
- [43] Zhou H, Yang G, Xiang Y, et al. A lightweight matrix factorization for recommendation with local differential privacy in big data. IEEE Transactions on Big Data, 2023, 9(1): 160-173
- [44] Duchi J. C., Jordan M. I., Wainwright M. J. Local privacy and statistical minimax rates.//Proceedings of the IEEE Symposium on Foundations of Computer Science, Berkeley, USA, 2013, 429-438
- [45] Tian H, Zeng C, Ren Z, et al. Sphinx: Enabling privacy-preserving online learning over the cloud.//Proceedings of the IEEE Symposium on Security and Privacy, San Francisco, USA, 2022, 2487-2501
- [46] Xu G, Li H, Liu S, et al. Verifynet: Secure and verifiable federated learning. IEEE Transactions on Information Forensics and Security, 2020, 15: 911-926
- [47] McMahan B, Moore E, Ramage D, et al. Communication-efficient learning of deep networks from decentralized data//Proceedings of the International Conference on Artificial Intelligence and Statistics, Fort Lauderdale, USA, 2017, 54: 1273-1282
- [48] Li L, Chen C, Pan L, et al. Sigd: A cross-session dataset for

- ppg-based user authentication in different demographic groups// Proceedings of the International Joint Conference on Neural Networks, Gold Coast, Australia, 2023, 1-8
- [49] Pimentel M. A. F, Johnson A. E. W, Charlton P, et al. Toward a robust estimation of respiratory rate from pulse oximeters. *IEEE Transactions on Biomedical Engineering*, 2017, 64(8): 1914-1923
- [50] Karlen W, Kobayashi K, Ansermino J M, et al. Respiratory rate estimation using respiratory sinus arrhythmia and photoplethysmogram. *IEEE Transactions on Biomedical Engineering*, 2013, 60(11): 3283-3290
- [51] Johnson A, Pollard T. Mark R. MIMIC-III clinical database (version 1.4). PhysioNet, 2016 <https://doi.org/10.13026/C2XW26>
- [52] Moody B, Moody G, Villarreal M, et al. MIMIC-III waveform database matched subset (version 1.0). PhysioNet, 2020 <https://doi.org/10.13026/c2294b>
- [53] Saeed M, Villarreal M, Reisner A. T, et al. Multiparameter Intelligent Monitoring in Intensive Care II (MIMIC-II): a public-access intensive care unit database. *Critical Care Medicine*, 2011, 39(5): 952
- [54] Lovisotto G, Turner H, Eberz S, et al. Seeing red: Ppg biometrics using smartphone cameras// Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops, Online, 2020, 3565-3574



ZHOU Hao, PhD., lecturer. His main research interests include federated learning, privacy preserving, verification, and differential privacy.

DAI Hua, Ph. D., professor. His main research interests include data management and security, database security, and behavior recognition.

Background

This paper addresses the challenges of privacy-preserving and verifiable federated learning (FL), a critical problem in secure machine learning and distributed data privacy. Federated learning allows multiple clients, such as mobile devices or edge servers, to collaboratively train machine learning models without sharing their private data. While FL offers solutions to data privacy concerns, it introduces new challenges in ensuring the security, robustness, and efficiency of the learning process. Specifically, privacy protection and integrity verification during training remain key research areas. Several challenges persist in federated learning. A major issue is developing robust privacy mechanisms to secure highly sensitive data, such as biometric information, against inference attacks. Existing methods, like FedAvg, lack privacy-preserving mechanisms, making them susceptible to privacy leaks. Approaches like VerifyNet and VeriFL offer privacy and security guarantees but at the cost of high computational and communication overhead due to zero-knowledge proofs or traditional homomorphic hashing and secure multiparty computation. These methods, while secure, can be computationally prohibitive, especially for clients with limited

YANG Geng, Ph. D., professor. His main research interests include computer networks and communications, information and network security, distributed and parallel computing, wireless sensor networks, and security.

HUANG Yu-Xian, Ph. D. candidate. His main research interests include privacy protection, Byzantine attacks, and federated learning.

WANG Zhou-Sheng, Ph. D. His main research interests include federated learning, biometric verification, and edge computing.

resources or high-dimensional data. Thus, there is a need for efficient solutions that balance privacy, verification, and computational efficiency.

This paper proposes a novel solution: the Biometric-based Privacy-Preserving Verifiable Federated Learning (BPPVFL) framework. BPPVFL leverages adaptive noise mechanisms tailored to biometric data and the CKKS homomorphic encryption scheme to achieve lightweight identity verification and data validation. By integrating biometric-based authentication, BPPVFL ensures robust identity verification while protecting sensitive biometric data. This approach significantly reduces verification overhead on both the client and server sides, making it practical for real-world applications. Experimental results on three biometric datasets (SigD, BIDMC, and TBME) demonstrate that BPPVFL achieves superior convergence rates and lower costs than VerifyNet and VeriFL, especially under varying privacy budgets and data dimensions.

In summary, this paper advances privacy-preserving federated learning by providing an efficient solution addressing privacy protection and verification overhead. BPPVFL offers a

new direction for federated learning applications in biometric data, ensuring secure and scalable training while maintaining high levels of privacy. This framework mitigates the trade-offs between privacy and performance and lays a foundation for future research in secure machine learning for sensitive data environments.

This work was supported in part by the National Natural

Science Foundation of China under Grant (Grant No. 62372244), the Natural Science Research Start-up Foundation of Recruiting Talents of Nanjing University of Posts and Telecommunications (Grant No. NY224058), the Open Project of Key Laboratory in Hospital Management Direction (Grant No. 2024LYKC003).