



基于静态假设的去中心化属性基内积函数加密

周彦伟¹, 刘现响¹, 张广进¹, 乔子芮^{2*}, 卢永辉^{1*}, 杨波¹, 朱天清³, 张明武⁴

1. 陕西师范大学人工智能与计算机学院, 西安 710119

2. 西安邮电大学网络空间安全学院, 西安 710121

3. 澳门城市大学数据科学学院, 澳门 999078

4. 湖北工业大学计算机学院, 武汉 430068

* 通信作者. E-mail: qzr@xupt.edu.cn, lyhui008@snnu.edu.cn

收稿日期: 2025-06-04; 修回日期: 2025-07-28; 接受日期: 2025-09-09; 网络出版日期: 2026-06-30

国家自然科学基金 (批准号: 62502380, 62272287, 62472150, U24B20149, 62436004, 62372317)、陕西省重点研发计划 (批准号: 2024GX-YBXM-074) 和中央高校基本科研业务费专项资金 (批准号: GK202301009) 资助项目

摘要 传统加密机制中, 数据需先解密才能获取具体内容, 对于隐私敏感的场景其无法为原始数据提供必要的隐私性, 从而削弱了加密方案的实用性. 近年来, 为完成内积运算的同时实现数据的隐私保护需求, 一种新的密码原语——内积函数加密机制被提出, 它能在不解密的前提下通过直接操作密文实现获得相应数据内积值的目的, 保证了数据的隐私性, 其中基于身份的内积函数加密机制缺乏对密文的细粒度访问控制; 基于素数阶群的属性基内积函数加密机制通常采用非静态的困难性假设证明其安全性, 该操作虽能保持常规攻击下的安全性, 但其安全性取决于敌手的询问次数, 导致相应构造易受到攻击; 此外, 现有构造普遍需要中心权威完成全局公私钥的生成, 中心化构造模式导致相应提议存在属性权威单点失效的隐患和可扩展性弱的性能瓶颈. 针对上述挑战, 本文基于素数阶群和静态安全性假设提出了去中心化的属性基内积函数加密机制, 基于决策的双线性 Diffie-Hellman 困难假设证明了它的安全性, 有效规避了非静态困难性假设带来的潜在安全风险, 显著提升了属性基内积函数加密机制的整体安全性. 此外, 本文所提出的方案在增强安全性的同时, 还具备支持大规模属性集合、去中心化的权威管理和抗合谋攻击等良好性能. 分析表明, 与现有的内积函数加密机制相比, 本文构造具有更优的计算效率, 更加适合在物联网等资源受限环境下的部署. 理想模型下被证明安全的密码原语, 由于普遍存在的边信道、冷启动等泄露攻击, 导致其在现实环境中不再具有原始的安全性, 为增强 AB-IPFE 机制的实用性, 本文设计了抗泄露的去中心化 AB-IPFE 机制, 并对该方案实现抗泄露攻击的关键技术和理论进行了详细分析.

关键词 内积函数加密, 属性基内积函数加密, 去中心化, 静态安全性假设, 泄露容忍性

1 引言

传统加密方案若需获取数据及其统计信息, 必须通过解密操作方能获得具体明文数据, 这在涉及敏感信息的应用场景中可能引发潜在的隐私泄露风险. 为此, 函数加密 (function encryption, FE) 提出了一种创新性的解决范式——其核心在于无需对密文进行完全解密即可实现数据处理. 为强化数据隐私保护机制, 内积函数加

引用格式: 周彦伟, 刘现响, 张广进, 等. 基于静态假设的去中心化属性基内积函数加密. 中国科学: 信息科学, 2026, 56: 1652–1664, doi: 10.1360/SSI-2025-0271

Zhou Y W, Liu X X, Zhang G J, et al. Decentralized attribute-based inner product function encryption scheme under static assumption. Sci Sin Inform, 2026, 56: 1652–1664, doi: 10.1360/SSI-2025-0271

密 (inner product function encryption, IPFE) 机制被作为 FE 的典型范式展开了系统性研究, 并相继提出了多个 IPFE 方案, 现有研究多聚焦于身份基内积函数加密机制^[1,2] 和分层身份基内积函数加密机制^[3,4].

1.1 研究现状

(1) **内积函数加密.** FE 由 Waters 等^[5] 首次提出, 其核心特性是允许在不解密原始数据的前提下获取指定计算结果. 作为 FE 的一种具体实现形式, IPFE 通过向量内积运算生成解密结果. Zhu 等^[6] 设计了基于身份分层架构的 IPFE 方案, 该方案在非静态困难性假设下达成安全性证明, 实现密文与密钥长度的恒定不变性, 并支持基于用户身份的一对一密文传输. 然而, 其依赖单一中心权威的架构存在单点故障风险, 制约了实际应用场景的适应性. 针对单一权威缺陷及共谋攻击的防御问题, Abdalla 等^[7] 基于 Lewko 与 Waters^[8] 提出的全局身份技术框架, 构建了去中心化 IPFE 机制, 能有效抵抗收买权威的共谋攻击行为, 增强了系统的安全性保障. 此后, Tseng 等^[9] 进一步提出具有固定密文与密钥长度的去中心化 IPFE 机制, 在存储与计算开销方面均实现了优化, 同时基于非静态困难性假设完成了严格的安全性论证. Han 等^[10] 近期亦提出一个新型的去中心化 IPFE 机制, 他们采用判定的双线性 Diffie-Hellman (decisional bilinear Diffie-Hellman, DBDH) 假设替代传统非静态困难性假设, 不仅完善了形式化安全证明体系, 而且确保方案的安全性与敌手的询问次数无关, 为 IPFE 机制安全强度的提升提供了新路径. 然而, 上述 IPFE 机制仅能实现基于身份的一对一信息传输, 无法支持一对多传输场景下的细粒度访问控制功能, 在一定程度上削弱了方案的实用性, 在涉及敏感信息处理的实际应用场景中具有局限性. 综上所述, 尽管现有的 IPFE 机制在隐私计算、去中心化架构及抗合谋攻击等方面展现出显著优势, 但其在细粒度访问控制功能的支持层面仍存在明显的局限性.

针对现有 IPFE 机制难以实现密文细粒度访问控制的缺陷, Chen 等^[11] 和 Belek 等^[12] 基于 Lewko 等的属性基加密 (attribute-based encryption, ABE) 机制^[13], 在合数阶群上设计了属性基内积函数加密 (attribute-based inner product function encryption, AB-IPFE) 机制的具体构造, 尽管上述方案为 IPFE 提供了细粒度访问控制的功能, 但计算开销较大导致其在实际应用中的适应性较差. 为进一步优化计算性能, Yang 等^[14] 基于 Bethencourt 等的 ABE 机制^[15], 在素数阶群中提出了 AB-IPFE 机制的一个新构造, 但其安全性是在非静态困难性假设下被证明的. 分析可知, 上述 AB-IPFE 机制虽具有细粒度访问控制的优势, 但存在属性集容量受限的问题, 使得实际部署时需通过反复执行系统初始化操作来扩展属性空间的大小; 同时存在权威中心单点故障及抗敌手合谋攻击能力弱的不足, 上述问题导致相应构造在计算效率、安全性及环境适应性等方面均受到了影响.

(2) **去中心化的细粒度访问控制.** 为解决传统 ABE 机制中属性权威单点失效对系统安全性所造成的危害, Lewko 与 Waters^[8] 率先提出去中心化的 ABE 机制的基础框架, 设计了多授权机构自主生成密钥对, 彻底消除了传统方案对中心化机构的依赖. 该方案通过整合全局身份标识与访问控制策略, 不仅增强了对合谋攻击的防御能力, 同时显著提升了实际部署的可行性. 在效率优化方面, Ye 等^[16] 基于素数阶群构建了新型的去中心化的 ABE 机制, 使系统同时满足语义安全性与操作不可伪造性的双重安全需求. 值得注意的是, 该方案缺乏严格的安全性证明过程. Zhang 等^[17] 基于素数阶群设计的去中心化的 ABE 机制引入在线/离线混合加密技术, 显著优化了加密阶段的运算效能. 该方案基于非静态困难性假设证明了相应的安全性, 使其在保持高效性的同时满足实际应用的安全需求. Cheng 等^[18] 则聚焦区块链融合方向, 提出区块链适配的去中心化的 ABE 机制, 利用分布式账本技术强化去中心化特性, 在区块链环境下展现出更优的系统兼容性与计算效率, 并通过非静态困难性假设证明建立了安全性理论保障. 现有去中心化的 ABE 机制虽已实现细粒度访问控制, 但存在密文信息泄漏风险, 攻击者可利用解密中间状态逆向推导原始明文数据, 导致数据隐私保障机制存在根本性漏洞; 此外, 现有方案均基于非静态困难性假设构建安全性证明, 这种弱安全基假设可能导致系统在面对关于敌手询问次数的攻击时, 存在潜在的安全威胁. 然而, 对于 IPFE 机制而言, 当前关于其去中心化细粒度访问控制功能的研究较少, 有关研究少有记载. 因此, 为增强 IPFE 机制的适用性, 亟须研究 IPFE 机制的去中心化细粒度访问控制技术.

(3) **IPFE 机制的泄露容忍性.** 在经典密码学理论框架下, 密码原语的安全性证明通常基于理想化的计算模型, 其安全性分析往往假设敌手仅能通过黑盒方式访问密码算法, 而无法获取系统运行过程中产生的任何额外信息. 然而, 现实物理系统的实现不可避免地存在信息泄露风险. 此类泄露攻击严重削弱了理论安全证明的有效

性,使得传统密码方案在实际部署中面临严峻挑战.泄露密码学的核心目标在于构建能够形式化抵抗部分密钥或内部状态泄露的密码方案,研究者已成功构建出具有可证明安全性的抗泄露公钥加密、抗泄露身份基加密以及抗泄露属性基加密等核心密码原语.部分研究^[19]基于 Goldreich-Levin 定理^[20]对 ABE 机制在辅助输入模型下的泄露容忍性进行了研究.尽管如此,现有抗泄露密码机制的研究范畴依然存在尚待解决的问题,特别是针对函数加密这一前沿领域,尤其是 IPFE 的泄露容忍性研究进展较缓慢,相关研究工作鲜有报道.为提升 IPFE 机制的实用性,需研究抗泄露的 IPFE 机制.

(4) 现有研究的局限性.根据上述分析,现有 IPFE 机制的研究不足本文总结如下. (i) 缺乏去中心化的细粒度访问控制.现有的去中心化 IPFE 机制^[7,9,10]基于身份实现信息传输的一对一控制功能,缺乏对数据传输的细粒度访问控制,容易使数据面临未经授权访问、潜在篡改等风险.虽然传统的 ABE 机制具有加密数据的细粒度控制能力,但先解密后使用的工作模式导致其无法实现对明文数据的保密计算功能. (ii) 缺乏静态困难性假设下的安全性论证.现有基于素数阶群构造的 AB-IPFE 机制^[6,7,9,10,14]均基于非静态困难性假设来证明相应构造的安全性,导致相关方案虽然可以抵御常规攻击,但仍可能遭受关于敌手询问次数的攻击,从而对方案的安全性构成较大的威胁. (iii) 缺乏对泄露容忍性的讨论.现有研究注重了内积的隐私运算过程,忽视了秘密信息的泄露对系统安全性所带来的隐患.由上述分析可知,基于静态假设构造去中心化的 AB-IPFE 机制已是 IPFE 领域当前亟须解决的关键问题;此外,AB-IPFE 机制的抗泄露性研究尚未引起研究者的关注,相关工作的报道较少.为充分说明上述问题研究的必要性,我们首先讨论下述两个疑问.

疑问 1: 能否将基于静态假设的去中心化 ABE 机制的构造方法通过自然转换得到基于静态假设的去中心化 AB-IPFE 机制? 该问题的答案是否定的.虽然文献^[21]提出了基于静态假设构造去中心化 ABE 机制的方法,但该方法无法直接迁移到去中心化 AB-IPFE 机制的构造,主要面临的问题有: (1) AB-IPFE 机制的用户私钥由向量和属性权威的主私钥组成,而上述 ABE 机制的构造方法中不涉及向量的相关运算; (2) 去中心化 AB-IPFE 机制的加密算法需考虑如何使用单个随机值实现对消息向量的隐藏,而上述 ABE 机制仅需使用随机值对单个消息进行隐藏,不涉及多消息的掩盖需求.

疑问 2: 能否采用传统公钥加密机制对泄露攻击的抵抗方法自然得到抗泄露的 AB-IPFE 机制? 该问题的答案同样是否定的.传统公钥加密的抗泄露处理方法一般采用哈希证明系统(或类似哈希证明系统的操作)生成具有一定最小熵的参数,在强随机性提取器的作用下,基于该参数生成隐藏明文的均匀随机密钥.由于 AB-IPFE 机制的加密算法不涉及类似哈希证明系统的运算,无法确保相应随机参数的熵能够满足强随机性提取器的要求.文献^[19]虽基于 Goldreich-Levin 定理,提出了辅助输入模型下抗泄露 ABE 机制的构造方法,但其无法直接应用到 AB-IPFE 机制,因为 Goldreich-Levin 定理实现抗泄露攻击的前提是加密算法需要两个向量进行内积求值计算,而去中心化的 AB-IPFE 机制的加密算法一般仅涉及单个向量的操作.

综上所述,现有 ABE 机制中基于静态假设的去中心化架构的构造方法^[21]及泄露攻击的抵抗技术^[19]很难直接应用,但它们的结论为研究基于静态安全性假设构造去中心化 AB-IPFE 机制及其泄露容忍性提供了技术指导 and 理论依据,本文受上述工作^[19,21]的启发,解决现有技术所面临的相关问题,研究基于静态假设的去中心化 AB-IPFE 机制及其泄露容忍性,并给出具体的方案构造.

1.2 本文的出发点及主要工作

当前,基于素数阶群构造的 IPFE 机制已成功应用于数据隐私保护领域,并在金融征信、医疗数据共享等敏感场景中展现出显著价值;然而,现有方案在系统适用性与安全模型方面仍存在多重关键性问题亟待解决.首先,在功能扩展性层面,主流去中心化 IPFE 方案仅支持身份维度的单向隐私计算(即一对一模式),缺乏一对多应用场景中细粒度访问控制需求的适应性;其次,在属性容量约束方面,经典 AB-IPFE 方案受限于系统公钥的代数结构特点,其属性空间维度存在理论上限,导致实际部署时需通过反复执行系统初始化来扩展属性集规模,这种重复初始化操作将引发多项式级的计算资源损耗.更值得注意的是,现有架构普遍依赖中心化权威机构进行全局公私钥的生成,致使系统面临单点失效风险,且多数素数阶群方案采用非静态困难性假设作为安全性基础,相应提议的安全性强度取决于敌手的询问次数,难以抵抗具备多项式计算能力的自适应敌手的攻击.

为突破上述技术瓶颈, 本文致力于构建基于静态困难性假设的去中心化 AB-IPFE 机制, 主要贡献如下.

- (1) 动态的细粒度访问控制架构. 通过引入线性秘密共享 (linear secret sharing scheme, LSSS) 策略重构密文解密权限管理体系, 实现支持多属性协作的灵活访问控制策略, 突破传统方案在细粒度授权机制方面的局限性.
- (2) AB-IPFE 机制的去中心化构造模式. 在素数阶群代数体系下设计新型去中心化 AB-IPFE 协议, 相较于现有方案具有双重优势: 一方面通过优化属性密钥的生成方式实现对大容量属性集的支持, 消除重复系统初始化带来的效率瓶颈; 另一方面采用去中心化的密钥生成框架提升了系统对单点失效与合谋攻击的防御能力.
- (3) 基于静态困难性假设的安全性证明体系. 基于 DBDH 困难性假设构建形式化的安全性证明过程, 通过改进模拟器构造方法实现良好的安全性归约.
- (4) 辅助输入模型下的泄露攻击抵抗. 基于 Goldreich-Levin 定理, 借助向量内积运算结果的均匀随机性对明文信息进行隐藏, 实现对属性私钥的泄露抵抗.

2 基础知识

2.1 访问结构及线性秘密共享

设 $\mathcal{P}$ 表示属性集合, 访问结构 $\mathbb{A}$ 定义为 $\mathcal{P}$ 的所有子集的一个集合, 即 $\mathbb{A} \subseteq 2^{\mathcal{P}}$. 若 $B \in \mathbb{A}$, 表示子集 B 是授权集, 拥有恢复秘密或解密密文的权限. 若 $B \notin \mathbb{A}$, 则表示子集 B 是非授权集, 无法获得秘密信息. 通常, 访问结构具有单调性, 即如果 $B \in \mathbb{A}$ 且 $B \subseteq C$, 那么必有 $C \in \mathbb{A}$.

定义1 (线性秘密共享方案) LSSS 是在一组属性集合 $\mathcal{P}$ 中安全共享秘密 $s \in \mathbb{Z}_q^*$ 的方法. 秘密 s 作为随机向量 $\mathbf{v}$ 在 $\mathbb{Z}_p^*$ 上的线性组合进行共享, 其中 $\mathbf{v} = (s, r_2, \dots, r_n)^\top$ 包括共享秘密 s 和 $n-1$ 个随机值 $r_2, \dots, r_n \in \mathbb{Z}_q^*$. 秘密共享矩阵 $\mathbf{M}$ 有 l 行和 n 列, 每一行 $\mathbf{M}_i$ 与一个属性 $\rho(i)$ 相关联, 其中 $\rho: \{1, 2, \dots, l\} \rightarrow \mathcal{P}$ 是一个单向映射函数. 每个属性的份额 λ_i 计算为 $\lambda_i = \mathbf{M}_i \cdot \mathbf{v}$, 其中 $\mathbf{M}_i$ 是 $\mathbf{M}$ 的第 i 行. 第 i 个份额 λ_i 被分配给属性 $\rho(i)$. 存在常数 $\{c_i \in \mathbb{Z}_q^* \mid i \in I_A\}$, 其中 I_A 索引 $\mathbf{M}$ 的行. 任何授权集 $A \subseteq \mathcal{P}$ 中的参与者子集都可以通过 $s = \sum_{i \in I_A} c_i \cdot \lambda_i$ 重构秘密 s . 任何未授权集 $U \not\subseteq \mathcal{P}$ 中的参与者子集都无法获得关于秘密 s 的任何相关信息.

2.2 困难性假设

设 G 和 G_T 是两个阶为素数 q 的乘法循环群, g 是 G 中的一个生成元, $e: G \times G \rightarrow G_T$ 是一个双线性映射.

定义2 (决策的双线性 Diffie-Hellman 假设^[22]) 给定元组 (g, g^a, g^b, g^c, T) , 其中 $a, b, c \in \mathbb{Z}_q^*$ 是随机选取的, 判定 T 是否等于 $e(g, g)^{abc}$. 任意概率多项式时间算法 $\mathcal{A}$ 解决该判定性问题的优势定义为 $\text{Adv}_{\mathcal{A}}^{\text{DBDH}} = |\Pr[\mathcal{A}(g, g^a, g^b, g^c, e(g, g)^{abc}) = 1] - \Pr[\mathcal{A}(g, g^a, g^b, g^c, R) = 1]| \leq \varepsilon$, 其中 R 是从 G_T 中均匀随机选取的元素, ε 是可忽略的值.

2.3 Goldreich-Levin 定理

定理1 (Goldreich-Levin 定理^[20]) 令 q 是一个大素数, H 是有限域 $GF(q)$ 的一个子集. 定义单向映射函数 $f: H^\ell \rightarrow \{0, 1\}^*$. 任意选取 $\mathbf{z} \leftarrow H^\ell$ 和 $\mathbf{s} \leftarrow GF(q)^\ell$, 计算 $y = f(\mathbf{z})$. 如果存在概率多项式时间 t 内的区分器 D 使得 $|\Pr[D(y, \mathbf{s}, \langle \mathbf{z}, \mathbf{s} \rangle) = 1] - \Pr[D(y, \mathbf{s}, u) = 1]| = \varepsilon$ 成立 (其中 $u \leftarrow_R GF(q)$ 和 $\ell = (3 \log q)^{1/\varepsilon}$), 则存在可逆器 U 在时间 $t' = t \cdot \text{poly}(\ell, |H|, 1/\varepsilon)$ 内使得 $\Pr[\mathbf{z} \leftarrow H^\ell, y \leftarrow f(\mathbf{z}): U(y) = \mathbf{z}] \geq \frac{\varepsilon^3}{512 \cdot \ell \cdot q^2}$ 成立.

3 去中心化的属性基内积函数加密机制

本节将介绍去中心化 AB-IPFE 机制的形式化定义, 正确性及其安全模型. 特别地, 为简化方案描述, 本文假设每个权威仅控制一个属性.

3.1 形式化定义

去中心化的 AB-IPFE 机制由下述 5 个算法组成, 具体叙述如下.

(1) **全局初始化 (GlobalSetup)**. 输入安全参数 κ 和支持的最大访问矩阵宽度 w_{max} , 系统运行初始化算法定义属性空间 $\mathcal{AU}$ 和全局身份 $\mathcal{GID}$, 生成双线性群 $\mathbb{G} = \langle q, G, G_T, g, e \rangle$ 以及哈希函数 $H : \mathcal{GID} \times [w_{max}] \rightarrow G$. 最终, 全局初始化算法公开全局参数 $GP = (\mathbb{G}, H, \mathcal{GID})$. 该算法过程表示为 $GP \leftarrow \text{GSetup}(1^\kappa, w_{max})$. 该算法主要用于代数结构的初始化, 并不生成全局公私钥.

(2) **权威初始化 (AuthoritySetup)**. 输入全局参数 GP 和属性权威 $u \in \mathcal{AU}$, 初始化算法输出权威公钥 PK_u 与私钥 MSK_u . 该算法过程表示为 $(PK_u, MSK_u) \leftarrow \text{ASetup}(GP, u)$.

(3) **密钥生成 (KeyGen)**. 密钥生成算法以全局参数 GP 、用户的全局身份 $GID \in \mathcal{GID}$ 、权威私钥 MSK_u 和谓词向量 $\mathbf{f}$ 为输入, 输出用户私钥 $SK_{GID,u}$. 该算法过程表示为 $SK_{GID,u} \leftarrow \text{KeyGen}(GP, GID, MSK_u, \mathbf{f})$.

(4) **加密 (Encryption)**. 加密算法以全局参数 GP 、权威公钥 $\{PK_u\}$ 、消息向量 $\mathbf{d} \in \mathcal{M}$ (其中 $\mathcal{M}$ 表示消息向量空间) 和访问策略 $\psi = (\mathbf{M}, \rho)$ 为输入, 输出加密密文 CT . 该算法过程表示为 $CT \leftarrow \text{Enc}(GP, \{PK_u\}, \mathbf{d}, \psi)$.

(5) **解密 (Decryption)**. 解密算法以密文 CT 、全局参数 GP 和用户私钥 $\{SK_{GID,u}\}$ 为输入, 输出解密结果 $\langle \mathbf{f}, \mathbf{d} \rangle$. 该算法过程表示为 $\langle \mathbf{f}, \mathbf{d} \rangle \leftarrow \text{Dec}(GP, CT, \{SK_{GID,u}\})$.

3.2 正确性

去中心化 AB-IPFE 机制 $\Pi = (\text{GSetup}, \text{ASetup}, \text{KeyGen}, \text{Enc}, \text{Dec})$ 对于每个安全参数 $\kappa \in \mathbb{N}$ 和消息向量 $\mathbf{d} \in \mathcal{M}$, 存在一个可忽略函数 $\text{negl}(\kappa)$, 使得下述关系式成立:

$$\Pr \left[\langle \mathbf{f}, \mathbf{d} \rangle = t \mid \begin{array}{l} GP \leftarrow \text{GSetup}(1^\kappa, w_{max}); \\ (PK_u, MSK_u) \leftarrow \text{ASetup}(GP, u); \\ SK_{GID,u} \leftarrow \text{KeyGen}(GP, GID, MSK_u, \mathbf{f}); \\ CT \leftarrow \text{Enc}(GP, \{PK_u\}, \mathbf{d}, \psi); \\ t \leftarrow \text{Dec}(GP, CT, \{SK_{GID,u}\}); \end{array} \right] \geq 1 - \text{negl}(\kappa),$$

其中对于用户密钥组成的属性/权威集合 Attr , 有 $\psi(\text{Attr}) = 1$ 成立, 这表明用户的私钥中包含的属性集合 Attr 满足密文提供的访问策略 $\psi = (\mathbf{M}, \rho)$.

3.3 安全模型

设 $\mathcal{A}$ 为一个概率多项式时间敌手. 为证明去中心化 AB-IPFE 机制在选择明文攻击 (chosen-plaintext attacks, CPA) 下满足语义安全性, 本文构建了由挑战者 $\mathcal{C}$ 与敌手 $\mathcal{A}$ 参与的交互式游戏.

全局初始化. 给定安全参数 κ 和支持的 LSSS 访问策略的最大矩阵宽度 w_{max} , 挑战者 $\mathcal{C}$ 定义属性空间 $\mathcal{AU}$, 并执行全局初始化算法 $\text{GSetup}(1^\kappa, w_{max})$ 生成全局参数 GP , 并将其返回给敌手 $\mathcal{A}$. 最后, 敌手 $\mathcal{A}$ 输出挑战访问策略 $\psi^* = (\mathbf{M}^*, \rho^*)$ 给挑战者 $\mathcal{C}$.

权威收买. 在全局初始化后, 敌手 $\mathcal{A}$ 可选择收买若干权威 u , 并向挑战者 $\mathcal{C}$ 公开这些被收买权威 u 的公钥 PK_u . 该询问模拟了现实环境中敌手能够通过收买获得个别属性权威 u 的私钥信息. 即在本文的安全模型中赋予了敌手 $\mathcal{A}$ 获知部分属性权威私钥的能力.

阶段一. 在此阶段, 敌手 $\mathcal{A}$ 适应性地进行下述询问, 但敌手 $\mathcal{A}$ 询问过私钥的属性和被收买的权威/属性 u 组成的属性集合不能满足挑战访问策略 ψ^* .

(1) **权威公钥提取询问.** 敌手 $\mathcal{A}$ 提交对于权威 u 的公钥提取询问, 挑战者 $\mathcal{C}$ 执行权威初始化算法 $\text{ASetup}(GP, u)$ 生成权威公钥 PK_u 并将相应的权威公钥 PK_u 返回给敌手 $\mathcal{A}$.

(2) **私钥提取询问.** 当敌手 $\mathcal{A}$ 提交一个对于 (GID, u) 的私钥提取询问, 挑战者 $\mathcal{C}$ 执行密钥生成算法生成私钥 $SK_{GID,u}$ 并将相应的私钥 $SK_{GID,u}$ 返回给敌手 $\mathcal{A}$.

挑战阶段. 敌手 $\mathcal{A}$ 向 $\mathcal{C}$ 提交两个等长的消息向量 $\{d_0, d_1\}$. 挑战者 $\mathcal{C}$ 随机选择 $\gamma \leftarrow \{0, 1\}$ 后, 执行加密算法将权威公钥 $\{PK_u\}$ 、消息向量 d_γ 和挑战访问策略 ψ^* 生成挑战密文 CT^* . 最后, $\mathcal{C}$ 将挑战密文 CT^* 返回给敌手 $\mathcal{A}$.

阶段二. 该阶段与阶段一类似, 敌手 $\mathcal{A}$ 适应性地进行权威公钥和私钥提取询问, 但敌手 $\mathcal{A}$ 询问过私钥的属性和被收买的权威属性组成的属性集合不能满足挑战访问策略 ψ^* .

猜测阶段. 敌手 $\mathcal{A}$ 输出对随机数 γ 的猜测 γ' . 若 $\gamma' = \gamma$, 则挑战者 $\mathcal{C}$ 输出 “1”, 表示敌手 $\mathcal{A}$ 在该游戏中获胜.

在上述游戏中, 敌手 $\mathcal{A}$ 获胜的优势定义为 $\text{Adv}_{\mathcal{A}}(\kappa) = |\Pr[\mathcal{A} \text{ wins}] - \frac{1}{2}|$, 其中敌手已掌握私钥的属性所组成的属性集合不能满足挑战访问策略 $\psi^* = (\mathbf{M}^*, \rho^*)$.

定义3 (去中心化 AB-IPFE 机制的 CPA 安全性) 对于任意的概率多项式时间敌手 $\mathcal{A}$, 若其在上述游戏中获胜的优势是可忽略的, 即有 $\text{Adv}_{\mathcal{A}}(\kappa) \leq \text{negl}(\kappa)$, 那么相应的去中心化 AB-IPFE 机制具有 CPA 安全性.

4 本文构造

本节将详细介绍本文提出的去中心化 AB-IPFE 机制的具体构造, 并详细分析该方案的正确性及安全性证明过程.

4.1 具体构造

(1) 全局初始化 (GlobalSetup). 给定安全参数 κ 和 AB-IPFE 机制所能支持的 LSSS 矩阵的最大宽度 $w_{max} = w_{max}(\kappa)$. 全局初始化算法定义属性空间 $\mathcal{AU}$, 并运行双线性群生成器 $\mathcal{G}(\kappa)$ 生成阶为素数 q 的双线性群 $\mathbb{G} = \langle g, G, G_T, g, e \rangle$, 其中 g 是群 G 的生成元. 此外, 算法设定全局身份空间 $\mathcal{GID}$ 和哈希函数 $H: \mathcal{GID} \times [w_{max}] \rightarrow G$ 将字符串 $(GID, i) \in \mathcal{GID} \times [w_{max}]$ 映射到 G 中的元素. 最后, 全局初始化算法公开全局参数 $GP = (\mathbb{G}, H, \mathcal{GID})$.

特别地, GlobalSetup 仅生成了 AB-IPFE 机制中各算法共用的代数结构、哈希函数等公开参数.

(2) 权威初始化 (AuthoritySetup). 给定全局参数 GP 和一个权威选定属性 $u \in \mathcal{AU}$, 权威初始化算法选取随机数 $\alpha_u, y_{u,2}, \dots, y_{u,w_{max}} \leftarrow \mathbb{Z}_q^{w_{max}}$ 并输出其相关的权威公钥和私钥:

$$PK_u = (e(g, g)^{\alpha_u}, g^{y_{u,2}}, \dots, g^{y_{u,w_{max}}}), \quad MSK_u = (\alpha_u, y_{u,2}, \dots, y_{u,w_{max}}).$$

特别地, 去中心化的特点主要体现在各属性权威分别独立完成各自的公私钥生成, 不再需要中心属性权威为系统生成全局的主公钥和主私钥.

(3) 属性密钥生成 (KeyGen). 密钥生成算法以全局参数 GP 、哈希函数 H 、用户的全局身份 $GID \in \mathcal{GID}$ 、权威私钥 MSK_u 和谓词向量 $\mathbf{f} = (f_1, f_2, \dots, f_n)$ 作为输入. 权威 u 为用户 GID 生成其私钥 $SK_{GID,u} = g^{\alpha_u \sum_{k \in [n]} f_k \prod_{j=2}^{w_{max}} H(GID \parallel j)^{y_{u,j}}}$. 最后, 输出用户私钥 $SK_{GID,u}$. 其中, f_k 和 n 分别是向量 $\mathbf{f}$ 的分量和长度. 属性权威为身份是 GID 的用户生成关于属性 u 的属性私钥.

(4) 加密 (Encryption). 加密算法以全局参数 GP 、哈希函数 H 、消息向量 $\mathbf{d} = (d_1, d_2, \dots, d_n) \in \mathcal{M}$ 、访问策略 $\psi = (\mathbf{M}, \rho)$ 和权威公钥 $\{PK_{\rho(i)}\}$ 作为算法的输入, 其中 $\mathbf{M}_{\ell \times w_{max}} = (\mathbf{M}_1, \dots, \mathbf{M}_\ell)^\top \in \mathbb{Z}_q^{\ell \times w_{max}}$ 且 $\rho: [\ell] \rightarrow u \in \mathcal{AU}$. 映射函数 ρ 将矩阵 $\mathbf{M}$ 的行与权威关联, 即一个权威/属性最多与矩阵 $\mathbf{M}$ 的一行关联. 算法选择随机数 $r_1, \dots, r_\ell \leftarrow \mathbb{Z}_q^\ell$; $\mathbf{v} = (z, v_2, \dots, v_{w_{max}}) \leftarrow \mathbb{Z}_q^{w_{max}}$, 以及 $\mathbf{x} = (x_2, \dots, x_{w_{max}}) \leftarrow \mathbb{Z}_q^{w_{max}-1}$ (其中向量 $\mathbf{v}$ 的作用是对加密随机数 z 进行秘密共享, 以确保属性集满足访问策略的用户能够解密. 向量 $\mathbf{x}$ 被用于抵抗合谋攻击, 仅当相关属性由同一用户 GID 所持有时, 才能将秘密共享的 “0” 恢复实现对盲化项的消除, 确保解密操作的正确性), 并计算

$$\begin{aligned} \forall i \in [n]: C_{0,i} &= e(g, g)^z e(g, g)^{d_i}, \\ \forall i \in [\ell]: C_{1,i} &= e(g, g)^{\mathbf{M}_i \cdot \mathbf{v}} e(g, g)^{\alpha_{\rho(i)} r_i}, \quad C_{2,i} = g^{r_i}, \end{aligned}$$

$$\forall i \in [\ell], j \in \{2, \dots, w_{max}\} : C_{3,i,j} = g^{y_{\rho(i),j} r_i} g^{M_{i,j} x_j}.$$

最后, 算法输出密文 $CT = ((M, \rho), \{C_{0,i}\}_{i \in [n]}, \{C_{1,i}\}_{i \in [\ell]}, \{C_{2,i}\}_{i \in [\ell]}, \{C_{3,i,j}\}_{i \in [\ell], j \in \{2, \dots, w_{max}\}})$.

(5) 解密 (Decryption). 解密算法以全局参数 GP 、访问策略 $\psi = (M, \rho)$ 、用户私钥 $\{SK_{GID,u}\}$ 以及密文 CT 作为算法的输入, 其中访问策略 $\psi = (M, \rho)$ 中 $M \in \mathbb{Z}_q^{\ell \times w_{max}}$ 且 $\rho: [\ell] \rightarrow u$ 为单向映射, 以及对应于矩阵 M 的行的子集的密钥 $\{SK_{GID,\rho(i)}\}_{i \in I}$ 与索引 $I \subseteq [\ell]$. 如果 $(1, 0, \dots, 0)$ 不在这些行的张成空间中, 则解密失败, 输出 $\perp$. 否则, 存在一个向量 $\{\theta_i\}_{i \in I}$ 使得 $(1, 0, \dots, 0) = \sum_{i \in I} \theta_i M_i$. 令 $\rho(i)$ 为 $i \in I$ 的对应权威, 计算并输出 $e(g, g)^{\langle d, f \rangle}$ 如下:

$$e(g, g)^{\langle d, f \rangle} = \prod_{k \in [n]} C_{0,k}^{f_k} / \prod_{i \in I} \left[\frac{\prod_{k \in [n]} C_{1,i}^{f_k} \cdot \prod_{j=2}^{w_{max}} e(H(GID \parallel j), C_{3,i,j})}{e(SK_{GID,\rho(i)}, C_{2,i})} \right]^{\theta_i}.$$

4.2 正确性

若用户密钥所拥有的属性集合 $Attr$ 满足相应密文的访问策略 ψ , 则可按如下步骤正确解密出目标密文:

$$\begin{aligned} & \prod_{i \in I} \left[\frac{\prod_{k \in [n]} C_{1,i}^{f_k} \cdot \prod_{j=2}^{w_{max}} e(H(GID \parallel j), C_{3,i,j})}{e(SK_{GID,\rho(i)}, C_{2,i})} \right]^{\theta_i} \\ &= \prod_{i \in I} \left[\frac{e(g, g)^{\sum_{k \in [n]} f_k \cdot M_i \cdot v} \cancel{e(g, g)^{\sum_{k \in [n]} f_k \cdot \rho(i) \cdot r_i}} \prod_{j=2}^{w_{max}} e(H(GID \parallel j), g^{y_{\rho(i),j} r_i}) e(H(GID \parallel j), g^{M_{i,j} x_j})}{\cancel{e(g, g)^{\sum_{k \in [n]} f_k \cdot \rho(i) \cdot r_i}} \prod_{j=2}^{w_{max}} e(H(GID \parallel j), g^{y_{\rho(i),j} r_i})} \right]^{\theta_i} \\ &= \prod_{i \in I} \left[e(g, g)^{\sum_{k \in [n]} f_k \cdot M_i \cdot v} \prod_{j=2}^{w_{max}} e(H(GID \parallel j), g)^{M_{i,j} x_j} \right]^{\theta_i} \\ &= \prod_{i \in I} e(g, g)^{\sum_{k \in [n]} f_k \cdot \theta_i M_i \cdot v} \prod_{j=2}^{w_{max}} \prod_{i \in I} e(H(GID \parallel j), g)^{\theta_i M_{i,j} x_j} \rightarrow \text{只有每一个 } H(GID \parallel j) \text{ 都相同时才能通} \\ & \quad \text{过下述步骤解密, 以此实现抵抗合谋攻击的目的.} \\ &= e(g, g)^{\sum_{i \in I} \sum_{k \in [n]} f_k \cdot \theta_i M_i \cdot v} \prod_{j=2}^{w_{max}} e(H(GID \parallel j), g)^{\sum_{i \in I} \theta_i M_{i,j} x_j} = e(g, g)^{z \sum_{k \in [n]} f_k} \\ & \quad \prod_{k \in [n]} C_{0,k}^{f_k} = \prod_{k \in [n]} \left(e(g, g)^z e(g, g)^{d_k} \right)^{f_k} = e(g, g)^{\langle d, f \rangle} \prod_{k \in [n]} e(g, g)^{z f_k} \\ &= e(g, g)^{\langle d, f \rangle} e(g, g)^{z \sum_{k \in [n]} f_k}, \end{aligned}$$

其中 $\sum_{i \in I} \theta_i M_i v = s$ 和 $\sum_{i \in I} \theta_i M_{i,j} x_j = 0$.

根据上述等式, 可知

$$e(g, g)^{\langle d, f \rangle} = \prod_{k \in [n]} C_{0,k}^{f_k} / \prod_{i \in I} \left[\frac{\prod_{k \in [n]} C_{1,i}^{f_k} \cdot \prod_{j=2}^{w_{max}} e(H(GID \parallel j), C_{3,i,j})}{e(SK_{GID,\rho(i)}, C_{2,i})} \right]^{\theta_i}.$$

4.3 安全性证明

定理2 如果存在敌手 $\mathcal{A}$ 能够以不可忽略的优势 ε 在多项式时间内攻破上述去中心化 AB-IPFE 机制的 CPA 安全性, 那么存在挑战者 $\mathcal{C}$ 能以不可忽略的优势 ε 解决 DBDH 困难性问题.

由于篇幅所限, 此处不再讨论上述定理的具体证明过程, 详见完整版¹⁾.

1) <https://github.com/Honey291/static>, 若此链接由于不可抗拒的原因导致无法访问, 请联系 zyw@snnu.edu.cn.

表 1 与现有相关方案的性能对比.

Table 1 Performance comparison with existing related schemes.

Scheme	Large universe	Collusion resistance	Fine-grained access control	Decentralization	Security assumption
Tseng et al. [9]	–	✓	×	✓	Non-static assumption
Han et al. [10]	–	✓	×	✓	Static assumption
Yang et al. [14]	×	×	✓	×	Non-static assumption
Our construction	✓	✓	✓	✓	Static assumption

表 2 与相应相关方案的计算效率对比.

Table 2 Computational efficiency comparison with existing related schemes.

Scheme	GlobalSetup	AuthoritySetup	KeyGen	Encryption	Decryption
Tseng et al. [9]	$\mathcal{O}(1)$	$\mathcal{O}(n\ell) G $	$\mathcal{O}(n\ell) G $	$\mathcal{O}(n) G_T + \mathcal{O}(n^2\ell) G $	$\mathcal{O}(n\ell) G_T $
Han et al. [10]	$\mathcal{O}(1)$	$\mathcal{O}(n\ell) G_T $	$\mathcal{O}(n\ell) G $	$\mathcal{O}(n\ell) G_T $	$\mathcal{O}(n\ell) G_T $
Yang et al. [14]	–	$\mathcal{O}(n^2) G $	$\mathcal{O}(3n + 2\ell) G $	$\mathcal{O}(n\ell) G_T + \mathcal{O}(2n) G $	$\mathcal{O}(n + 6\ell) G_T $
Our construction	$\mathcal{O}(1)$	$\mathcal{O}(\ell) G $	$\mathcal{O}(n + \ell) G $	$\mathcal{O}(n) G_T + \mathcal{O}(n\ell) G $	$\mathcal{O}(n\ell) G_T $

4.4 性能比较及仿真模拟

本节将所提出的去中心化 AB-IPFE 机制与其他相关方案^[9,10,14]进行功能和计算效率的详细对比,以证明本文方案能够更加适应实际场景的应用需求.

(1) **性能分析.** 本节将本文方案与现有构造^[9,10,14]所具备的安全属性及特点进行比较,比较结果如表 1 所示,其中,“✓”表示相关构造具备相应的性能;“×”表示相关构造不具备相应的性能;“–”表示相关构造是基于身份基密码体制的,不存在属性集的概念. 由表 1 可知, Tseng 等^[9]与 Han 等^[10]的方案虽实现去中心化架构,却缺乏细粒度访问控制机制,存在数据越权访问风险及安全隐患. Yang 等^[14]的方案虽引入细粒度访问控制,但未兼容大属性集与去中心化特性,属性空间扩展时需反复执行系统重初始化流程,引发显著计算开销. 此外, Tseng 等^[9]与 Yang 等^[14]的方案均基于非静态安全模型,难以抵御敌手提出的与询问次数相关的攻击. 相较而言,本文构建的去中心化 AB-IPFE 机制在支持大属性集扩展与分布式管理的同时,基于静态安全性假设实现了强安全性防护.

(2) **计算效率的理论分析.** 本文将所提出的方案与现有研究方案进行计算效率的理论分析对比,结果如表 2 所示,其中 n 是参与运算的向量长度, ℓ 是权威的数量, $|G|$ 是指其算法在群 G 上的运算,而 $|G_T|$ 指的是在群 G_T 上的运算. 具体讲,在 AB-IPFE 机制中 ℓ 是访问策略中访问矩阵的行数,在去中心化的身份基 IPFE 机制中 ℓ 是分布式体系中密钥生成中心的数量;“–”表示相关构造是基于中心权威的密码体制,将其初始化算法归入权威初始化阶段,因此不包含全局初始化算法. 由于全局初始化仅建立了群结构 $\mathbb{G}$ 和哈希函数 H 的代数描述,无实际运算开销,故时间复杂度为 $\mathcal{O}(1)$. 由表 2 可知,本文方案在权威初始化和密钥生成阶段的时间复杂度分别为 $\mathcal{O}(\ell)|G|$ 和 $\mathcal{O}(n + \ell)|G|$ 远优于 Tseng 等^[9]和 Han 等^[10]所提出方案的 $\mathcal{O}(n^2)|G|$, $\mathcal{O}(n\ell)|G|$ 以及 $\mathcal{O}(n\ell)|G_T|$. 这种线性时间复杂度特性有效保障了系统在动态扩展权威节点及用户规模时的操作响应性,为构建可伸缩的大规模分布式权威体系提供了理论支撑. 在加密与解密效率维度,本文方案与现有 IPFE 机制保持一致. 需特别指出, Yang 等的方案^[14]虽在小规模属性集场景下具有较优的计算效率,但其中心化架构隐含权威单点失效的风险,且在属性集维度扩张时需频繁执行系统初始化操作,导致计算开销呈超线性增长. 相较之下,本文方案在维持优于现有去中心化方案^[9,10]计算效率优势的同时,通过创新的细粒度访问控制机制,实现了安全性与可扩展性的均衡,使其在大规模分布式应用场景中更具实际部署价值.

(3) **计算效率的模拟比较.** 为验证本文方案在受限计算环境下的实际效能,基于 PBC (pairing-based cryptography) 密码库构建仿真实验平台,在树莓派开发板上实现原型方案的模拟验证. 通过将方案向量维度初始设定为 10,然后将属性空间规模从 10 动态扩展到 100,分别评估算法各阶段计算效能随属性规模增长的变化规律.

```

SUCCESS - Decryption matches expected result!
=== Number of attributes: 60, Vector size: 10 ===
Global Setup Time: 0.00 ms
Authority Setup Time (for 60 authorities): 32.00 ms
Key Generation Time (for 60 keys): 111.00 ms
Encryption Time: 81.00 ms
Decryption Time: 113.00 ms
Total Time: 337.00 ms

SUCCESS - Decryption matches expected result!
=== Number of attributes: 70, Vector size: 10 ===
Global Setup Time: 0.00 ms
Authority Setup Time (for 70 authorities): 47.00 ms
Key Generation Time (for 70 keys): 131.00 ms
Encryption Time: 81.00 ms
Decryption Time: 128.00 ms
Total Time: 387.00 ms

```

图1 资源受限环境中 PBC 库代码的执行结果.

Figure 1 Execution results on PBC library in resource-constrained environments.

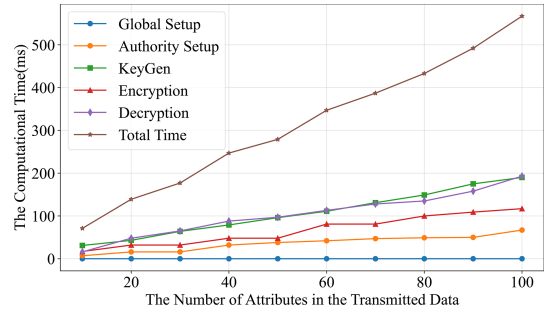


图2 (网络版彩图) 资源受限环境下属性数量与方案计算开销间关系.

Figure 2 (Color online) Time cost of the scheme with increasing attributes in resource-constrained environments.

实验结果表明 (如图 1 所示), 全局初始化阶段是亚毫秒级开销 (<0.005 ms), 因低于数据采集精度阈值而被忽略. 值得注意的是, 当属性空间扩展至 70 维时, 本文 AB-IPFE 机制的整体运算耗时仍可维持在 400 ms 量级内, 展现出良好的可扩展性. 为清晰呈现不同算法的运算效率趋势, 本文采用线性图示方法展示属性规模扩张对计算效率的影响. 如图 2 所示, 当属性基数扩展至 100 维时, 各算法的时间开销均严格限定于 200 ms 内. 值得注意的是, 实验数据显示各算法时间开销的增长率与属性集的线性扩张保持同步平稳态势. 这一量化结果不仅验证了方案设计的计算优越性, 更通过严格的时间复杂度变化图像, 为系统在现实部署环境中的可行性提供了实证依据. 为方便实验结果复现, 本文已将所使用的代码公开²⁾.

5 抗泄露的去中心化 AB-IPFE 机制

随着现代计算环境中侧信道攻击、冷启动攻击等非常规攻击方式的广泛存在, 传统基于理想黑箱假设的密码学安全模型已呈现出显著局限性. 上述理论研究与实际需求间的不一致性, 推动了泄露弹性密码学研究范式的兴起, 其核心目标在于构建能够抵御部分密钥信息泄露的密码协议体系. 上述去中心化 AB-IPFE 机制不具备抵抗泄露攻击的能力, 那么在存在泄露的环境中上述提议无法满足其所声称的安全性. 为进一步增强 AB-IPFE 机制的实用性, 本节将在上述构造的基础上提出抗泄露的去中心化 AB-IPFE 机制的具体构造.

5.1 具体构造

由上文分析可知, Goldreich-Levin 定理要求加密算法需涉及两个向量的内积求值运算才能提供辅助输入模型下的泄露容忍性, 本节在上文去中心化 AB-IPFE 机制的基础上, 修改加密算法的结构, 使其满足上述要求. 权威初始化 (AuthoritySetup) 和属性密钥生成 (KeyGen) 算法与第 4 节的 AB-IPFE 机制相同, 此处不再赘述.

(1) **全局初始化 (GlobalSetup).** 给定安全参数 κ , AB-IPFE 机制所能支持的 LSSS 矩阵的最大宽度 $w_{max} = w_{max}(\kappa)$ 以及最大行数 $\ell = (3 \log q)^{1/\varepsilon}$ (ℓ 的取值由定理 1 确定). 全局初始化算法定义属性空间 $\mathcal{AU}$, 并运行双线性群生成器 $\mathcal{G}(\kappa)$ 生成阶为素数 q 的双线性群 $\mathbb{G} = \langle g, G, G_T, g, e \rangle$, 其中 g 是群 G 的生成元. 此外, 算法设定全局身份空间 $\mathcal{GID}$ 和哈希函数 $H: \mathcal{GID} \times [w_{max}] \rightarrow G$ 将字符串 $(GID, i) \in \mathcal{GID} \times [w_{max}]$ 映射到 G 中的元素. 最后, 全局初始化算法公开全局参数 $GP = (\mathbb{G}, H, \mathcal{GID})$.

(2) **权威初始化 (AuthoritySetup).** 给定全局参数 GP 和一个权威选定属性 $u \in \mathcal{AU}$, 权威初始化算法选取随机数 $\alpha_u, y_{u,2}, y_{u,3}, \dots, y_{u,w_{max}} \leftarrow \mathbb{Z}_q^{w_{max}}$ 并输出其相关的权威公钥和私钥:

$$PK_u = (g^{\alpha_u}, g^{y_{u,2}}, \dots, g^{y_{u,w_{max}}}), \quad MSK_u = (\alpha_u, y_{u,2}, \dots, y_{u,w_{max}}).$$

2) <https://github.com/Honeyme291/DE-AB-IPFE>.

(2) 加密 (Encryption). 加密算法以全局参数 GP 、哈希函数 H 、消息向量 $\mathbf{d} \in \mathcal{M}$ 、LSSS 访问策略 $\psi = (\mathbf{M}, \rho)$ 和权威公钥 $\{PK_{\rho(i)}\}$ 作为算法的输入, 其中 $\mathbf{M}_{\ell \times w_{max}} = (\mathbf{M}_1, \dots, \mathbf{M}_\ell)^\top \in \mathbb{Z}_q^{\ell \times w_{max}}$ 且 $\rho: [\ell] \rightarrow u \in \mathcal{AU}$. 映射函数 ρ 将矩阵 $\mathbf{M}$ 的行与权威关联, 即一个权威/属性最多与矩阵 $\mathbf{M}$ 的一行关联. 算法选择随机数 $z \leftarrow \mathbb{Z}_q^*$, $r_1, \dots, r_\ell \leftarrow \mathbb{Z}_q^\ell$, $s_1, \dots, s_\ell \leftarrow \mathbb{Z}_q^\ell$, $v_2, \dots, v_{w_{max}} \leftarrow \mathbb{Z}_q^{w_{max}-1}$ 和 $\mathbf{x} = (x_2, \dots, x_{w_{max}}) \leftarrow \mathbb{Z}_q^{w_{max}-1}$, 设 $\mathbf{v} = (z, v_2, \dots, v_{w_{max}})$, 并计算

$$\begin{aligned} \forall i \in [n]: C_{0,i} &= e(g, g)^{d_i} \prod_{j=1}^{\ell} e(g, g)^{\alpha_{\rho(j)} s_j z}, \\ \forall i \in [\ell]: C_{1,i} &= e(g, g)^{\alpha_{\rho(i)} s_i \mathbf{M}_i \cdot \mathbf{v}} e(g, g)^{\alpha_{\rho(i)} r_i}, \quad C_{2,i} = g^{r_i}, \\ \forall i \in [\ell], j \in \{2, \dots, w_{max}\}: C_{3,i,j} &= g^{y_{\rho(i),j} r_i} g^{M_{i,j} x_j}. \end{aligned}$$

最后, 算法输出密文 $CT = ((\mathbf{M}, \rho), \{C_{0,i}\}_{i \in [n]}, \{C_{1,i}\}_{i \in [\ell]}, \{C_{2,i}\}_{i \in [\ell]}, \{C_{3,i,j}\}_{i \in [\ell], j \in \{2, \dots, w_{max}\}})$.

(3) 解密 (Decryption). 解密算法以全局参数 GP 、访问策略 $\psi = (\mathbf{M}, \rho)$ 、用户私钥 $\{SK_{GID,u}\}$ 以及密文 CT 作为算法的输入, 其中访问策略 $\psi = (\mathbf{M}, \rho)$ 中 $\mathbf{M} \in \mathbb{Z}_q^{\ell \times w_{max}}$ 且 $\rho: [\ell] \rightarrow u$ 为单向映射, 以及对应于矩阵 $\mathbf{M}$ 的行的子集的密钥 $\{SK_{GID,\rho(i)}\}_{i \in I}$ 与索引 $I \subseteq [\ell]$. 如果 $(1, 0, \dots, 0)$ 不在这些行的张成空间中, 则解密失败, 输出 $\perp$. 否则, 存在一个向量 $\{\theta_i\}_{i \in I}$ 使得 $(1, 0, \dots, 0) = \sum_{i \in I} \theta_i \mathbf{M}_i$. 令 $\rho(i)$ 为 $i \in I$ 的对应权威, 通过下述计算输出 $e(g, g)^{\langle \mathbf{d}, \mathbf{f} \rangle}$:

$$e(g, g)^{\langle \mathbf{d}, \mathbf{f} \rangle} = \prod_{k \in [n]} C_{0,k}^{f_k} / \prod_{i \in I} \left[\frac{\prod_{k \in [n]} C_{1,i}^{f_k} \cdot \prod_{j=2}^{w_{max}} e(H(GID \parallel j), C_{3,i,j})}{e(SK_{GID,\rho(i)}, C_{2,i})} \right]^{\theta_i}.$$

5.2 正确性

若用户密钥所对应的属性集合满足相关密文的访问策略, 则可按如下步骤解密出正确的内积值:

$$\begin{aligned} & \prod_{i \in I} \left[\frac{\prod_{k \in [n]} C_{1,i}^{f_k} \cdot \prod_{j=2}^{w_{max}} e(H(GID \parallel j), C_{3,i,j})}{e(SK_{GID,\rho(i)}, C_{2,i})} \right]^{\theta_i} \\ &= \prod_{i \in I} \left[\frac{e(g, g)^{\alpha_{\rho(i)} s_i \sum_{k \in [n]} f_k \cdot \mathbf{M}_i \cdot \mathbf{v}} \prod_{k \in [n]} \cancel{e(g, g)^{f_k \alpha_{\rho(i)} r_i}} \prod_{j=2}^{w_{max}} \cancel{e(H(GID \parallel j), g^{y_{\rho(i),j} r_i})} e(H(GID \parallel j), g^{M_{i,j} x_j})}{\prod_{k \in [n]} \cancel{e(g, g)^{f_k \alpha_{\rho(i)} r_i}} \prod_{j=2}^{w_{max}} \cancel{e(H(GID \parallel j), g^{y_{\rho(i),j} r_i})}} \right]^{\theta_i} \\ &= \prod_{i \in I} \left[e(g, g)^{\alpha_{\rho(i)} s_i \sum_{k \in [n]} f_k \cdot \mathbf{M}_i \cdot \mathbf{v}} \prod_{j=2}^{w_{max}} e(H(GID \parallel j), g^{M_{i,j} x_j}) \right]^{\theta_i} \\ &= \prod_{i \in I} e(g, g)^{\alpha_{\rho(i)} s_i \sum_{k \in [n]} f_k \cdot \theta_i \mathbf{M}_i \cdot \mathbf{v}} \prod_{j=2}^{w_{max}} \prod_{i \in I} e(H(GID \parallel j), g)^{\theta_i M_{i,j} x_j} \\ &= e(g, g)^{\sum_{i \in I} \sum_{k \in [n]} f_k \alpha_{\rho(i)} s_i \theta_i \mathbf{M}_i \cdot \mathbf{v}} \prod_{j=2}^{w_{max}} e(H(GID \parallel j), g)^{\sum_{i \in I} \theta_i M_{i,j} x_j} = e(g, g)^{z \sum_{i \in I} \alpha_{\rho(i)} s_i \sum_{k \in [n]} f_k} \\ &= \prod_{k \in [n]} C_{0,k}^{f_k} = \prod_{k \in [n]} \left(\prod_{i=1}^{\ell} e(g, g)^{z \alpha_{\rho(i)} s_i} e(g, g)^{d_k} \right)^{f_k} = e(g, g)^{\langle \mathbf{d}, \mathbf{f} \rangle} \prod_{k \in [n]} \prod_{i \in [\ell]} e(g, g)^{z \alpha_{\rho(i)} s_i f_k} \\ &= e(g, g)^{\langle \mathbf{d}, \mathbf{f} \rangle} e(g, g)^{z \sum_{i \in [\ell]} \alpha_{\rho(i)} s_i \sum_{k \in [n]} f_k}, \end{aligned}$$

其中 $\sum_{i \in I} \theta_i \mathbf{M}_i \mathbf{v} = z$ 和 $\sum_{i \in I} \theta_i M_{i,j} x_j = 0$.

综上所述, 本文抗泄露的去中心化 AB-IPFE 机制满足相应的正确性要求.

5.3 安全性

定理3 若存在敌手 $\mathcal{A}$ 能够以不可忽略的优势 ε 在多项式时间内攻破上述抗泄露的去中心化 AB-IPFE 机制的 CPA 安全性, 那么存在挑战者 $\mathcal{C}$ 能以相同的优势 ε 解决 DBDH 困难性问题.

特别地, 上述方案的安全性证明过程与定理 2 的方法基本一致. 对于抗泄露性的证明, 详见文献 [19] 中的相关描述. 鉴于两者证明思路的相似性, 本文对此不再赘述.

5.4 抗泄露性分析

在辅助输入模型下, 攻击者被赋予访问特定单向函数的能力, 这类函数作用于私钥 sk 后生成的函数映射 $f(sk)$ 可用于仿真不同密钥泄露场景. 值得注意的是, 即便此类函数泄露的信息量达到信息论层面的彻底泄漏阈值, 攻击方仍无法从 $f(sk)$ 的任意组合中逆向重构初始私钥 sk . 本小节将重点论证该 AB-IPFE 机制在辅助输入框架下对密钥泄漏的容忍特性. 由定理 1 可知, 对于有限域 $GF(q)$ 上的两个长度为 ℓ 的向量 $\mathbf{z}$ 和 $\mathbf{s}$, 若存在一个多项式时间的区分器 D 能以不可忽略的优势区分以下两个元组: $(y, \mathbf{s}, \langle \mathbf{z}, \mathbf{s} \rangle)$ 与 $(y, \mathbf{s}, u)$, 其中 $y = f(\mathbf{z})$ 且 $u \in GF(q)$ 为均匀随机变量, 那么存在一个可逆器, 能够以不可忽略的优势攻破函数 f 的单向性. 若内积 $\langle \mathbf{z}, \mathbf{s} \rangle$ 和随机值 u 可以在已知 $f(\mathbf{z})$ 和 $\mathbf{s}$ 的前提下被区分, 则意味函数 f 的某些位包含可提取的信息, 从而危及其单向性. 进一步讲, 区分器 D 在已知秘密 $\mathbf{z}$ 的泄露信息 y 和辅助向量 $\mathbf{s}$ 的前提下, 内积值 $\langle \mathbf{z}, \mathbf{s} \rangle$ 对于其依然是均匀随机的. 在上面结论的基础上, 我们可得到一个新结论: 即当辅助向量 $\mathbf{s} = (s_1, s_2, \dots, s_\ell)$ 不公开时, 内积值 $\langle \mathbf{z}, \mathbf{s} \rangle$ 同样与 $GF(q)$ 上的均匀随机值是不可区分的.

对于上述 AB-IPFE 机制, 敌手 $\mathcal{A}$ 可通过概率多项式次的泄露询问获得了关于属性集合 S 对应私钥 $sk = \{SK_{u, GID}\}_{u \in S}$ 的泄露信息 $f(sk)$, 由于 sk 的核心信息是相应属性权威的主私钥 $\alpha = (\alpha_{\rho(1)}, \alpha_{\rho(2)}, \dots, \alpha_{\rho(\ell)})$, 因此敌手 $\mathcal{A}$ 通过泄露询问获得了关于相应属性权威主私钥 α 的泄露信息 $\tilde{f}(\alpha)$. 一般情况下, 泄露信息 $\tilde{f}(\alpha)$ 的获得会增加敌手 $\mathcal{A}$ 在上述游戏中获胜的优势, 除非泄露信息不影响对加密算法中隐藏明文操作所使用元素的随机性, 下面将分析主私钥的泄露信息 $\tilde{f}(\alpha)$ 对用于隐藏明文的元素 $\prod_{j=1}^{\ell} e(g, g)^{\alpha_{\rho(j)} s_j z}$ 的随机性是没有影响的.

密文中对 $\mathbf{d}$ 分量 d_i 的隐藏操作为 $C_{0,i} = e(g, g)^{d_i} \prod_{j=1}^{\ell} e(g, g)^{\alpha_{\rho(j)} s_j z} = e(g, g)^{d_i} e(g, g)^{\sum_{j=1}^{\ell} \alpha_{\rho(j)} s_j z}$, 由定理 1 可知, 已知相应的参数为 $\alpha = (\alpha_{\rho(1)}, \alpha_{\rho(2)}, \dots, \alpha_{\rho(\ell)})$ 和 $\mathbf{s} = (s_1, s_2, \dots, s_\ell)$, 对于任意的概率多项式时间敌手 $\mathcal{A}$, 有关系 $|\Pr[\mathcal{A}(y, \mathbf{s}, \langle \alpha, \mathbf{s} \rangle) = 1] - \Pr[\mathcal{A}(y, \mathbf{s}, u) = 1]| = \varepsilon$ 成立, 其中 ε 是安全参数上可忽略的值, u 是一个均匀随机值和 $y = \tilde{f}(\alpha)$ 为关于相应属性权威主私钥的泄露信息. 根据上述关系可知 $|\Pr[\mathcal{A}(y, \langle \alpha, \mathbf{s} \rangle) = 1] - \Pr[\mathcal{A}(y, u) = 1]| = \varepsilon$ 成立. 由于 $\prod_{i=1}^{\ell} e(g, g)^{\alpha_{\rho(i)} s_i z}$ 与群 G_T 上的均匀随机元素是不可区分的, 因此实现了对消息向量 $\mathbf{d}$ 的完美隐藏, 即使敌手获得了关于相关权威主私钥 $\alpha = (\alpha_{\rho(1)}, \alpha_{\rho(2)}, \dots, \alpha_{\rho(\ell)})$ 的泄露信息 $y = \tilde{f}(\alpha)$, 密文 $C_{0,i} = e(g, g)^{d_i} \prod_{j=1}^{\ell} e(g, g)^{\alpha_{\rho(j)} s_j z}$ 依然保持与均匀随机值间的不可区分性. 因此在辅助输入模型中 (定理 1 保证了对本文去中心化 AB-IPFE 机制私钥泄露信息求逆的困难性), 即使敌手 $\mathcal{A}$ 获得了相应的泄露信息, 上述 AB-IPFE 机制依然保持其所声称的安全性. 因此, 本文构造在辅助输入模型下具有抵抗泄露攻击的能力.

6 结论

本文针对当前 IPFE 机制普遍存在的动态复杂性假设依赖性及数据细粒度访问控制缺失导致的安全性缺陷, 建构了一种基于标准静态假设的去中心化 AB-IPFE 机制. 该机制通过引入 LSSS 访问结构, 构建多属性协同的自适应访问控制策略, 有效克服了传统机制存在的策略表达局限性. 在安全性层面, 本文提议采用分布式权威架构替代中心化信任实体, 既消除了中心化权威节点的单点瓶颈问题, 又具备对合谋攻击的抵抗能力, 强化了数据共享过程的安全性保障. 模拟仿真表明, 本文机制在功能完备性与计算效率层面均展现出较优的适配性, 可为复杂应用场景的部署需求提供可行性解决方案. 为进一步提升本文 AB-IPFE 机制的实用性, 在上述构造的基础上, 本文提出了抗泄露的去中心化 AB-IPFE 机制, 并详细分析了该方案抵抗泄露攻击的基本原理.

本文虽在静态困难性假设下提出了去中心化 AB-IPFE 机制的具体构造, 避免了实际部署时单点失效对系统造成的影响, 但上述研究依然存在下述不足. (1) 仅达到了 CPA 安全性, 即攻击者只能获取密文信息, 而无法

发起适应性的解密查询. 然而, 实际中的攻击者通常具备更强的攻击能力. (2) 存在密钥托管问题, 即系统权威掌握用户的完整属性私钥. 针对上述问题, 本文下一个阶段将在去中心化 AB-IPFE 机制的研究基础上, 研究选择密文攻击安全的去中心化 AB-IPFE 机制.

参考文献

- 1 Wang J, Zhou Y, Zhu Y, et al. Bidirectional identity-based inner-product functional re-encryption in vaccine data sharing. *IEEE Trans Cloud Comput*, 2025, 13: 617–628
- 2 Zhang M, He C, Shen G. Identity-based key verifiable inner product functional encryption scheme. In: *Proceedings of Blockchain Technology and Emerging Applications*, 2024. 3–22
- 3 Song G, Deng Y, Huang Q, et al. Hierarchical identity-based inner product functional encryption. *Inf Sci*, 2021, 573: 332–344
- 4 Belel A, Dutta R, Mukhopadhyay S. Hierarchical identity-based inner product functional encryption for unbounded hierarchical depth. In: *Proceedings of Stabilization, Safety, and Security of Distributed Systems*, 2023. 274–288
- 5 Rouselakis Y, Waters B. Efficient statically-secure large-universe multi-authority attribute-based encryption. In: *Proceedings of Financial Cryptography and Data Security*, 2015. 315–332
- 6 Zhu Y, Zhou Y, Wang J, et al. Revocable-hierarchical-identity-based inner product function encryption in smart healthcare. *IEEE Internet Things J*, 2025, 12: 15319–15332
- 7 Abdalla M, Benhamouda F, Kohlweiss M, et al. Decentralizing inner-product functional encryption. In: *Proceedings of Public-Key Cryptography*, 2019. 128–157
- 8 Lewko A, Waters B. Decentralizing attribute-based encryption. In: *Proceedings of Advances in Cryptology—EUROCRYPT 2011*, 2011. 568–588
- 9 Tseng Y F, Gao S J. Decentralized inner-product encryption with constant-size ciphertext. *Appl Sci*, 2022, 12: 636
- 10 Han J, Chen L, Hu A, et al. Privacy-preserving decentralized functional encryption for inner product. *IEEE Trans Dependable Secure Comput*, 2024, 21: 1680–1694
- 11 Chen Y, Zhang L, Yiu S M. Practical attribute based inner product functional encryption from simple assumptions. *Cryptology ePrint Archive*, 2019. <https://eprint.iacr.org/2019/846>
- 12 Belel A, Dutta R. Attribute-based inner product functional encryption in key-policy setting from pairing. In: *Proceedings of Advances in Information and Computer Security*, 2024. 101–121
- 13 Lewko A, Okamoto T, Sahai A, et al. Fully secure functional encryption: attribute-based encryption and (hierarchical) inner product encryption. In: *Proceedings of Advances in Cryptology—EUROCRYPT 2010*, 2010. 62–91
- 14 Yang H, Peng C. CP-IPFE: ciphertext-policy based inner product functional encryption. *IEEE Trans Inform Forensic Secur*, 2024, 19: 5419–5433
- 15 Bethencourt J, Sahai A, Waters B. Ciphertext-policy attribute-based encryption. In: *Proceedings of 2007 IEEE Symposium on Security and Privacy*, 2007. 321–334
- 16 Ye Y, Zhang L, You W, et al. Secure decentralized access control policy for data sharing in smart grid. In: *Proceedings of IEEE Conference on Computer Communications Workshops (INFOCOM WKSHPS)*, 2021. 1–6
- 17 Zhang L, Zhang T, Wu Q, et al. Secure decentralized attribute-based sharing of personal health records with blockchain. *IEEE Internet Things J*, 2022, 9: 12482–12496
- 18 Cheng H, Lo S L, Lu J. A blockchain-enabled decentralized access control scheme using multi-authority attribute-based encryption for edge-assisted Internet of Things. *Internet Things*, 2024, 26: 101220
- 19 Zhou Y W, Xu R, Qiao Z R, et al. Leakage-resilient attribute-based encryption scheme with large universe. *Sci Sin Inform*, 2024, 54: 2761–2777 [周彦伟, 徐然, 乔子芮, 等. 支持大属性集合的抗泄露属性基加密机制. *中国科学: 信息科学*, 2024, 54: 2761–2777]
- 20 Dodis Y, Goldwasser S, Kalai Y T, et al. Public-key encryption schemes with auxiliary inputs. In: *Proceedings of Theory of Cryptography*, 2010. 361–381
- 21 Datta P, Komargodski I, Waters B. Decentralized multi-authority ABE for NC^1 from BDH. *J Cryptol*, 2023, 36: 6
- 22 Qiao Z, Zhu Y, Zhou Y, et al. A continuous leakage-resilient CCA secure identity-based key encapsulation mechanism in the standard model. *J Syst Architecture*, 2025, 162: 103388

Decentralized attribute-based inner product function encryption scheme under static assumption

Yanwei ZHOU¹, Xianxiang LIU¹, Guangjin ZHANG¹, Zirui QIAO^{2*}, Yonghui LU^{1*}, Bo YANG¹,
Tianqing ZHU³ & Mingwu ZHANG⁴

1. School of Artificial Intelligence and Computer Science, Shaanxi Normal University, Xi'an 710119, China

2. School of Cyberspace Security, Xi'an University of Posts and Telecommunications, Xi'an 710121, China

3. Faculty of Data Science, City University of Macau, Macao 999078, China

4. School of Computer, Hubei University of Technology, Wuhan 430068, China

* Corresponding author. E-mail: qzr@xupt.edu.cn, lyhui008@snnu.edu.cn

Abstract In traditional cryptographic systems, data must undergo decryption to access its original form. This methodology fails to ensure sufficient privacy preservation for raw information in sensitive scenarios, consequently compromising the practical applicability of conventional encryption mechanisms. Recent advancements in cryptographic primitives have introduced inner product functional encryption (IPFE) to enable simultaneous computation of inner products and preservation of data confidentiality through direct ciphertext operations, eliminating decryption requirements. However, identity-based IPFE implementations exhibit limitations in achieving fine-grained access control for ciphertext. Meanwhile, attribute-based IPFE constructions over prime-order groups predominantly depend on non-static assumptions for security proofs. Although such assumptions theoretically maintain resistance against standard adversarial attacks, their dependence on quantifying adversarial queries introduces vulnerabilities to specialized attack vectors, and existing frameworks predominantly employ centralized authority structures for generating global cryptographic keys, creating single points of failure in attribute management systems and inherent scalability constraints. Therefore, this paper proposes a decentralized attribute-based IPFE (AB-IPFE) scheme constructed over prime-order groups under static assumptions to address the above limitations. The security is proved based on the Decisional Bilinear Diffie-Hellman assumption, eliminating dependency on non-static assumptions and substantially strengthening the security foundation of AB-IPFE systems. Comparative performance evaluations demonstrate superior operational efficiency compared to existing IPFE implementations and are advantageous for resource-constrained IoT deployments. Notably, cryptographic schemes proven secure in idealized models frequently succumb to practical vulnerabilities through implementation attacks. To further address this discrepancy, we develop a leakage-resilient variant of our decentralized AB-IPFE scheme.

Keywords inner product function encryption, attributed-based inner product encryption, decentralization, static hardness assumption, leakage resilience