



Delft University of Technology

抵抗恶意服务器的口令增强加密方案

Zhao, Yi; Liu, Hang; Liang, Kaitai; Ming, Yang; Zhao, Xiang Mo; Yang, Bo

DOI

[10.13328/j.cnki.jos.006440](https://doi.org/10.13328/j.cnki.jos.006440)

Publication date

2023

Document Version

Accepted author manuscript

Published in

Ruan Jian Xue Bao/Journal of Software

Citation (APA)

Zhao, Y., Liu, H., Liang, K., Ming, Y., Zhao, X. M., & Yang, B. (2023). 抵抗恶意服务器的口令增强加密方案. *Ruan Jian Xue Bao/Journal of Software*, 34(5), 2482-2493. <https://doi.org/10.13328/j.cnki.jos.006440>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

抵抗恶意服务器的口令增强加密方案*

赵一¹, 刘行¹, LIANG Kaitai², 明洋¹, 赵祥模¹, 杨波³

¹(长安大学 信息工程学院, 陕西 西安 710064)

²(Faculty of Electrical Engineering, Mathematics & Computer Science, Delft University of Technology, Delft, the Netherlands)

³(陕西师范大学 计算机科学学院, 陕西 西安 710119)

通信作者: 赵一, E-mail: yizhao@chd.edu.cn



摘要: 口令增强加密是一个近年来新出现的原语, 可以通过增加一个第三方密码服务提供商承担辅助解密的功能, 抵抗已有的服务器猜测低熵口令即可解密带来的恶意离线攻击风险, 即实现了对口令认证进行增强并增加加密的功能. 结合近年来新出现的算法替换攻击威胁, 对提出该原语工作中的方案给出了一种服务器积极攻击的方法, 该攻击具有不可检测性且可以让服务器仍然能实施离线攻击, 从而证明原方案不具备其声称的抵抗恶意服务器的功能. 接着讨论与总结能够抵抗恶意服务器实施算法替换攻击的方案应当具备的性质与构造特点; 随后, 给出一个能够真正抵抗恶意服务器算法替换攻击的方案并给出了仿真结果; 最后, 对于复杂交互式协议受到算法替换攻击时的安全性影响需要的系统性研究进行了展望.

关键词: 口令增强加密; 算法替换攻击; 不可检测性; 抵抗恶意服务器

中图法分类号: TP306

中文引用格式: 赵一, 刘行, LIANG Kaitai, 明洋, 赵祥模, 杨波. 抵抗恶意服务器的口令增强加密方案. 软件学报, 2023, 34(5): 2482–2493. <http://www.jos.org.cn/1000-9825/6440.htm>

英文引用格式: Zhao Y, Liu H, Liang KT, Ming Y, Zhao XM, Yang B. Password Hardening Encryption Services Against Malicious Server. Ruan Jian Xue Bao/Journal of Software, 2023, 34(5): 2482–2493 (in Chinese). <http://www.jos.org.cn/1000-9825/6440.htm>

Password Hardening Encryption Services Against Malicious Server

ZHAO Yi¹, LIU Hang¹, LIANG Kaitai², MING Yang¹, ZHAO Xiang-Mo¹, YANG Bo³

¹(School of Information Engineering, Chang'an University, Xi'an 710064, China)

²(Faculty of Electrical Engineering, Mathematics & Computer Science, Delft University of Technology, Delft, The Netherlands)

³(School of Computer Science, Shaanxi Normal University, Xi'an 710119, China)

Abstract: Password hardening encryption (PHE) is an emerging primitive in recent years. It can resist offline attack brought by keyword guessing attack from server via adding a third party with crypto services joining the decryption process. This primitive enhances the password authentication protocol and adds encryption functionality. This paper presents an active attack from server in the first scheme that introduced this primitive. This attack combines the idea from a cutting-edge threat called algorithm substitution attack which is undetectable and makes the server capable of launching offline attack. This result shows that the original PHE scheme can not resist attacks from malicious server. Then this study tries to summarize the property that an algorithm substitution attack resistant scheme should have. After that this paper presents a PHE scheme that can resist such kind of attacks from malicious server with simulation results. Finally, this study concludes the result and gives some expectation for future systematic research on interactive protocols under algorithm substitution attack.

* 基金项目: 国家重点研发计划 (2017YFB0802000); 国家自然科学基金 (62072054, U2001205, 61772326, 61802241, 61802242); 陕西省重点研发计划 (2021GY-047); 长安大学中央高校基本科研业务费专项 (300102240102)

收稿时间: 2021-04-15; 修改时间: 2021-06-01, 2021-06-29; 采用时间: 2021-08-26; jos 在线出版时间: 2022-09-30

CNKI 网络首发时间: 2022-11-16

Key words: password hardening encryption (PHE); algorithm substitution attack; undetectable; malicious server

使用口令进行认证是当前网络认证中最广泛采用的方法,绝大多数服务、应用都使用基于口令认证的方式实现对服务器的访问控制.现在服务器端已经广泛的采用存储口令的加盐哈希值取代直接存储口令本身的方法增强口令存储的安全性,但仍不能应对层出不穷的攻击方式.其中最为常见的,就是存储口令哈希值的数据库泄露后,使用离线暴力破解猜出口令,然后登录服务器获取数据.另一方面,由于口令往往较短,是低熵字符串,所以很容易用暴力破解的方法测试哈希值获得.现实中诸如 CSDN、天涯和世纪佳缘等知名网站都发生过口令数据库被入侵的事件,产生了极大的不良影响与损失.为应对上述威胁,出现了使用增加一个密码服务器的方式实现的口令增强协议 (password hardening services)^[1-3].随后, Lai 等人发现提供加密存储服务的服务器也存在风险,即密钥放在本地时,数据库泄露的同时会泄露解密密钥,这将导致即使加密存储也不能阻止口令被猜测出后的密文被解密.因此他们提出了口令增强加密 (password hardening encryption, PHE) 的概念,在口令增强协议的基础上令密码服务器也参与加密和解密过程,使得即使口令与数据库被敌手获得,服务器被敌手控制,离线攻击也难以实施^[4].由于该密码服务器使得口令认证和解密都需要在线进行,降低了敌手攻击的频率,因此被称为频率限制器 (rate limiter).

然而,一个能获得口令的敌手,其实能够“合法”的实施比离线攻击更强大的攻击方式,上述工作仅考虑抵抗离线猜测攻击仍不足以保障用户的数据安全.近年来,随着对斯诺登事件背后密码学原理研究的进展,算法替换攻击等后门攻击手段进入了研究人员的视野^[5].该类攻击利用在算法中嵌入后门或者改变参数选取方式的方法,可以在不影响协议正常功能的前提下,使得持有后门信息的敌手获得额外的信息.本文发现,如果上述方案中,能够收买服务器的敌手只要实施算法替换攻击,依然可以在离线情况下获得用户的数据.因此需要设计新的能够抵御该类更加积极的攻击方式的方案.而通用的抵抗算法替换攻击的手段如逆向防火墙,用于口令增强加密的场景则太过繁琐且效率低下.本文的具体贡献如下.

(1) 对 Lai 等人^[4]的口令增强加密方案提出了一种能够收买服务器的敌手可以实施的具备不可检测性的攻击方法,在完成正常协议功能,其他参与方都无法发现的基础上获得用户的数据.

(2) 分析该方案不能抵抗上述攻击的原因,给出能够抵抗该算法替换攻击的方案应当具备的基本特点.

(3) 基于 Lai 等人^[4]的方案提出了新的能够抵御上述积极攻击的方案,并对其安全性进行了证明与效率分析.

本文第 1 节介绍口令增强加密和算法替换攻击的研究现状并分析其安全特性.第 2 节介绍需要用到的基础知识和安全模型.第 3 节提出恶意服务器可以实施的新的积极攻击方法.第 4 节分析容易遭受该类攻击方案的原因以及要抵抗该类攻击应该具备的安全特性.第 5 节详细介绍本文提出的新的口令增强加密方案并证明其安全特性.第 6 节对本文方案进行试验仿真分析,展示方案的实用性.最后总结全文,并对口令增强加密方案未来的发展进行探讨.

1 相关研究

1.1 口令增强的相关研究

2015 年, Everspaugh 等人首先提出用外置的密码服务来增强口令认证协议的安全性^[1].他们提出一个新的密码学原语-可验证部分不经意的伪随机函数 (partially oblivious, PRF) 并用其作为外置的密码服务,同时给出了部分不经意性的形式化定义.然而可验证的伪随机函数需要较强的交互式假设才能构造,且用到的对运算效率较低,因此 Schneider 等人提出了使用部分不经意的承诺来降低假设强度以及提升方案效率^[2].Lai 等人指出 Schneider 等人的方案不能抵抗离线猜测攻击,并提出了效率 and 安全性都更强的方案^[3].随后 Lai 等人提出口令增强加密的概念^[4],考虑提供加密存储的服务器,将加密功能加入到口令增强协议中,他们用改良的方案保护短密钥,然后用密钥加密长密文. Brost 等人给出了门限的口令增强加密方案^[6].

上述方案的构造,均是在用户和服务器之外增加一个半诚实的密码服务器-频率限制器,使得服务器无论进行口令认证或者解密都需要在线完成,降低了猜测的频率,因此可以用于抵抗离线攻击.该密码服务器由密码专业人士运营提供辅助服务,不易受到攻击,但是对服务器的用户内容是好奇的.这些工作均假设服务器被收买后,依然

是从正常协议流程中尝试获得信息,是被动的敌手,因此不能抵抗更加积极的攻击类型.

1.2 算法替换攻击的相关研究

斯诺登事件之后,对后门攻击的研究进入了学界的视野.算法替换攻击是后门攻击的一种重要类型,指的是敌手可以通过木马等入侵方式替换用户正在使用的算法.成功的算法替换攻击要求攻击具有不可检测性,即用户和其他协议参与方仍然能够完成方案的既定功能,然而具有陷门的敌手可以从交互消息中得到正常运行无法得到的信息.2014年,Bellare等人提出了算法替换攻击的形式化定义与模型^[7],并给出了针对对称加密算法进行算法替换攻击的方法.随后几年里,针对签名^[8,9]、消息鉴别码^[10,11]、密钥生成函数以及单向函数等^[12-14]原语,学界都提出了对应的算法替换攻击方法.通常抵抗该类攻击的手段主要有两种,一是使用确定性算法,然而这种方法会导致传统安全性的下降.另一种是使用密码学逆向防火墙^[15],目前并没有能够高效的通用方案或者构造,都是针对特定原语构造适合原语的高效防火墙^[16-18].

上述两个领域的发展情况表明,口令增强方案考虑的安全性局限于抵抗传统的敌手,在敌手收买服务器的情况下,没有考虑敌手偏离协议算法进行攻击的可能,实际上安全性只针对半诚实敌手设计.该方案场景下,算法替换攻击的方式以及抵抗的方法仍缺少研究.因此本文将对已有的口令增强加密方案给出服务器的算法替换攻击方法,并研究如何抵抗该类攻击.

2 定义与安全模型

2.1 口令增强加密的定义

一个口令增强加密方案 PHE 由 7 个算法组成.

- $Setup(1^\lambda)$: 初始化算法输入安全参数,输出公开参数 pp .
- $Kgen_S(pp, 1^\lambda)$: 服务器密钥生成算法,输出服务器的公私钥对 (pk_S, sk_S) .
- $Kgen_R(pp, 1^\lambda)$: 频率限制器密钥生成算法,输出频率限制器的公私钥对 (pk_R, sk_R) .
- $Enc(pw, sk_S, m, sk_R, l)$: 用户使用口令 pw 登录服务器并上传明文消息 m ,服务器和频率限制器以公共输入 $l = (l_S, l_R)$ 联合执行(带标签的)加密协议,服务器得到密文 C 以及更新过的标签 (l'_S, l'_R) .
- $Dec(pw, sk_S, C, sk_R, l_S, l_R)$: 用户使用口令 pw 登录服务器,服务器追溯密文 C 以及标签 (l_S, l_R) ,然后和频率限制器联合执行(带标签的)解密协议,服务器输出消息 m 或者 $\perp$.
- $Rotate(sk_S, sk_R)$: 服务器和频率限制器联合更新密钥,输入双方的私钥,服务器输出更新后的公私钥对和将用于密文更新的令牌 (pk'_S, sk'_S, τ) ,频率限制器输出更新后的公私钥对 (pk'_R, sk'_R) .
- $Update(\tau, C, l)$: 密文更新算法,输入更新令牌 τ 以及密文 C 和对应的标签 $l = (l_S, l_R)$,输出更新后的密文 C' .

2.2 口令增强加密的经典安全模型

口令增强加密方案可以拥有的安全性质有消息隐藏性、部分不经意性、可靠性以及前向安全性,下面具体给出每种安全性质的形式化定义.

- 消息隐藏性

该安全性质主要刻画恶意服务器在猜测出口令之前不能通过离线猜测攻击获得明文的任何信息,由如下游戏 $Hide_{PHE,A}^b$ 来定义.

- (1) $pp \leftarrow Setup(1^\lambda), (pk_R, sk_R) \leftarrow Kgen_R(pp, 1^\lambda)$
- (2) $O := \{P \langle \cdot, R(sk_R, \dots) \rangle\}, P \in \{Enc, Dec, Rotate\}$
- (3) $(sk_S^*, pw^*, m_0^*, m_1^*) \leftarrow A_1^O(pp, pk_R)$
- (4) $(C^*, l^*) \leftarrow Enc(pw^*, sk_S^*, m_b^*, sk_R, l)$
- (5) $b' \leftarrow A_2^O(C^*, l^*)$
- (6) return $b = b'$

其中, A_2 不能向 $O := \{Dec(\cdot, R(sk_R))\}$ 询问挑战口令 pw^* 和 sk^* .

如果对任意敌手 $A=(A_1, A_2)$, 存在一个可忽略函数 $negl(\cdot)$ 在上述游戏中, 有 $|\Pr[Hide_{PHE,A}^0(1^\lambda)=1] - \Pr[Hide_{PHE,A}^1(1^\lambda)=1]| < negl(\lambda)$, 则该口令增强加密方案是消息隐藏的.

- 部分不经意性

该性质主要刻画频率限制器不能获得明文的任意信息, 部分的意思是频率限制器可能会关联同一用户的两次操作, 由如下游戏 $POb_{PHE,A}^b$ 来定义.

- (1) $pp \leftarrow Setup(1^\lambda), (pk_S, sk_S) \leftarrow Kgen_S(pp, 1^\lambda)$
- (2) $O := \{P(S(sk_S), \cdot), \cdot\}, P \in \{Enc, Dec, Rotate\}$
- (3) $(pw_0^*, m_0^*, pw_1^*, m_1^*) \leftarrow A_1^O(pp, pk_S)$
- (4) $(C^*, l^*, st) \leftarrow Enc(sk_S^*, pw_b^*, m_b^*)$
- (5) $b' \leftarrow A_2^O(C^*, l^*, st)$
- (6) return $b = b'$

其中, A_2 不能向 $O := \{Dec(S(sk_S), \cdot)\}$ 询问挑战口令 (pw_0^*, l^*) 和 (pw_1^*, l^*) .

如果对任意敌手 $A=(A_1, A_2)$, 存在一个可忽略函数 $negl(\cdot)$ 在上述游戏中, 有 $|\Pr[POb_{PHE,A}^0(1^\lambda)=1] - \Pr[POb_{PHE,A}^1(1^\lambda)=1]| < negl(\lambda)$, 则该口令增强加密方案是部分不经意的.

- 可靠性

一个口令增强加密方案具备可靠性意味着使用错误的口令能够解密合法生成密文的概率是可忽略的.

一个口令增强加密方案具备强可靠性意味着即使是非法的密文, 恶意频率限制器不按照算法规定行动时一定会被检测到, 偏离算法规定的行动包括让服务器使用相同标签和口令解密相同密文输出不同结果, 或者是把非法密文当合法密文正常解密.

- 前向安全性

一个口令增强加密方案具备前向安全性意味着密钥和密文更新后, 旧的密钥无法解密更新后的密文, 更新后的密钥和密文依然能够按照正常协议解密.

2.3 对口令增强加密服务器的算法替换攻击

一个算法 PHE 的替换算法 $\overline{PHE} = (\overline{Setup}, \overline{Kgen_S}, \overline{Kgen_R}, \overline{Enc}, \overline{Dec}, \overline{Rotate}, \overline{Update})$ 定义如下:

- $\overline{Kgen_S}(pp, 1^\lambda)$: 替换后的服务器密钥生成算法, 输出服务器的带后门密钥 $subk$.
- $\overline{Enc}(pw, subk, m, sk_R, l)$: 用户使用口令 pw 登录服务器并上传明文消息 m , 服务器和频率限制器以公共输入 $l = (l_S, l_R)$ 联合执行 (带标签的) 加密协议, 服务器得到密文 C 以及更新过的标签 (l'_S, l'_R) .
- $\overline{Dec}(pw, subk, C, sk_R, l_S, l_R)$: 用户使用口令 pw 登录服务器, 服务器追溯密文 C 以及标签 (l_S, l_R) , 然后和频率限制器联合执行 (带标签的) 解密协议, 服务器输出消息 m 或者 $\perp$.
- $\overline{Rotate}(subk, sk_R)$: 服务器和频率限制器联合更新密钥, 输入双方的私钥, 服务器输出更新后的公私钥对和将用于密文更新的令牌 $(subk', \tau)$, 频率限制器输出更新后的公私钥对 (pk'_R, sk'_R) .
- $\overline{Update}(\tau, C, l)$: 密文更新算法, 输入更新令牌 τ 以及密文 C 和对应的标签 $l = (l_S, l_R)$, 输出更新后的密文 C' .

一个算法替换攻击方案如果满足 $View_{PHE}^A(1^\lambda) \stackrel{c}{\approx} View_{\overline{PHE}}^A(1^\lambda)$, 即区分的敌手在正常方案中和算法替换方案中的视图是不可区分的, 则称其为不可检测的. 这里区分的敌手可以是用户和频率限制器.

2.4 抵抗对服务器算法替换攻击的口令增强加密

本文要研究抵抗对服务器算法替换攻击的口令增强加密方案, 需要给出对应的安全模型. 由于频率限制器在相关文献中被设定为由密码专业人员运营的服务器, 因此自身可以有恶意行为, 但是不容易受密码学攻击影响, 因此本文不考虑对频率限制器的算法替换攻击. 在口令增强加密的经典安全性中, 部分不经意性和可靠性都是针对恶意频率限制器的, 故此处不给出对应抵抗算法替换攻击的安全模型. 下面给出定义抵抗服务器算法替换攻击的

消息隐藏性的安全游戏 $Hide_{PHE, \bar{A}}^b$:

- (1) $pp \leftarrow Setup(1^\lambda), (pk_R, sk_R) \leftarrow Kgen_R(pp, 1^\lambda), \overline{PHE} \leftarrow \bar{A} = (\bar{A}_1, \bar{A}_2)$
- (2) $O := \{P \langle \cdot, R(sk_R, \dots) \rangle\}, P \in \{\overline{Enc}, \overline{Dec}, \overline{Rotate}, \overline{Update}\}$
- (3) $(subk = sk_S^*, pw^*, m_0^*, m_1^*) \leftarrow \bar{A}_1^O(pp, pk_R)$
- (4) $(C^*, l^*) \leftarrow Enc(pw^*, sk_S^*, m_b^*, sk_R, l)$
- (5) $b' \leftarrow \bar{A}_2^O(C^*, l^*)$
- (6) return $b = b'$

其中, $\overline{PHE} = (Setup, \overline{Kgen_S}, Kgen_R, \overline{Enc}, \overline{Dec}, \overline{Rotate}, \overline{Update})$ 是第 2.3 节中给出的 PHE 的视图不可区分的替换算法. $\bar{A}_2$ 不能向 $O := \{\overline{Dec}(S(sk_S), \cdot)\}$ 询问挑战口令 pw^* 和 sk^* .

如果对任意敌手 $\bar{A} = (\bar{A}_1, \bar{A}_2)$ 存在一个可忽略函数 $negl(\cdot)$ 在上述游戏中, 有 $|\Pr[Hide_{PHE, \bar{A}}^0(1^\lambda)=1] - \Pr[Hide_{PHE, \bar{A}}^1(1^\lambda)=1]| < negl(\lambda)$, 则该口令增强加密方案在算法替换攻击下依然是消息隐藏的.

3 恶意服务器的算法替换攻击

本节将给出文献 [4] 中方案的恶意服务器可以实施导致离线攻击依然发生的算法替换攻击, 然后证明该攻击是不可检测的. 为了方便读者进行比较, 文献中的原方案将同时给出.

3.1 文献 [4] 方案

给定算法描述 $AES = (G, E, D)$, 零知识证明协议 $POK = (Gen_{crs}, P, V)$.

• $Setup(1^\lambda) = (H_S, H_R, g, G), |G| = q, H_S, H_R : \{0, 1\}^* \rightarrow G$, 其中 G 是 q 阶乘法循环群.

• $Kgen_S(1^\lambda) = (pk_S = \perp, sk_S = y \in \mathbb{Z}_p)$

• $Kgen_R(1^\lambda) = (pk_R = g^x, sk_R = x)$

• $Enc(pw, sk_S, m, sk_R, r_S, r_R) = (C_0, C_1, C_m)$ 其中 $C_0 = H_R^x(r_R, 0)H_S^y(pw, r_S, 0)$, r_S, r_R 分别是服务器和频率限制器生成的随机数, $K \leftarrow G, C_1 = H_R^x(r_R, 1)H_S^y(pw, r_S, 1)K^y, C_m = E_K(m)$.

交互式加密过程是频率限制器先计算 $(H_R^x(r_R, 0), H_R^x(r_R, 1))$ 和一个零知识证明 $\pi = POK(stmt, x)$, 其中 $stmt = \{\log Y_1 = \log Y_2\}, (Y_1 = H_R^x(r_R, 0), Y_2 = H_R^x(r_R, 1))$. 然后将 (Y_1, Y_2, r_R, π) 发送给服务器, 服务器验证 π 是有效的证明后, 计算 $C = (C_0, C_1, C_m) = (H_R^x(r_R, 0)H_S^y(pw, r_S, 0), H_R^x(r_R, 1)H_S^y(pw, r_S, 1)K^y, C_m)$, 存储 (C, r_S, r_R) , 丢弃 m 和 π .

• $Dec(pw, sk_S, C, sk_R, r_S, r_R) = m$, 交互式解密过程如下.

服务器首先计算 $T_0 = C_0H_S^{-y}(pw, r_S, 0)$ 并和 r_R 一起发给频率限制器.

频率限制器检查是否 $T_0 = H_R^x(r_R, 0)$, 如果成立则计算 $T_1 = H_R^x(r_R, 1)$ 和证明 $\pi = POK(stmt, x)$, $stmt = \{\log T_0 = \log T_1\}$, 将 (T_1, π) 发送给服务器.

服务器检查 π 是否成立, 如果成立, 则计算 $K = (T_1^{-1}H_S^{-y}(pw, r_S, 1))^{1/y}$, $m = D_K(C_m)$, 否则输出 $\perp$.

• $Rotate(sk_S, sk_R)$: 频率限制器首先随机选择 (α, β) 然后发给服务器, 服务器更新自己的密钥为 $y' = \alpha y$, 频率限制器更新私钥为 $x' = \alpha x + \beta$, 计算公钥 $y' = g^{x'}$ 并发公开.

• $Update(C, r_S, r_R)$: 密文更新算法, 服务器计算更新后的密文 $C' = (C'_0, C'_1, C'_m) = (C_0^\alpha H_R^\beta(r_R, 0), C_1^\alpha H_R^\beta(r_R, 1), C_m)$.

如后文图 1 所示, 文献 [4] 方案在注册过程中, 第 1 步选择的对称密钥没有任何承诺或者证明绑定, 恶意的服务器完全可以一开始不选择对称密钥, 然后根据频率限制器发来的消息再生成对应的密钥, 就能在外界没有察觉的情况下破坏协议的安全性. 该方案的登录流程如后文图 2.

3.2 算法替换攻击的方案

本文方案只针对服务器进行攻击, 因此服务器不参与的算法方案和原方案相同, 此处不再赘述, 只给出服务器参与算法的替换版本. 本文方案的注册和登录流程如图 3 和图 4 所示.

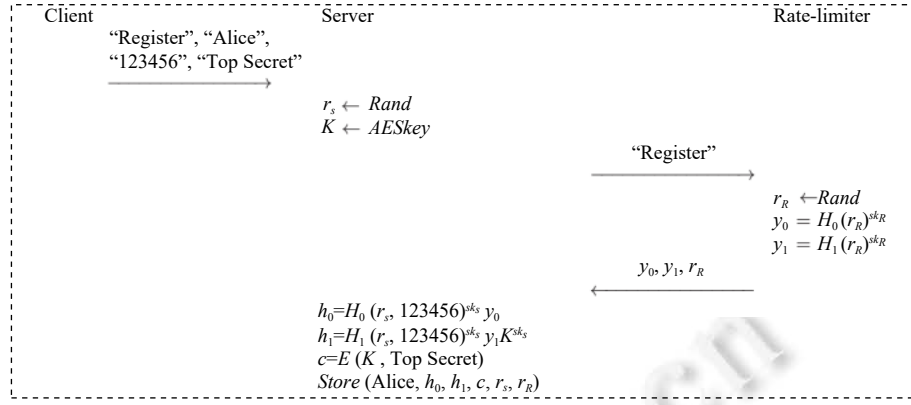


图 1 文献 [4] 方案注册流程示意图

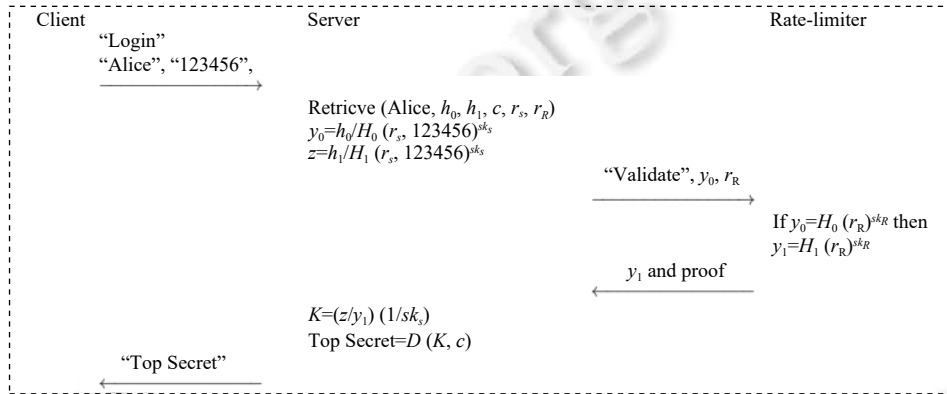


图 2 文献 [4] 方案登录流程示意图

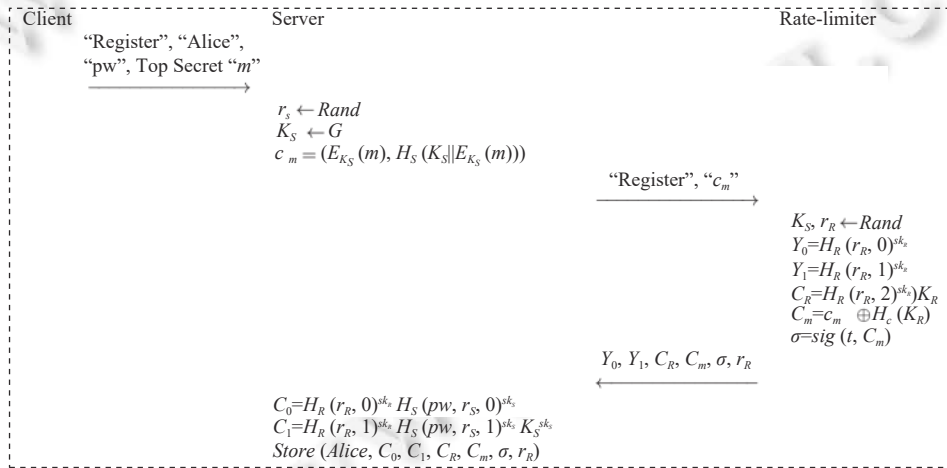


图 3 本文方案注册流程示意图

- $\overline{Kgen_S}(1^\lambda) = Kgen_S(1^\lambda)$
- $\overline{Enc}(pw, sk_S, m, sk_R, r_S, r_R) = (C_0, C_1, C_m)$, 其中, $C_0 = H_R^x(r_R, 0) H_S^y(pw, r_S, 0)$, $K \leftarrow H_S(r_S, sk_S)$, $C_1 = H_R^x(r_R, 1) H_S^y(pw, r_S, 1) K^{sk_s}$, $C_m = E_K(m)$.

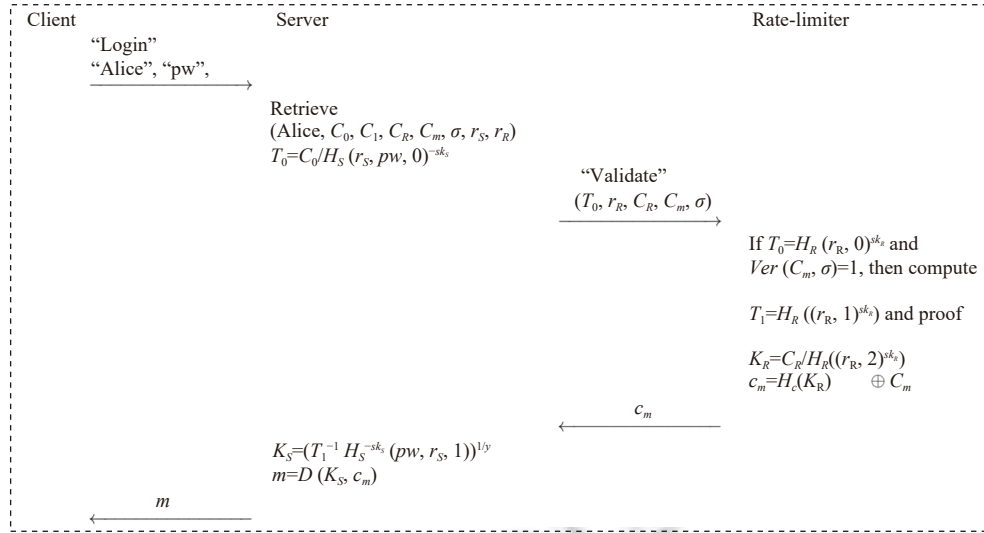


图4 本文方案登录流程示意图

交互式加密过程是频率限制器先计算 $(H_R^x(r_R, 0), H_R^x(r_R, 1))$ 和一个零知识证明 $\pi = \text{POK}(stmt, x)$, 其中 $stmt = \{\log Y_1 = \log Y_2\}$, $(Y_1 = H_R^x(r_R, 0), Y_2 = H_R^x(r_R, 1))$, 服务器验证发来的元组确实是合法生成后, 计算 $C = (C_0, C_1, C_m) = (H_R^x(r_R, 0) H_S^y(pw, r_S, 0), H_R^x(r_R, 1) H_S^y(pw, r_S, 1) K^y, C_m)$, 存储 C_m , 丢弃 m 和 π .

• 解密算法有两种工作模式, 如果是用户正常在线登录执行解密, 则 $\overline{Dec}(pw, sk_S, C, sk_R, r_S, r_R) = Dec(pw, sk_S, C, sk_R, r_S, r_R)$.

如果服务器要在离线情况下解密, 则计算 $K = H_S(r_S, sk_S)$, $m = D_K(C_m)$.

• $\overline{Rotate}(sk_S, sk_R) = Rotate(sk_S, sk_R)$.

3.3 不可检测性分析

可以看到, 上述替换攻击算法中, 加密算法仅改变了原算法中 AES 密钥的选取算法, 因此替换算法方案的正确性可以保持. 另外上述方案的安全性基于随机谕言机模型证明, 由哈希函数的伪随机性, $\Pr[K = K^* : K^* \leftarrow H_S(r_S, sk_S)] = \Pr[K = K^* : K^* \leftarrow \text{AES.Gen}] = 1/q$, 即两种密钥生成算法的分布是相同的, 且加密步骤的其他协议也保持不变, 因此加密替换算法的输出和原算法是不可区分的.

替换后的解密算法有两个工作模式, 在线的工作模式和原算法完全相同, 因此在正常合法执行协议的过程中, 用户和频率限制器的视图和正常方案也是完全相同的. 密钥和密文更新算法都和原方案一样. 离线模式下, 密文形式保持不变且与其他参与方交互. 因此, 该替换攻击算法也可以抵御时间分析检测.

综上所述, 该替换攻击算法具有不可检测性.

4 分析与讨论

4.1 易遭受算法替换攻击的原因分析

上述算法替换攻击, 只改变随机选取 AES 密钥的算法, 就达到了攻击效果. 结合已知的针对加密和签名方案的成功攻击方法可以发现, 都是通过对方案中随机数产生的过程进行篡改来达到攻击的目的. 替换攻击算法做到不可检测性的关键, 就在于替换随机数产生的算法和原方案随机数生成的算法是不可区分的, 文献 [7] 中攻击对称加密算法就用的是替换随机的初始向量的方法. 文献 [8] 攻击签名的方法, 是用伪随机函数的输出替换签名使用的随机数, 而伪随机函数的输入是敌手可以公开获得的信息, 从而敌手可以知道部分签名使用的随机数, 从而计算出签名的私钥.

从这类攻击算法中可以发现, 容易遭受算法替换攻击的方案, 随机数只是提供了攻击的渠道, 根本性的原因在于, 无论签名还是加密, 都是执行签名和加密的用户单方面完全决定的, 那么该用户无论出于何种原因, 只要存在不可检测的替换算法, 就可以不被察觉的运行替换算法而不是真实算法. 文献 [4] 的研究动机之一, 在于口令增强加密算法中, 如果一个服务器有完全的解密能力, 就不能避免离线猜测攻击, 因此解密过程需要频率限制器参与使得只能在线解密. 而本文提出这样一个论断, 只要服务器具有完全的加密能力, 就不能避免算法替换攻击. 文献 [4] 中方案会遭受如上的攻击, 就在于 AES 密文的密钥完全是由服务器自己决定的, 虽然密钥本身通过频率限制器的加入进行了加密, 但是 AES 的密文依然完全由服务器自身决定, 这就给了服务器可乘之机.

4.2 抵抗算法替换攻击的方法讨论

一般性的随机化算法抵抗算法替换攻击的方式是使用逆向防火墙. 在文献 [16] 中, 给出了适用于消息传输协议的逆向防火墙构造. 值得强调的是, 为了高效的处理长明文, 不能简单地使用可再随机化加密的方式直接再随机化密文, 使用传统的混合加密的方式就会遇到和文献 [4] 方案一样的会收到算法替换攻击的困境, 因此文献 [16] 使用了带逆向防火墙的密钥协商方案来解决对称加密方案使用密钥的问题, 即如果协议正常执行, 发送方和接收方会算出相同的密钥, 因此如果加密者替换对称密钥, 会导致接收者无法解密, 从而被检测到替换攻击行为, 因此只能按照方案规定流程进行.

然而, 上述解决方法在口令增强加密的环境下是不可行的. 因为口令增强加密中, 频率限制器也只是半诚实的, 并不是消息的接收方. 并且方案要求频率限制器满足部分不经意性, 因此用于加密明文的密钥是不能让频率限制器获得的, 这样让服务器和频率限制器进行密钥协商来限制服务器行为的方法将损害原语本身要求的安全性.

综上所述, 口令增强方案抵抗服务器算法替换攻击的挑战在于既不能使用混合加密, 也不能使用密钥协商, 由第 4.1 节的讨论可知, 既不让服务器具有完全的加密能力, 又需要频率限制器参与加密过程, 现有的方法均不能直接解决该问题.

4.3 本文方案的构想

基于以上讨论, 本文将尝试使用重加密作为一个既让频率限制器参与加密, 又不会让其知道明文和密钥信息的方式来达到抵抗服务器算法替换攻击的目的. 即令服务器在自己加密消息之后, 将密文也发给频率限制器, 由频率限制器再度对密文进行加密并签名发回给服务器进行存储. 这样, 即使服务器实施算法替换攻击, 其存储的密文是被频率限制器再加密过的密文, 没有频率限制器参与是不能解密的, 这样即使服务器猜中口令也无法实施离线攻击.

对于频率限制器一侧, 文献 [4] 提出的场景设定该外置服务是由密码学专家管理的服务器, 不易受到攻击, 且其方案已经用协议中需要做出证明的过程, 对频率限制器可能的恶意行为进行了限制. 即使协议改为可以收到服务器加密过的密文, 对于密文的内容也是不能获取的, 部分不经意性依然成立. 而且, 提供外置密码服务是有盈利动机的, 在不能获得信息的情况下恶意扭曲协议不是理性的敌手会采取的行动. 因此本文对抵抗频率限制器的算法替换攻击不做考虑. 这也是本文不使用通用的逆向防火墙解决方案的原因之一, 对口令增强加密的服务器的算法替换攻击, 可以有针对性的高效防御手段.

5 抵抗恶意服务器的口令增强加密方案

5.1 构造

给定算法描述 $AES = (G_A, E, D)$, Schnorr 签名方案 $Schn = (Gen, Sig, Ver)$, 零知识证明协议 $POK = (Gen_{crs}, P, V)$.

• $Setup(1^\lambda) = (H_S, H_R, H_c, g, G)$, $|G| = q$, 其中 G 是 q 阶乘法循环群, $H_S, H_R : \{0, 1\}^* \rightarrow G$, $H_c : \{0, 1\}^* \rightarrow \{0, 1\}^{l_m}$, 这里哈希函数 H_c 的输出长度根据密文长度改变, 使用一个普通的哈希函数 (如 SHA-256), 以递归迭代的方式可以得到.

$$\bullet Kgen_S(1^\lambda) = (pk_S = \perp, sk_S = y \in \mathbb{Z}_p)$$

$$\bullet Kgen_R(1^\lambda) = (pk_R = (g^x, g^t), sk_R = (x, t))$$

$$\bullet \text{Enc}(pw, sk_S, m, sk_R, r_S, r_R) = (C_0, C_1, C_R, C_m, \sigma)$$

其中, $C_0 = H_R^x(r_R, 0)H_S^y(pw, r_S, 0)$, r_S 、 r_R 分别是服务器和频率限制器生成的随机数; $C_1 = H_R^x(r_R, 1)H_S^y(pw, r_S, 1)K_S^y$, $K_S \leftarrow G_A$; $C_R = H_R^x(r_R, 2)K_R$, $K_R \leftarrow G$; $C_m = E_{K_S}(m) \oplus H_c(K_R)$; $\sigma = \text{Sig}(t, C_R, C_m)$.

交互式加密过程如下:

服务器选择 $\sigma = \text{Sig}(t, C_R, C_m)$, 计算 $K_S \leftarrow G$, 然后将 $c_m = (E_{K_S}(m), H_S(K_S \parallel E_{K_S}(m)))$ 发送给频率限制器.

频率限制器先随机选择 $K_R, r_R \in \{0, 1\}^*$, 计算 $C_m = c_m \oplus H_c(K_R)$, $(H_R^x(r_R, 0), H_R^x(r_R, 1), C_R = H_R^x(r_R, 2)K_R)$, 然后计算 $\sigma = \text{sig}(t, C_m) = (r, s)$, $\pi = \text{POK}(stmt, x)$, 其中 $stmt = \{\log Y_1 = \log Y_2 = \log g^x\}$, $Y_1 = H_R^x(r_R, 0)$ 将 $Y_2 = H_R^x(r_R, 1)$ 发送给服务器.

服务器验证 π 是有效的证明, 签名 σ 是合法的后, 计算 $C = (C_0, C_1, C_R, C_m) = (H_R^x(r_R, 0)H_S^y(pw, r_S, 0), H_R^x(r_R, 1)H_S^y(pw, r_S, 1)K_S^y, C_R, C_m)$, 存储 (C, σ, r_S, r_R) , 丢弃 m 和 π .

$\bullet \text{Dec}(pw, sk_S, C, sk_R, r_S, r_R) = m$, 交互式解密过程如下:

服务器首先计算 $T_0 = C_0 H_S^{-y}(pw, r_S, 0)$ 并发送 $(T_0, r_R, C_R, C_m, \sigma)$ 给频率限制器.

频率限制器检查是否 $T_0 = H_R^x(r_R, 0)$ 且 σ 是有效的签名, 如果成立则计算 $T_1 = H_R^x(r_R, 1)$, $\pi = \text{POK}(stmt, x)$, $stmt = \{\log T_0 = \log T_1\}$, $K_R = C_R / H_R^x(r_R, 2)$, $c_m = H_c(K_R) \oplus C_m$, 将 c_m 发送给服务器.

服务器验证 c_m 是否通过, 如果是, 计算 $K_S = (T_1^{-1} H_S^{-y}(pw, r_S, 1))^{1/y}$, $m = D_{K_S}(c_m)$, 否则输出 $\perp$.

$\bullet \text{Rotate}(sk_S, sk_R)$: 频率限制器首先随机选择 (α, β) 然后发给服务器, 服务器更新自己的密钥为 $y' = \alpha y$, 频率限制器更新私钥为 $x' = \alpha x + \beta$, $t' = \alpha t$, 计算公钥 $pk_R' = (g^x, g^y)$ 并发公开.

$\bullet \text{Update}(C, \sigma, r_S, r_R)$: 密文更新算法, 服务器计算更新后的密文 $C' = (C'_0, C'_1, C'_R, C'_m) = (C_0^\alpha H_R^\beta(r_R, 0), C_1^\alpha H_R^\beta(r_R, 1), C_R^\alpha H_R^\beta(r_R, 2), C_m)$, $\sigma' = (r', s') = (r^\alpha, \alpha s)$

5.2 安全性分析

定理 1. 消息隐藏: 如果 DDH 假设成立, 则以上方案对算法替换攻击依然是消息隐藏的.

证明: 我们用以下游戏序列来进行论证:

Game0: 初始的符合第 2.3 节算法替换攻击的游戏, 挑战者模拟频率限制器, 和敌手进行算法元组 $\overline{\text{PHE}} = (\text{Setup}, \overline{\text{Kgen}_S}, \overline{\text{Kgen}_R}, \overline{\text{Enc}}, \overline{\text{Dec}}, \overline{\text{Rotate}}, \overline{\text{Update}})$ 的交互.

Game1: 与 Game0 不同的是, 挑战者给出挑战密文中的 C_R 时, 不使用 K_R 计算而是随机选择一个 U_R , 计算 $C_R = H_R^x(r_R, 2)U_R$, 其他步骤都和 Game0 一样.

Game2: 与 Game1 不同的是, 计算 C_m 时也使用随机选择的 U_R 来计算, $C_m = E_{K_S}(m) \oplus H_c(U_R)$, 其他步骤都和 Game1 一样.

引理 1. 如果 DDH 假设成立, 那么 Game0 和 Game1 对于算法替换攻击的敌手是不可区分的.

证明: Game0 中, 敌手在交互式加密的步骤中收到 $(Y_1, Y_2, C_R, C_m, \sigma, \pi)$, 其中 $Y_1 = H_R^x(r_R, 0)$, $Y_2 = H_R^x(r_R, 1)$, (Y_1, C_R) 或 (Y_2, C_R) 都构成一个以 K_R 为明文的 ElGamal 加密方案. 该密文完全由频率限制器生成, 不受对服务器进行算法替换攻击的敌手影响.

因此若存在一个敌手能够区分 Game0 和 Game1, 那么依据 ElGamal 加密的证明过程, 挑战者可以构造一个算法 B 攻破 DDH 问题. 结论得证.

引理 2. Game1 和 Game2 是统计不可区分的.

证明: Game1 中, C_R 中的明文已经是随机选取的 U_R , 且 C_R 是 CPA 安全的, 因此从 C_R 中不能得到 K_R 的信息, 对于敌手来说, $\Pr[K = K_R] = 1/q$, 与随机选取的 U_R 分布相同, 因此 Game1 和 Game2 是统计不可区分的.

引理 3. 在 Game2 中, 方案是消息隐藏的.

证明: Game2 中, 挑战密文中的 $C_m = E_{K_S}(m) \oplus H_c(U_R)$, 由于 U_R 是均匀分布的, 因此 $H_c(U_R)$ 也是均匀分布的, 那么无论算法替换攻击的敌手如何影响加密算法 E , C_m 都是均匀分布的字符串, 敌手不能区分 C_{m_0} 和 C_{m_1} .

结合引理 1-引理 3, 定理 1 得证.

需要注意的是, 本文方案的安全性完全不依赖于 K_S 的加密情况, 因此服务器即使猜中口令甚至于知道 K_S , 也不能进行离线解密, 而文献 [4] 方案完全依赖于服务器独立对 K_S 的操作, 有被算法替换攻击的风险且缺少预防机制.

定理 2. 部分不经意性. 如果 DDH 假设成立, 则以上方案在随机谕言机模型下对频率限制器是部分不经意的.

证明: 概述: 本文方案保留了文献 [4] 方案等对频率限制器的操作, 因此自然地保留了部分不经意性的性质. 该部分证明与文献 [4] 类似.

定理 3. 可靠性. 如果 POK 是可靠的且是知识的证明, 则以上方案在随机谕言机模型下具有强可靠性.

证明: 概述: 要证明可靠性, 即敌手不能使服务器接受错误的论断或者从合法的密文中解密出另一个明文. 前者可直接由 POK 的可靠性得到. 上述方案与文献 [4] 方案不同就在于后者, 由于频率限制器参与了加密, 因此当频率限制器成为敌手的时候, 它可能尝试让服务器解密出不同的明文, 即回复一个 $c_{m'}$ 使得服务器得到 m' . 然而在不知道 K_S 的情况下, 伪造标签 $t = H_S(K_S \| E_{K_S}(m'))$ 成功的概率是可忽略的. 综上可靠性成立.

6 仿真与测试

为了进行详细的评估, 本文使用 Charm-Crypto 框架实现了方案, 以下测试在虚拟机 (4 GB 内存, 8 核处理器, 20 GB 磁盘存储的 64 位 Linux 操作系统) 中完成, 真机配置为 Intel(R) Core(TM) i7-7700HQ CPU @ 2.80 GHz (8 CPUs), 2.8 GHz, 8 GB 内存, Windows 10 家庭中文版 64 bit. 由于交互过程中不可预测因素太多, 如网络带宽、传输方式、数据流方向以及数据传输模式等, 故先不考虑交互过程.

首先评估替换攻击算法和原算法的时间成本, 本文通过使用 10000 轮的平均时间成本来评估算法的时间成本.

如图 5 所示, 在服务器加密算法中, 替换攻击算法略快, 这是因为直接哈希的时间消耗小于随机选取群元素的时间消耗, 后门攻击算法可以增加冗余计算使得实际消耗的时间接近, 从而抵御时间分析检测.

接下来进行本文能够抵抗恶意服务器的方案的效率与文献 [4] 方案效率的对比. 同样地, 通过使用 10000 轮的平均时间成本分别来评估服务器端加密与解密算法的时间成本.

如图 6 所示, 在服务器加密算法中, 本文方案具备了抵抗恶意服务器算法替换攻击的能力, 但是在服务器端增加的计算代价却很少, 在明文长度为 2 Mb 时, 服务器端加密的时间开销仅增加了 2.125 ms.

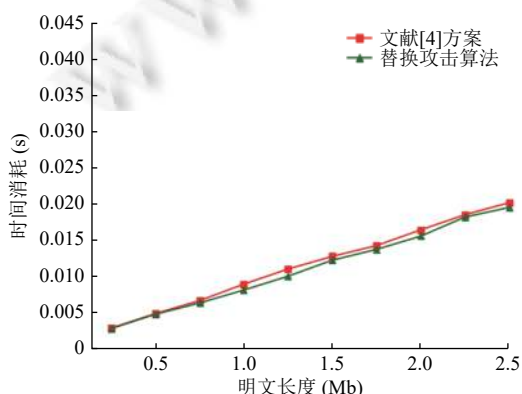


图 5 替换攻击方案和文献 [4] 方案服务器加密的时间开销对比

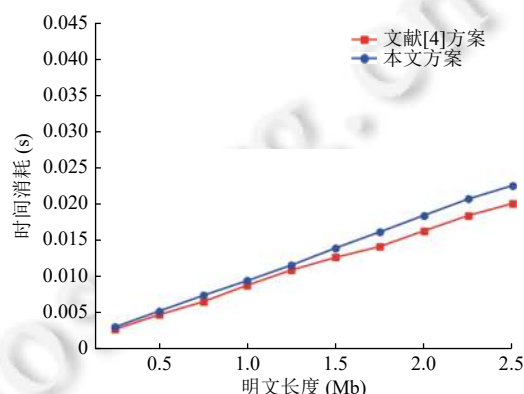


图 6 本文方案和文献 [4] 方案服务器加密的时间开销对比

如图 7 所示, 本文的方案中服务器端在解密时间开销相比于文献 [4] 方案中服务器端的解密时间开销相差不多, 在明文长度为 2 Mb 时, 服务器端解密的时间开销仅增加了 0.935 ms.

关于频率限制器, 在文献 [4] 方案中频率限制器的操作很少, 同样带来的安全性增强效果并不高. 我们将一些操作分配给了频率限制器, 将提升安全性所带来的代价大部分都转移到了频率限制器上, 使得网络结构的变化对于原有设备的负担并没有明显的增加, 能够增强用户使用密码增强服务的动机.

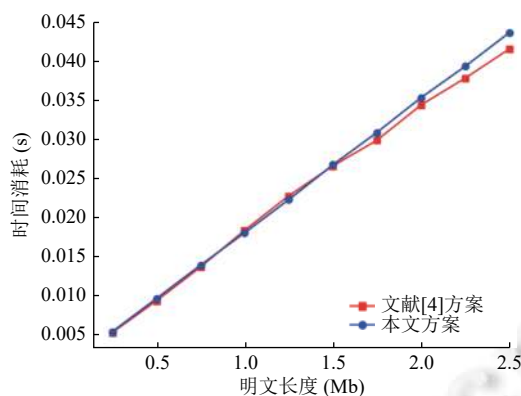


图7 本文方案和文献[4]方案服务器解密的时间开销对比

7 结 论

本文对口令增强加密这一服务中恶意服务器可能采取的攻击手段进行了分析,提出了对已有方案的替换攻击算法,使得已有方案不能够达到口令增强的效果.基于以上结果,本文进一步提出了能够抵御恶意服务器离线攻击的方案,同时给出了仿真测试结果,表明可以以较小的代价抵御恶意服务器的算法替换攻击.在这一研究中值得注意的经验是,有加密能力的用户总是能生成恶意的密文,因此需要另一方参与进行交互式加密.进一步地,本文推测,只要协议有一个参与方能够确保不会受到算法替换攻击,那么就可以设计一个不使用逆向防火墙的抵抗其他参与方受到算法替换攻击的协议版本.后续工作将进一步对该推断进行验证.

References:

- [1] Everspaugh A, Chatterjee R, Scott S, Juels A, Ristenpart T. The Pythia PRF service. In: Proc. of the 24th USENIX Security Symp. Washington: USENIX Association, 2015. 547–562.
- [2] Schneider J, Fleischhacker N, Schröder D, Backes M. Efficient cryptographic password hardening services from partially oblivious commitments. In: Proc. of the 2016 ACM SIGSAC Conf. on Computer and Communications Security. Vienna: ACM, 2016. 1192–1203. [doi: 10.1145/2976749.2978375]
- [3] Lai RWF, Egger C, Schröder D, Chow SSM. Phoenix: Rebirth of a cryptographic password-hardening service. In: Proc. of the 26th USENIX Conf. on Security Symp. Vancouver: USENIX Association, 2017. 899–916.
- [4] Lai RWF, Egger C, Reinert M, Chow SSM, Maffei M, Schröder D. Simple password-hardened encryption services. In: Proc. of the 27th USENIX Security Symp. Baltimore: USENIX Association, 2018. 1405–1421.
- [5] Li G, Liu JW, Zhang ZY. An overview on cryptography against mass surveillance. Journal of Cryptologic Research, 2019, 6(3): 269–282 (in Chinese with English abstract). [doi: 10.13868/j.cnki.jcr.000301]
- [6] Brost J, Egger C, Lai RWF, Schmid F, Schröder D, Zoppelt M. Threshold password-hardened encryption services. In: Proc. of the 2020 ACM SIGSAC Conf. on Computer and Communications Security. Virtual Event: ACM, 2020. 409–424. [doi: 10.1145/3372297.3417266]
- [7] Bellare M, Paterson KG, Rogaway P. Security of symmetric encryption against mass surveillance. In: Proc. of the 34th Annual Cryptology Conf. Santa Barbara: Springer, 2014. 1–19. [doi: 10.1007/978-3-662-44371-2_1]
- [8] Ateniese G, Magri B, Venturi D. Subversion-resilient signatures: Definitions, constructions and applications. Theoretical Computer Science, 2020, 820: 91–122. [doi: 10.1016/j.tcs.2020.03.021]
- [9] Liu C, Chen R, Wang Y, et al. Asymmetric subversion attacks on signature schemes. In: Proc. of the 23rd Australasian Conf. on Information Security and Privacy. Wollongong: Springer, 2018. 376–395. [doi: 10.1007/978-3-319-93638-3_22]
- [10] Al Mansoori F, Baek J, Salah K. Subverting MAC: How authentication in mobile environment can be undermined. In: Proc. of the 2016 IEEE Conf. on Computer Communications Workshops (INFOCOM WKSHPS). San Francisco: IEEE, 2016. 870–874. [doi: 10.1109/INFOCOMW.2016.7562200]
- [11] Armour M, Poettering B. Substitution attacks against message authentication. IACR Trans. on Symmetric Cryptology, 2019, (3): 152–168.

- [12] Russell A, Tang Q, Yung M, Zhou HS. Clipping the power of kleptographic attacks. In: Proc. of the 22nd Int'l Conf. on the Theory and Application of Cryptology and Information Security. Hanoi: Springer, 2016. 34–64. [doi: [10.1007/978-3-662-53890-6_2](https://doi.org/10.1007/978-3-662-53890-6_2)]
- [13] Russell A, Tang Q, Yung M, Zhou HS. Generic semantic security against a kleptographic adversary. In: Proc. of the 2017 ACM SIGSAC Conf. on Computer and Communications Security. Texas: ACM, 2017. 907–922. [doi: [10.1145/3133956.3133993](https://doi.org/10.1145/3133956.3133993)]
- [14] Russell A, Tang Q, Yung M, Zhou HS. Correcting subverted random oracles. In: Proc. of the 38th Annual Int'l Cryptology Conf. Santa Barbara: Springer, 2018. 241–271. [doi: [10.1007/978-3-319-96881-0_9](https://doi.org/10.1007/978-3-319-96881-0_9)]
- [15] Mironov I, Stephens-Davidowitz N. Cryptographic reverse firewalls. In: Proc. of the 34th Annual Int'l Conf. on the Theory and Applications of Cryptographic Techniques. Sofia: Springer, 2015. 657–686. [doi: [10.1007/978-3-662-46803-6_22](https://doi.org/10.1007/978-3-662-46803-6_22)]
- [16] Dodis Y, Mironov I, Stephens-Davidowitz N. Message transmission with reverse firewalls —Secure communication on corrupted machines. In: Proc. of the 36th Annual Int'l Cryptology Conf. Santa Barbara: Springer, 2016. 341–372. [doi: [10.1007/978-3-662-53018-4_13](https://doi.org/10.1007/978-3-662-53018-4_13)]
- [17] Chen RM, Mu Y, Yang GM, Susilo W, Guo FC, Zhang MW. Cryptographic reverse firewall via malleable smooth projective hash functions. In: Proc. of the 22nd Int'l Conf. on the Theory and Application of Cryptology and Information Security. Hanoi: Springer, 2016. 844–876. [doi: [10.1007/978-3-662-53887-6_31](https://doi.org/10.1007/978-3-662-53887-6_31)]
- [18] Zhou YY, Guo J, Li FG. Certificateless public key encryption with cryptographic reverse firewalls. Journal of Systems Architecture, 2020, 109: 101754. [doi: [10.1016/j.sysarc.2020.101754](https://doi.org/10.1016/j.sysarc.2020.101754)]

附中文参考文献:

- [5] 李耕, 刘建伟, 张宗洋. 抗大规模监视密码学研究综述. 密码学报, 2019, 6(3): 269–282. [doi: [10.13868/j.cnki.jcr.000301](https://doi.org/10.13868/j.cnki.jcr.000301)]



赵一(1985—), 男, 博士, 讲师, 主要研究领域为公钥密码学, 抗后门攻击, 隐私保护.



明洋(1979—), 男, 博士, 教授, 博士生导师, CCF 专业会员, 主要研究领域为密码学, 网络安全.



刘行(1999—), 男, 博士生, 主要研究领域为公钥密码学, 区块链技术.



赵祥模(1966—), 男, 博士, 教授, 博士生导师, CCF 高级会员, 主要研究领域为车路协同与自动驾驶技术, 车联网与智能网联汽车测试技术.



LIANG Kaitai(1985—), 男, 助理教授, 主要研究领域为信息安全与密码学, 协议安全分析.



杨波(1963—), 男, 博士, 教授, 博士生导师, 主要研究领域为密码学, 信息安全.