

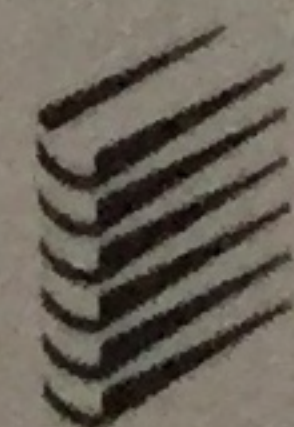
2018

金海 —— 主编
徐鹏 邹德清 —— 副主编

中国网络空间安全 前沿科技发展报告 2018



中国工信出版集团



人民邮电出版社
POSTS & TELECOM PRESS

本书编委会

主 编 金 海

副主编 徐 鹏 邹德清

委 员 (按姓氏拼音排序)

陈 洁 陈 晶 陈 恺 陈荣茂 陈 玺 陈晓峰 陈宇飞

程 鹏 段海新 高 飞 葛春鹏 巩俊卿 韩劲松 韩伟力

何道敬 何德彪 胡红钢 胡 志 黄 琼 黄欣沂 纪守领

冀晓宇 赖俊祚 李 进 李 琦 李 珍 林璟铎 刘 炅

刘 哲 罗向阳 屈龙江 任 奎 沈 超 沈 剑 孙 兵

汪 定 汪京培 王 聪 王 磊 温金明 翁 健 肖 亮

徐文渊 许封元 杨 珉 禹 勇 郁 昱 张 超 张 甲

张 江 张 磊 张明武 赵运磊 郑佳佳 周亚金 朱浩瑾

诸葛建伟

推荐序

网络空间安全是国家安全的重要组成部分，人才是网络空间安全学科发展的重要基础。为加强网络空间安全学科建设、加快网络空间安全人才培养、促进网络空间安全科学技术发展，华中科技大学、武汉市网信办、国家网络安全人才与创新基地和武汉网络安全战略与发展研究院共同创办了“网络空间安全喻园青年科学家论坛”。该论坛已经连续举办两届，共邀请来自国内知名高校的 83 位杰出青年学者参与论坛并做特邀报告。论坛具有受邀嘉宾年轻化、演讲主题前沿化、研讨内容专业化、研究方向多样化、交流讨论对等化等鲜明特色，为广大青年学者提供一个有利于把握学科建设需求、有利于展示科学研究成果、有利于探讨学术发展方向、有利于多安全方向交叉的互动平台。

《中国网络空间安全前沿科技发展报告》作为“网络空间安全喻园青年科学家论坛”的成果与总结，由来自全国 27 所高校（科学院）的 59 位杰出青年学者联合倾力编写完成，内容丰富，涵盖了可信计算、密码学、网络安全、物联网安全、云计算与云数据安全、大数据安全、区块链和人工智能等领域的前沿科技。报告中的每篇文章都凝聚了一位或多位优秀科研工作者近年来的科研心得，提纲挈领地阐述了相关领域内的当前研究背景、简明扼要地分析了相关技术在国内外的研究现状，深入浅出地介绍了学者在各自领域的研究成果。每篇文章的最后还启发式地给出了各位优秀学者对各自领域的研究展望与宝贵建议，这为网络空间安全领域的广大学者或学子们提供了很好的借鉴价值。同时，该报告对于各级国家与地方主管部门参考制定网络空间安全发展战略、明确网络空间安全领域的人才培养方向也具有重要参考价值。

报告内涵丰富，既有科普性，也具启发性。阅读本报告，可以了解目前网络空间安全领域面临的各种科学问题与挑战以及众多前沿科技的国内外研究现状、发展特色。

中国工程院院士

沈昌祥

2018 年 11 月 20 日于北京

伪随机函数与密码算法.....	1
伪随机函数研究.....	3
对称密码算法研究.....	7
可修订的数字签名研究.....	17
功能加密的理论的实现.....	21
全同态加密若干关键问题的研究.....	24
基于格的后量子密码研究.....	27
量子计算密码分析.....	40
恶意系统或算法下的安全防护.....	45
旁路攻击防护技术研究.....	47
可证明安全的程序混淆.....	51
抗颠覆攻击的若干关键技术研究.....	55
通用计算平台的密钥安全机制.....	59
数据驱动的口令定向猜测技术研究.....	64
网络安全.....	69
网络空间中网络和系统安全攻防对抗核心技术.....	71
频谱安全技术研究现状与建议.....	76
软件定义的网络防御技术.....	81
网络空间测绘技术.....	85
基于强化学习的无线网络安全技术研究.....	88
物联网系统安全.....	91
基于模拟态信号的物联网安全检测和认证研究.....	93
内生安全的工业控制系统协同防护技术研究.....	97
物联网软件系统安全关键技术.....	101
面向星载综合电子系统的入侵检测技术研究.....	105
信息物理融合系统综合安全.....	110

物联网可信认证、保密传输与终端安全

人机物一体化物联网可信安全认证

人机行为特征识别与终端身份安全

车联网中的信息安全传输

面向移动智能系统的安全漏洞挖掘方法

移动终端感知安全

云计算与云数据安全

云计算安全

复杂网络环境中数据安全关键技术研究

外包数据完整性保障技术研究

基于代理重加密的云数据安全共享方法研究

云环境下数据安全保护关键技术研究

大数据隐私保护与区块链

大数据环境下的用户隐私保护技术

面向大数据的安全加密搜索

可搜索加密技术的研究现状与建议

区块链隐私保护机制研究

基于区块链的统一身份认证与管理

人工智能安全

深度学习模型安全

机器学习系统安全

机器学习中的隐私保护与样本安全研究

语音领域的人工智能安全

人工智能系统推理安全

软件源代码漏洞智能检测

数据驱动的口令定向猜测技术研究

汪 定

北京大学

1. 研究背景

身份认证是保障信息系统安全的第一道防线，口令是应用最为广泛的身份认证方法^[1,2]。随着网络技术的深入发展，一方面越来越多的网络服务需要口令保护，另一方面人类大脑认知能力有限，只能记忆 5~7 个口令^[3]，导致用户不可避免地使用低信息熵的弱口令^[4]，在口令中使用个人相关信息^[5]，在多个网站中重用同一口令^[6]，带来严重的安全威胁。2004 年，比尔·盖茨^[7]对外宣告口令将消亡，微软公司将使用多因子认证替代纯口令认证。后续一系列研究(如[8,9])也分别从口令无法抵抗离线猜测攻击、口令过期策略无法保证更新后的口令的不可预测性等方面论证了口令认证技术的不可持续性。与此同时，大量替代口令的认证技术不断被提出，如图形口令、行为认证、多因子认证等等。

出乎意料的是，时至今日，**口令的地位在工业界不仅丝毫没有被动摇，并且在越来越多的信息系统应用中得到加强**。这一现象吸引了越来越多的学者，开始引起学术界的反思。研究发现，虽然这些替代型方案有的在安全性方面优于口令，有的在可用性方面胜过口令，但几乎都在可部署性方面劣于口令，并且各自存在一些固有的缺陷^[10]。比如，基于硬件的认证技术成本高昂，使用不方便；基于生物的认证技术不具有可撤销性，且存在隐私泄露问题。因此，学术界逐渐开始形成一个共识^[11,12]：**在可预见的未来，基于口令的身份认证技术仍将是主要的用户认证方式**。

当前，口令认证系统面临的**最大安全威胁是口令定向猜测攻击**。根据 VeriZon 的最新报告^[13]，2017 年发生的数以千计的已经确认数据泄露事件中，21.95%是因为攻击者获取了攻击目标的口令，并且这一数字呈逐年上升趋势。根据攻击过程中是否利用用户个人信息，口令猜测算法可分为漫步攻击和定向攻击^[14]。漫步猜测主要利用了用户倾向于选择流行口令的行为，而定向猜测不仅利用普通用户使用流行口令的脆弱行为，而且会利用用户使用个人信息构造口令和重用口令的脆弱行为。由于定向猜测下允许的攻击者能力更强大，在相同猜测次数下，定向猜测会大大提高猜测成功率。**随着大规模个人信息泄露事件的不断发生(如[15-17])，各种类型的个人身份信息和用户在其它网站使用的口令都越来越容易被攻击者获取，定向猜测逐渐成为一个十分严峻的现实威胁**。

2. 该研究存在的科学问题

(1) 口令中异构个人身份信息的感知问题

用户的个人身份信息 PII 不仅多种多样，而且在结构和功能方面严重异构。比如，姓名和爱好由字母组成，生日一般由数字组成；姓名可直接作为口令的组成部分，而性

别会影响口令但不直接构成口令组成部分。PII 的多样性要求定向猜测算法具有自主感知能力和可扩展性，而 PII 的异构性要求算法具有隐语义感知能力和自适应性。

(2) 用户口令修改行为的上下文环境感知问题

用户在修改现有口令时，有一系列各异的变换规则可供使用，如插入、删除、大小写变换、字符跳变(如 password→passw0rd)等等，这些变换规则也可以综合使用(如 password→Passw0rd1)。用户选择什么样的变换规则，往往与上下文环境（服务类型、口令策略等）相关。针对每个不同用户，如何将这些规则进行优先级排序是困难的。

(3) 小训练样本下的口令定向猜测模型构建问题

移动设备、网银、虚拟货币交易帐户、WiFi 接入等典型特殊应用环境下的口令与普通 Web 应用环境下的口令，有明显差异。这些特殊应用环境下的口令数据由于本地存储（如移动设备 PIN 码）、分散存储（如 WiFi 接入点的口令）和强健安全存储（如网银口令）等原因，极少像 Web 应用那样大规模地泄露，导致相关的公开口令数据集严重不足。如何在真实口令数据匮乏的情况下，设计适于特殊应用环境特点的、高效的、自适应的、可识别 PII 语义的定向口令猜测模型是一个关键问题。

3. 国际研究现状

国际上，关于定向口令猜测的研究刚刚起步，仅有少数几个算法，较知名为 2016 年 Li 等提出的 Personal-PCFG 算法。Personal-PCFG 基于 Weir 等提出的漫步算法 PCFG。

PCFG。2009 年，Weir 等提出了基于概率上下文无关文法(PCFG)的漫步口令猜测算法。该算法的核心假设是，口令的字母段 L、数字段 D 和特殊字符段 S 是相互独立的。它首先将口令根据前述三种字符类型进行切分，比如“wang123!”被切分为 L₄: wang, D₃: 123 和 S₁!:，L₄D₃S₁ 被称为该口令的模式。与 Markov 类似，该算法也分训练和猜测集生成两个阶段。

在训练阶段，最关键的是统计出口令模式频率表和字符组件（语义）频率表。基于上下文无关文法，对泄漏口令进行统计，得到各种模式的频率，和模式中数字组件、特殊字符组件的频率。比如针对 L₄D₃S₁，统计在全部口令中以 L₄D₃S₁ 为模式的口令频率，以及“wang”在长为 4 的字母段的频率，“123”在长为 3 的数字串中的频率和“!”在长为 1 的特殊字符串中的频率。

在猜测集生成阶段，依据上面获得的模式频率表和语义频率表，生成一个带频率猜测的集合，以模拟现实中口令的概率分布。比如，猜测“wang123!”的概率计算为 $P(\text{wang123!}) = P(\$ \rightarrow L_4 D_3 S_1) * P(L_4 \rightarrow \text{wang}) * P(D_3 \rightarrow 123) * P(S_1 \rightarrow !)$ = 0.15*0.3* 0.6*0.3 = 0.0081。这表明，“wang123!”的可猜测度为 0.0081。

Personal-PCFG。2016 年，Li 等提出了定向攻击猜测算法 Personal-PCFG。该算法基于漫步 PCFG 攻击模型，基本思想与 PCFG 攻击模型完全相同：将口令按字符类型按

长度进行切分。为实现这一思想, Personal-PCFG 首先将 PII 分为 6 大类(即用户名 A、邮箱前缀 E、姓名 N、生日 B、手机号 P 和身份证 G), 并将这 6 种 PII 字符类型视为与漫步 PCFG 模型里的 L、D、S 同等地位, 这样 Personal-PCFG 中有 9 种类型字符; 接着, 在训练过程中, 将训练集中每个口令如同漫步 PCFG 攻击模型那样, 按相应字符类型及其长度进行分段, 比如“wang123!”被切分为 $N_4D_3S_1$ 分段, 因为“wang”属于姓名 N 这一大类, 长度为 4。剩余训练过程与漫步 PCFG 模型类似。

猜测集生成阶段分两步。第一步, 运行漫步 PCFG 模型的猜测集生成过程, 产生中间猜测集。该猜测集既包含“123456”这样的可以直接使用的猜测, 也会包含带 PI 类型字符的“中间猜测”(如 N_1 , N_2123); 第二步将“中间猜测”里的 PII 类型字符替换为攻击对象的相应长度的 PII 信息, 如将 N_3123 替换为 wan123, ang123, ang123, $\dots$, wei123(假设攻击目标姓名为“wang lei”)。

4. 我国针对该研究的一流成果描述

我国关于定向口令猜测的研究, 虽然起步较晚, 但 TarGuess 攻击理论^[5]于 ACM CCS 2016 会议上一经提出, 立即产生了积极而广泛的国际影响力。TarGuess 是一个口令定向猜测攻击理论框架, 包含 7 个理论攻击模型 TarGuess I~VII, 以刻画 7 种不同现实能力的定向猜测攻击者。限于篇幅, 这里仅以 TarGuess-I 为例说明。

TarGuess-I 的目标是通过利用用户 U 的一些 Type-1 PII(如姓名和生日), 对 U 的口令进行猜测。它建立在前述 Weir 等提出的 PCFG 模型之上。TarGuess-I 的基本思想是: 人群中有多少比例使用某种可人可 PII, 那么攻击对象也将有同样可能性(比例)使用该 PI。为实现这一思想, 精确感知口令中的 PII 语义, 除了 PCFG 模型中定义的 L、D、S 标签外, TarGuess-I 还定义了一系列基于个人信息类型的 PII 标签(如 $N_1 \sim N_7$ 和 $B_1 \sim B_{10}$)。

对于每一个 PII 标签, 它的下标数字代表这个类型的 PII 用法的一个细分(即子类型标签), 而不是像 L、D、S 标签或者 Personal-PCFG 算法中定义的 6 种 PII 类型字符那样, 下标数字表示相应的长度。比如, N 代表姓名信息的用法, 而 N_1 表示姓名全称, N_2 表示姓名全称的缩写(如“wang lei”缩写为 lw), $\dots$; B 代表生日信息的用法, B_1 表示以年月日的格式使用生日(如 19820607), B_2 表示以月日年的格式使用生日, $\dots$ 。这样, 就产生了一个可感知 PII 的上下文无关文法。

在训练阶段, 首先按照最长前缀(the longest prefix)的原则对训练集里每个口令中所含有的 PII 信息进行匹配并替换为相应的 PII 标签, 并且只考虑长度大于 2 的 PII 字段。训练阶段的剩余步骤与漫步 PCFG 模型相同。

猜测集生成阶段分两步。第一步, 运行漫步 PCFG 模型的猜测集生成过程, 产生中间猜测集, 该猜测集既包含“123456”这样的可以直接使用的猜测, 也会包含带 PII 标签的“中间猜测”(如 N_1 , N_2123); 第二步将“中间猜测”里的 PII 基本字符替换为攻

击对象的相应 PII 信息。比如，将 N₂123 替换为 wang123（假设攻击对象姓名为“wang lei”），因为 N₂ 表示姓氏“wang”，而不是像 Personal-PCFG 替换为 wa123, an123, ng123, ... ei123。

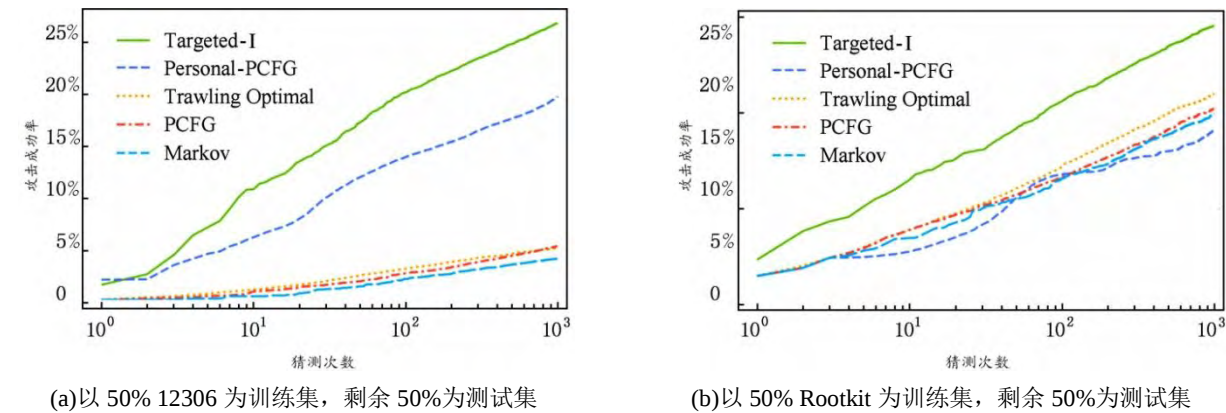


图 1 定向猜测攻击算法对比^[2]

图 1 对 Targeted-I 和 Personal-PCFG 这两个定向口令猜测攻击算法进行了对比。同时，为了更好显示定向攻击算法在在线猜测方面的优势，我们将定向攻击算法与漫步攻击算法（PCFG、Markov 和 Trawling-optimal）也进行了对比。结果显示，在猜测数为 10~1000 时，Targeted-I 的攻击效率比 Personal-PCFG 高出 37.18%~73.29%，比理想漫步算法（Trawling-optimal）高出 412%~740%。

5. 我国未来可进一步开展的研究建议

当前，国内外关于口令定向猜测虽然有一些阶段性成果，但仍属于起步阶段，在很多重要方向上亟待系统性、规律性、原则性的研究。其中，有下面三个首要的方向：

（1）可自主学习的基于 PII 的定向猜测

现有的定向猜测算法都需要人工对各类异构的 PII 信息进行识别并枚举定义，因此仅仅关注了姓名、生日等少数显式的 PII 类别。并且，PII 分类的粒度和 PII 标签定义的正确性严重依赖于算法设计者的攻击经验。因此，研究如何自主学习（识别）口令中各类显式和隐式 PII 的使用频率，实现自适应的定向口令猜测模型，是一个核心问题。

2）可感知上下文环境的基于多个姊妹口令的定向猜测

一方面，随着大规模泄露事件的不断发生，越来越多的用户会泄露两个以上口令。另一方面，用户在重用口令时，有一系列各异的变换规则可供用户选择，如插入、删除、大小写变换等等。用户会选择什么样的口令变换规则，往往与上下文环境相关，如网站的口令生成策略。针对口令变换规则的多样性和上下文相关性，设计基于多个姊妹口令的设计跨站猜测模型具有十分重要的现实意义。

3）针对特殊应用环境下的定向在线口令猜测

移动设备、网银、虚拟货币交易帐户、WiFi 接入等典型特殊应用环境下的口令与普通 Web 应用环境下的口令,有明显差异。比如,移动设备认证往往使用纯数字组成的 PIN 码;网银口令保护极端敏感资源,但往往与额外认证因子(如 U 盾)联合使用;虚拟货币交易帐户口令虽保护极端敏感资源,但往往像传统 Web 服务那样单独使用;WiFi 接入点往往要求口令长度达 12 位,且口令往往需要与他人共享使用。但另一方面,这些特殊应用环境下的口令数据极少像 Web 应用那样大规模地泄露,导致当前关于这些特殊应用环境下的公开口令数据集却严重不足。如何在特殊应用环境中真实口令数据匮乏的情况下,设计高效的、自适应的定向猜测模型是一个有前景的研究方向。

参考文献

- [1] Joseph Bonneau, Cormac Herley, Paul C. van Oorschot, Frank Stajano. Passwords and the evolution of imperfect authentication[J]. Communications of the ACM, 2015, 58(7): 78-87.
- [2] 王平, 汪定, 黄欣沂. 口令安全研究进展[J]. 计算机研究与发展, 2016, 53(10): 2173-2188.
- [3] Mark Keith, Benjamin Shao, Paul John Steinbart. The usability of passphrases for authentication: An empirical field study[J]. International Journal of Human-Computer Studies, 2007, 65(1): 17-28
- [4] Joseph Bonneau. The Science of Guessing: Analyzing an Anonymized Corpus of 70 Million Passwords[C]. Proc. IEEE S&P 2012, pp. 538-552
- [5] Ding Wang, Zijian Zhang, Ping Wang, Jeff Yan, Xinyi Huang. Targeted Online Password Guessing: An Underestimated Threat[C]. Proc. ACM CCS 2016, pp. 1242-1254
- [6] APal, Bijeta, Tal Daniel, Rahul Chatterjee, and Thomas Ristenpart. Beyond Credential Stuffing: Password Similarity Models using Neural Networks. Proc. IEEE S&P 2019, pp. 1-18
- [7] Gates predicts death of the password[EB/OL], Feb. 2004. <http://www.cnet.com/news/gates-predicts-death-of-the-password/>
- [8] Luke S., Lisa J., William E., Matthew P., Patrick T., Patrick M., Trent J.. Password exhaustion: Predicting the end of password usefulness[C]. Proc. ICISS 2006, pp. 37-55
- [9] Yinqian Zhang, Fabian Monrose, Michael Reiter. The security of modern password expiration: an algorithmic framework and empirical analysis[C]. Proc. ACM CCS 2010, pp. 176-186
- [10] Joseph Bonneau, Cormac Herley, et al. The Quest to Replace Passwords: A Framework for Comparative Evaluation of Web Authentication Schemes[C]. Proc. IEEE S&P 2012, pp. 553-567
- [11] Nasir Memon. How biometric authentication poses new challenges to our security and privacy[J]. IEEE Signal Processing Magazine, 2017, 34(4): 196-194.
- [12] Cormac Herley, Paul Van Oorschot. A research agenda acknowledging the persistence of passwords[J]. IEEE Security & Privacy, 2012, 10(1): 28-36
- [13] 2018 Data Breach Investigations Report [EB/OL], April 2019. http://www.documentwereld.nl/files/2018/Verizon-DBIR_2018-Main_report.pdf
- [14] Ding Wang, Debiao He, Haibo Cheng, Ping Wang. fuzzyPSM: A New Password Strength Meter Using Fuzzy Probabilistic Context-Free Grammars. Proc. of the 46th IEEE/IFIP International Conference on Dependable Systems and Networks (IEEE/IFIP DSN 2016), pp. 695-606.
- [15] AJ Dellinger. Understanding The First American Financial Data Leak: How Did It Happen And What Does It Mean? May 2019. <https://www.forbes.com/sites/ajdellinger/2019/05/26/understanding-the-first-american-financial-data-leak-how-did-it-happen-and-what-does-it-mean/#7774cc00567f>
- [16] MORE THAN 20 MILLION INDIVIDUALS IMPACTED BY AMCA DATA BREACH.. June 2019. <https://www.hipaaguide.net/more-than-20-million-individuals-impacted-by-amca-data-breach/>
- [17] <https://breachalarm.com/sources>