

# 针对品牌的网络钓鱼扩线与检测方案

文伟平, 朱一帆, 吕子晗, 刘成杰

(北京大学软件与微电子学院, 北京 100080)

**摘要:**近年来, 无论是钓鱼攻击的数量还是其造成的损失都在不断增加, 网络钓鱼攻击成为人们面临的主要网络安全威胁之一。当前, 已有许多网络钓鱼检测方法被相继提出以抵御网络钓鱼攻击, 但现有钓鱼检测方法多为被动检测, 且容易引起大量误报。针对上述问题, 文章提出一种钓鱼扩线方法。首先, 根据钓鱼网站的信息从多维度进行分析, 得到与之相关的其他网站, 以找到更多还未被发现的钓鱼网站; 然后, 针对钓鱼网站的视觉仿冒特性, 提出一种基于深度学习的网络钓鱼检测方法, 将截图进行切割, 得到判定为 logo 的区域, 再使用 EfficientNetV2 挖掘视觉仿冒特性; 最后, 对疑似钓鱼网站进行综合评价, 以降低误报率。通过对现有钓鱼网站进行实验验证, 证明了文章所提方法的有效性。

**关键词:** 网络钓鱼检测; 深度学习; 图片切割; 主动扩线

**中图分类号:** TP309 **文献标志码:** A **文章编号:** 1671-1122 (2023) 12-0001-09

中文引用格式: 文伟平, 朱一帆, 吕子晗, 等. 针对品牌的网络钓鱼扩线与检测方案 [J]. 信息安全, 2023, 23 (12): 1-9.

英文引用格式: WEN Weiping, ZHU Yifan, LYU Zihan, et al. Brand-Specific Phishing Expansion and Detection Solutions[J]. Netinfo Security, 2023, 23(12): 1-9.

## Brand-Specific Phishing Expansion and Detection Solutions

WEN Weiping, ZHU Yifan, LYU Zihan, LIU Chengjie

(School of Software and Microelectronics, Peking University, Beijing 100080, China)

**Abstract:** In recent years, both the number of phishing attacks and the losses caused by them have been increasing, and phishing attacks have become one of the main network security threats that people face. Currently, many phishing detection methods have been proposed to defend against phishing attacks, but most of the known phishing detection methods are passive detection and are prone to cause a large number of false positives. In response to the above issues, this paper proposed a phishing expansion method. Firstly, according to the phishing website information, it was analyzed in a multi-dimensional manner, and other related websites were obtained, so as to find more phishing websites that have not been discovered yet. Then, aiming at the visual counterfeiting characteristics of phishing websites, this paper proposed

收稿日期: 2023-08-15

基金项目: 国家自然科学基金 [61872011]

作者简介: 文伟平 (1976—), 男, 湖南, 教授, 博士, 主要研究方向为系统与网络安全、大数据与云安全、智能计算安全; 朱一帆 (1997—), 男, 湖北, 硕士研究生, 主要研究方向为网络与系统安全; 吕子晗 (1996—), 男, 河北, 硕士研究生, 主要研究方向为网络流量分析和恶意行为识别; 刘成杰 (1998—), 男, 湖南, 博士研究生, 主要研究方向为软件安全、漏洞挖掘和入侵检测。

通信作者: 文伟平 weipingwen@pku.edu.cn

a phishing detection method based on deep learning, cutting the screenshots to obtain the area judged as a logo, and using EfficientNetV2 to mine visual counterfeiting characteristic. Finally, conducted a comprehensive evaluation of suspected phishing websites to reduce the false positive rate. The effectiveness of the method proposed in this paper was proved by the experimental verification of the existing phishing websites.

**Key words:** phishing detection; deep learning; picture cutting; expansion of phishing sites

## 0 引言

互联网已经成为人们生活中不可或缺的一部分,截至2020年,互联网活跃用户数约有53亿人。随着互联网技术的广泛应用,网络攻击的频率也不断增加。反网络钓鱼工作组APWG (Anti-Phishing Working Group) 发布的《2022年第四季度网络钓鱼活动趋势报告》显示,2022年攻击事件超过470万次。2020年,Verizon数据泄露调查报告显示,67%的数据泄露事件是由网络钓鱼或电子邮件网络钓鱼等社会工程攻击造成的<sup>[1]</sup>。社会工程是指在信息安全方面利用人的心理,使其采取行动或泄露机密信息<sup>[2]</sup>,以获取用户数据甚至骗取财产。网络钓鱼攻击是一种典型的社会工程攻击,网络钓鱼者通过模仿网站来冒充受信任的品牌,并要求受害者提供个人敏感信息,使用户产生巨大的经济损失。

现有的网络钓鱼检测技术方案主要基于黑名单、分类和截图设计解决方案<sup>[3]</sup>。基于黑名单的解决方案(如Google Safe Browsing<sup>[4]</sup>和OpenPhish),使用动态黑名单跟踪报告的钓鱼网址。基于分类的解决方案,在收集的网络钓鱼数据集上进行训练,使用从给定网页及其统一资源定位器(Uniform Resource Locator, URL)中提取的特征对网络钓鱼进行二进制预测<sup>[5-8]</sup>。虽然VERMA<sup>[5]</sup>等人认为基于URL特征的二进制预测钓鱼网站有望实现。但是在现实中,随着网络钓鱼工具包的不断发展,钓鱼工具包的创建通常是自动化的。用于训练的黑名单和收集的钓鱼网页很快就会过时,因此,大量没有模仿目标品牌URL的钓鱼网站无法被基于URL特征的二进制预测方法检测。

本文提出一种基于截图的网络钓鱼检测解决方案,

网络钓鱼攻击的主要手段是用视觉上类似于合法网站的页面来欺骗用户。因此,一些研究者提出一种基于截图相似性的检测方法,早期的网络钓鱼识别方案将给定网页的屏幕截图与参考数据库中所有网页的截图<sup>[9,10]</sup>进行比较。例如,FU<sup>[9]</sup>等人提出使用地球移动器距离(Earth Mover's Distance, EMD)技术来计算两个网页截图的相似度,但这种方法受限于网页内容频繁更新,且许多钓鱼网站与合法网站的整体截图相似性不高,导致准确率较低。因此,最近的研究转向使用品牌logo进行网络钓鱼识别<sup>[11-14]</sup>。本文采用Detectron2框架下的Faster R-CNN模型,对截图进行切割识别,获取截图中的logo部分,并将其与品牌列表中及已知钓鱼网站的截图进行对比,以提高检测的准确率和效率。

考虑受保护品牌列表变化时分类模型需要重新进行训练,本文采用训练速度快、准确率高的EfficientNetV2模型,以适应实际应用需求,降低时间成本。基于logo的方法可能会在某些场景下产生误报,例如,一个良性网页使用淘宝logo但该logo不在其域名列表中,可能被误判为钓鱼网站。针对这种情况,本文设计了一种网页评价方法,对疑似钓鱼网站进行多角度评价,包括Page Rank值、Alexa排名等,以防止误报。

## 1 相关工作

### 1.1 网络钓鱼

网络钓鱼诈骗包括利用欺骗性电子邮件、网站和短信等手段,旨在激发受害者的急迫感或好奇心,诱使受害者访问仿冒的钓鱼网站或打开含恶意软件的附件,进而透露敏感信息。这些钓鱼网站在外观和内容上与信誉良好的网站极为相似,导致受害者输入敏感信息。一旦信息被提交,将会泄露给攻击者。

使用钓鱼网站的主要攻击类型包括网络钓鱼、偷渡式下载和垃圾邮件。网络钓鱼通过伪装的网页欺骗用户泄露私人信息或敏感信息；偷渡式下载是用户在访问链接时无意中下载了恶意软件；垃圾邮件是未经请求发送的、出于广告或网络钓鱼目的的邮件。钓鱼网站的存在对网络信息安全构成了巨大威胁，对用户的隐私安全造成了严重损害，因此研究出能够及时准确地检测到钓鱼网站的方法显得尤为迫切<sup>[15]</sup>。

在反钓鱼手段方面，研究者也在不断开发多种方法以对抗网络钓鱼的威胁。现有的网络钓鱼检测方法大致可以分为非机器学习方法和机器学习方法。非机器学习方法包括黑名单方法、视觉相似性检测等；机器学习方法包括AdaBoost、朴素贝叶斯等，其中深度学习是机器学习领域中一个新的研究方向，其涵盖卷积神经网络(Convolutional Neural Network, CNN)、循环神经网络等。

## 1.2 基于深度学习的钓鱼网站识别技术

随着深度学习技术的发展，一些研究人员开始探索基于深度学习的网络钓鱼检测方法<sup>[16]</sup>。例如，BU<sup>[17]</sup>等人提出一种深度自动编码器模型，用于检测零日网络钓鱼攻击，该模型主要从URL字符串中提取字符级特征。SOMESHA<sup>[18]</sup>等人引入深度学习模型，使用从HTML页面和第三方服务中提取的10个特征来检测网络钓鱼网站，另外对比了3种不同的深度学习模型，并对18个特征进行权重计算。ADEBOWALE<sup>[19]</sup>等人采用CNN和长短期记忆(Long Short-Term Memory, LSTM)算法来构建一个混合分类模型，该模型从Phish Tank和Common Crawl收集URL，并从网站中提取图像、文本和框架等网站内容。其中，图像特征被用于离线CNN模型，而文本特征被用于LSTM分类器。该方案的创新之处在于它结合了图像和文本的特点，为网络钓鱼检测提供了一个多维度的视角。

## 2 针对品牌的网络钓鱼防护方案

本文提出的方案主要包括钓鱼网站扩线模块、钓鱼网站识别模块和钓鱼网站评价模块，整个过程如图1所示。当收到报告的钓鱼网站URL列表后，首先，通

过分析钓鱼网站扩线模块，找出更多可疑网站；然后，通过钓鱼网站识别模块对扩线模块输出的可疑网站列表进行筛选，识别出更可能是钓鱼网站的候选；最后，钓鱼网站评价模块对识别出的可疑网站进行进一步的分析和评估，确定哪些是真正的钓鱼网站。对于新发现的钓鱼网站，系统将再次启动钓鱼网站扩线模块，进行新一轮的钓鱼网站发现和确认。这种设计使得钓鱼网站的检测更加全面和系统，可以有效提高识别准确率和效率。

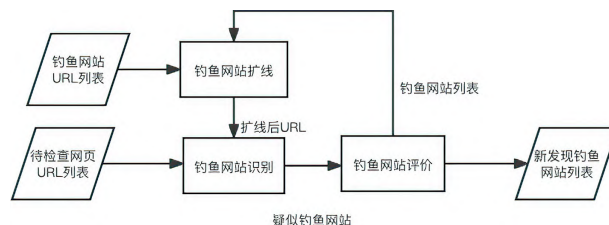


图1 网络钓鱼防护方案

### 2.1 钓鱼网站扩线

本文所设计的钓鱼网站扩线模块算法如算法1所示，旨在通过分析已知钓鱼网站的特征来发现新的钓鱼网站，输入为已知的钓鱼网站URL列表。在获取每个钓鱼网站的标题和网址信息后，进行以下3个步骤。步骤一：URL分析。如果URL是IP地址，则通过如virtual total和奇安信等工具来分析和获取相关的域名；如果URL是域名，则获取其对应的IP地址，从而找到关联域名和子域名。步骤二：whois信息分析。获取钓鱼网站的whois信息，如果注册邮箱和注册人不在白名单内（即不属于品牌列表下的URL的注册邮箱和注册人，或者通用的注册邮箱和注册人），则通过whois反查，获取同一注册邮箱和注册人注册的其他域名和IP，进而获取相关域名和子域名。步骤三：fofa查询。使用fofa查询功能，根据钓鱼网站的标题搜索其他IP地址和域名。将从步骤一至步骤三中获取的IP、域名以及它们的标题和截图输出，为下一步的筛选工作做准备。钓鱼网站扩线模块的核心在于通过分析和对比钓鱼网站的特征，寻找具有相似特征的其他网站，从而有效识别出新的钓鱼网站。

#### 算法1 钓鱼网站扩线模块算法



```

Input :InurlList
Output: OuturlList Title Pic
for each url in InurlList do //step1:url 分为 ip 和域名两种情况
    gettitle // 获取标题
    if url==url.ip then // 如果为 ip, 则获取相同 ip 下其他相关
        联的域名及其子域名
        OuturlList ← FindDomain,FindSub
    else if url == url.domain then // 如果为 domain, 获取相应
        ip, 然后得到其他相关联域名和子域名
        OuturlList ← FindDomain,FindSub,Findip
    if (Email, Person)=getwhois(url) not in whitelist then //step2:
        获取 url 的 whois 信息, 找到相关联 url
        OuturlList ← getsame(whoisEmail,whoisPerson)
if title not null then
    OuturlList ← fofa(title)//step3: 通过 fofa 找到与钓鱼网
    站拥有相同 title 的 url
    getPic
    return OuturlList Title Pic
    
```

## 2.2 钓鱼网站识别

钓鱼网站识别模块如图2所示,该模块将URL、描述合法品牌logo及其网络域的目标品牌列表作为输入,将生成的可疑钓鱼网站(如果URL在识别模块被判断为网络钓鱼)作为输出。网络钓鱼识别任务可以被分解为目标切割任务和图像识别任务。首先,使用目标切割模型,将截图中的logo切割出来;然后,通过CNN模型将检测到的logo与目标品牌列表中的logo进行比较,从而识别出网络钓鱼目标;最后,如果检测到的logo与目标品牌列表中的logo匹配,表明其已通过白名单筛选,可将其URL视为可疑钓鱼网站输出。

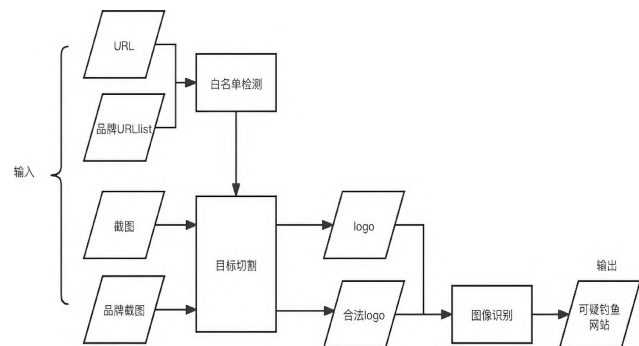


图2 钓鱼网站识别模块

### 2.2.1 目标切割任务

本文选择Detectron2平台的Faster R-CNN模型实现目标切割功能,图3展示了Faster R-CNN模型的结构。模型的输入为一个屏幕截图,该截图首先通过ResNet50<sup>[20]</sup>转换成形状为 $M \times M \times c$ 的特征图,其中 $M$ 表示特征图的大小, $c$ 表示通道数。ResNet50网络采用残差网络结构,该结构通过引入残差连接来解决深层网络中的退化问题,使得模型能够在保持网络深度的同时,提升训练的精度。

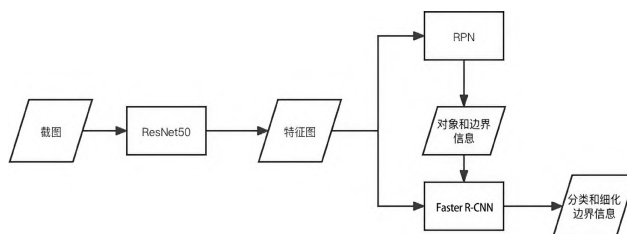


图3 Faster R-CNN 模型

在Faster R-CNN中,区域建议网络(Region Proposal Network,RPN)用于预测输入屏幕截图上的一组边界框。RPN通过表示特征图来估计包含的对象(如logo)的存在概率和形状。RPN模型如图4所示,在每个滑动窗口位置,RPN同时预测多个区域提案,每个位置最多可预测 $k$ 个提案。特征图上的每个点都作为一个锚点,RPN会生成多种尺度和宽高比的锚点框,这些锚点框的坐标基于原始图像坐标。锚点框被输入到分类层(cls层)和回归层(reg层)两个网络层,cls层输出的概率得分为 $2k$ ,评估每个提议作为对象或非对象的概率,即判断特征图中的锚点框是否属于 $k$ 类对象中的一种。reg层输出4个位置坐标,共有 $4k$ 个输出,用于编码 $k$ 个锚点框的坐标<sup>[21]</sup>。传统的检测方法在生成检测框时耗时较长,如R-CNN中的选择性搜索(Selective Search,SS)。相比之下,Faster R-CNN直接利用RPN生成检测框,显著提升了生成速度,是其一大优势。

本文采用Faster R-CNN模型来预测对象类别,并细化每个对象的形状和大小,输入为特征图和边界框。总体而言,给定一张屏幕截图,Faster R-CNN模型能够报告一系列候选标志,每个标志附有一个置信度得分。

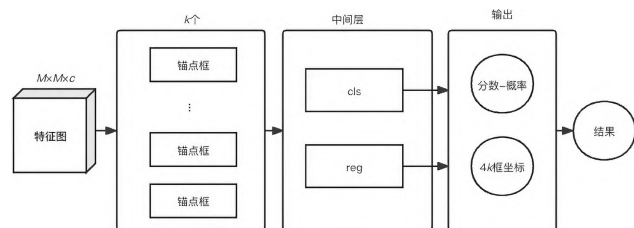


图4 RPN模型

根据置信度得分对这些标志进行排名，排名最高的标志被视为潜在的身份标识。

### 2.2.2 图像识别任务

目标品牌列表中包括中国银行、茅台等多个品牌。针对每个品牌，本文维护一个包含合法网站和已确认为钓鱼网站的logo图像列表，合法网站的logo一般通过在HTML中提取或进行图片搜索获得。为了提高准确率，本文将品牌合法网站的截图通过目标切割任务处理后被认定为logo的矩形区域作为合法logo。这一做法旨在防止通过目标切割得到的截图中的矩形区域被错误识别为logo，或者与正规logo相混淆，从而导致误判。钓鱼网站通常批量制作，一旦发现新的钓鱼网站，就将其识别为logo的矩形区域，若其与logo列表中的其他图像相似度较低，就将该图像加入logo列表，以此提高识别效率。

给定一个logo，如果其与目标品牌列表中的logo相似度超过预设阈值，则将相应品牌标记为潜在的网络钓鱼目标。本文使用基于CNN的EfficientNetV2模型进行图像分类，将图像分为受保护的品牌和无关图像。对于被分类为受保护品牌的图像，若具有模仿特征，则被视为可疑的钓鱼网站。在图像处理阶段，对训练集的图像进行处理，由于训练集图像为logo图，无需进行旋转处理，仅进行大小调整和归一化处理，以加快训练速度。

EfficientNetV2模型自2021年发布以来，在参数效率和训练速度方面表现均优于现有技术<sup>[22]</sup>。由于受保护的列表经常变动，如果添加一个新品牌就要求重新训练整个网络。因此，相比于VggNet、GoogLeNet等常用的图像分类模型，EfficientNetV2<sup>[23]</sup>的快速训练优

势显得尤为重要。虽然通过训练过程中逐渐增加图像大小可以进一步加速训练，但这可能导致准确率下降。为了补偿在准确率方面的损失，EfficientNetV2提出一种改进的渐进学习方法，通过自适应地调整图像大小和正则化（如数据增强）来优化学习过程。通过这种渐进学习方法，EfficientNetV2的性能显著优于其他模型，有助于提高本文系统在识别钓鱼网站方面的准确率。

### 2.3 钓鱼网站评价

钓鱼网站视觉上与合法品牌网站相似，其目的是窃取用户提交的敏感信息。因此，识别钓鱼网站需要提取两类关键特征：网页相似度相关特征和窃取功能特征。本文通过图像切割任务和图像识别任务输出的网页，满足了网页相似度相关的特征需求<sup>[24]</sup>。然而，对于窃取功能特征的判断尚未涉及，这可能导致如下情况。一个合法网页，若在介绍或引用保护品牌列表中的品牌（如中国银行）时使用了该品牌的logo，且该网页不在白名单中，该网页的截图也可能被错误地认定为钓鱼网站。

为解决上述问题，本文对图像识别任务输出的图片进行进一步评估，以判断其是否为钓鱼网站。特征空间下的评价特征如表1所示，针对窃取功能特征和信用度特征进行判断。窃取功能特征指与窃取敏感信息相关的特征，如账号和密码，根据HTML中是否包含<form>、<password>、<submission>等标签进行判断。因此，在评估网站时，本文将获取网站的文档对象模型（Document Object Model, DOM）树，检查是否包含这些窃取功能特征，并根据这些特征的标签使用频率以及获取信息的相关性进行加权评估，以综合判断网页是否存在窃取功能特征。

表1 特征空间下的评价特征

| 特征空间   | 特征                                            |
|--------|-----------------------------------------------|
| 窃取功能特征 | <form> 标签<br><password> 标签<br><submission> 标签 |
| 信用度特征  | Page Rank<br>Alexa 排名<br>域名存在时间               |

信用度特征是指判断网站可受信任程度的特征,对于钓鱼网站经常使用子品牌和第三方品牌元素的情况,本文发现现有的目标检测和图像识别方法尚不能准确地将此类网站识别为良性网站。因此,在评估一个网站的信誉度时,需要考虑诸如域名存活时间、Page Rank 值、Alexa 排名等因素。对于各个评价特征,判定为钓鱼网站的评价分数如公式(1)所示。

$$F(x, y) = g(x) \times h(y) \quad (1)$$

其中,  $g(x)$  为窃取功能特征评价分数,  $h(y)$  为信用度特征。窃取功能特征评价分数  $g(x)$  的计算方法如公式(2)所示。

$$g(x) = \alpha(x_1) + \alpha(x_2) + \alpha(x_3) \quad (2)$$

其中,  $\alpha$  值的计算方法如公式(3)所示。

$$\alpha(x) = \begin{cases} 0, & \text{网页不存在 } m \text{ 标签} \\ 1, & \text{网页存在 } m \text{ 标签} \end{cases} \quad (3)$$

信用特征评价分数  $h(y)$  的计算方法如公式(4)所示。

$$h(y) = \beta(y_1) \times \gamma(y_2) \times \delta(y_3) \quad (4)$$

其中,  $\beta(y)$ 、 $\gamma(y)$  和  $\delta(y)$  是信用度特征判别函数,当  $y$  超过阈值时记为 1,不超过时记为 0。即判定为钓鱼网站需要同时满足具有窃取功能、网站信用度低以及域名存在时间短等要求,以降低误报率。

### 3 实验设计与结果分析

#### 3.1 模型训练及数据来源

在 Faster R-CNN 模型方面,本文基于 Detectron2 框架训练 Faster R-CNN 模型。Detectron2 是一个适用于对象检测、分割以及其他视觉识别任务的平台,其提供了丰富的基线结果和预训练模型。Faster R-CNN 主要用于目标切割和从截图中提取 logo,由于本文不需要使用保护品牌列表中的网站截图进行训练,因此可以从含有 logo 的大量网站中获得数据来源。本文使用 Phishpedia<sup>[25]</sup> 提供的、经过 Logo-2K+ 数据集训练的 Faster R-CNN 模型<sup>[26]</sup>,并以本文的数据进行进一步的训练和测试。

对于 EfficientNetV2 模型,本文通过 PyTorch 框架进行训练。在数据来源方面,本文从某单位获取了部分

仿冒保护品牌的钓鱼网站数据。在确认网站为钓鱼网站后,本文保存了其 URL、HTML 的截图,以防止网站失效,并使用这些截图对 Faster R-CNN 模型再次进行训练和测试。为了训练品牌分类,本文将钓鱼网站按仿冒品牌进行手动分类,通过训练后的 Faster R-CNN 模型对官方网站及其子网站进行图像切割,获取相应的 logo。经扩线后,满足要求的 118 个品牌的钓鱼网站数量为 2548 个,经过图形切割后获取的 logo 数量为 3312 个,保护品牌网站截取的 logo 数量为 367 个。其他良性网站经过过去重后再随机选取 1000 个网站,得到 logo 数量为 1356 个。在目标切割时,有时会选取多个区域为 logo 区,因此 logo 数量大于网页数量。由于各个品牌的钓鱼网站数量不同,在训练分类模型时,为防止数据集过少导致的训练结果不佳,本文选择其中数量最多的 8 个保护品牌网站进行分类训练。

对于数据集,钓鱼目标切割任务的数据集包括 2548 个钓鱼网站截图和 1000 个良性网站截图,将这些截图按照 9:1 的比例随机分配到训练集和验证集中。对于图像识别任务,本文进行了两类实验,一类是使用网站截图(包括钓鱼网站和品牌网站截图共 2216 个,以及非品牌网站截图 2000 个),另一类是使用切割后的 logo 图像。钓鱼网站和品牌网站的切割后 logo 总数为 3521 个,按照 9:1 的比例分配到训练集和验证集中,而非保护品牌的 logo 中选取 900 个作为训练集,100 个作为验证集。

#### 3.2 指标的定义

在性能评估方面,本文采用的指标包括召回率、假阳率、精确率和  $F1$  值。召回率表示正确预测的网络钓鱼数量与所有网络钓鱼数量的比例,反映了所有钓鱼数据中被正确预测为阳性样本的概率。假阳率表示错误预测为网络钓鱼的网站数量与所有合法网站数量的比例,显示了所有合法网站中被错误预测为阳性样本的概率。精确率指正确预测的网络钓鱼数量与预测的网络钓鱼总数的比例,显示了被预测为阳性样本中实际上是阳性的比例。 $F1$  值是精度和召回率的加权调



和平均值,其同时考虑准确率和召回率,在类别不均衡的情况下比单一的准确率更有参考价值。另外,评估指标还包括真阳率、真阴率和假阴率。真阳率是对阳性样本进行正确预测的比例;真阴率是对阴性样本进行正确预测的比例;假阴率是对阴性样本预测错误的比例。

### 3.3 实验环境

本文的实验环境设置如表2所示。

表 2 实验环境

| 实验环境   | 具体信息                                           |
|--------|------------------------------------------------|
| 操作系统   | Ubuntu                                         |
| CPU    | Intel(R) Xeon(R) Platinum 8369B CPU @ 2.90 GHz |
| GPU    | A100(80 GB)                                    |
| 内存     | 128 GB                                         |
| 开发语言   | Python 3.9                                     |
| 深度学习框架 | PyTorch 2.0.1                                  |

### 3.4 钓鱼扩线效果

本文从某机构获取部分钓鱼网站数据,并对这些网站进行扩线,经过人工审核后得到新的钓鱼网站,表3展示了2023年1~5月的钓鱼扩线情况。扩线方法包括寻找与原网站相关的IP或域名相关的新网页。实验结果表明,表3中新钓鱼网站在扩线后数量占比小,扩线后网站中包含大量非目标网站。本文的扩线方法主要分析了钓鱼网站的各类特征,寻找具有相同特征的其他网站。上述情况表明,虽然本文的扩线方法可以发现具有相同特征的网站,但只有一部分是钓鱼网站。因此,扩线的有效性受到限制。

表 3 钓鱼扩线

| 日期     | 供扩线数量 / 个 | 扩线后数量 / 个 | 新钓鱼网站数量 / 个 | 后续 fofa 发现新网站数量 / 个 |
|--------|-----------|-----------|-------------|---------------------|
| 202301 | 1196      | 21848     | 105         | 331                 |
| 202302 | 4445      | 61113     | 367         | 113                 |
| 202303 | 1188      | 19708     | 178         | 147                 |
| 202304 | 2013      | 36514     | 187         | 104                 |
| 202305 | 1387      | 21934     | 96          | 65                  |

本文基于网站标题(title)的扩线方法,使用fofa引擎根据特定的标题进行搜索。例如,以“浦发银行在线客服”为标题,在fofa上进行搜索,得到64条匹配结果(涵盖4个独立IP),其中大多数为钓鱼网站。这种

基于标题的扩线方法的特点是搜索结果中网页内容高度同质化,截图基本相似甚至一致,钓鱼网站占比较大。然而,这种方法的局限性在于依赖钓鱼网站的标题信息,而拥有标题信息的钓鱼网站数量有限,且存在大量重复标题。目前,钓鱼网站扩线工作尚处于初级阶段,尚未发现其他具有类似扩线成果的研究,因此缺乏对比数据。下一步,本文计划使用fofa的功能来扩展包含特定标题的钓鱼网站,进一步提高扩线的有效性和准确率。

### 3.5 钓鱼目标切割任务效果

本文利用预训练的Faster R-CNN模型进行实验,尽管只使用了少量样本进行训练,但效果显著。图5展示了通过扩线方法找到的伪装成中国银行的钓鱼网站,经过模型切割后,得到3个被判定为logo区域的矩形区间,并将其裁剪出来,结果如图6所示。图6中的3张图片均具有明显的仿冒特征,用于误导用户相信该网站是真实的中国银行网站。



图 5 钓鱼网站目标切割 - 中国银行

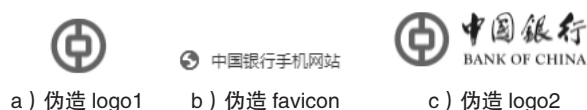


图 6 钓鱼网站切割 logo- 中国银行

### 3.6 图像识别任务效果

本文对8个保护品牌的钓鱼网站截图进行目标切割,将所得到的截图和logo按照每个品牌9:1的比例随机分为训练集和验证集。为了验证图像识别任务的效果,本文进行了对比实验,分别使用EMD、Alexnet和EfficientNetV2模型进行测试,分类对比实验结果如表4所示。实验结果表明,使用logo作为分类目标时,召回率和精确率都有显著提升,其中EfficientNetV2在召回率和准确率上相比于Alexnet表现更佳。

对实验结果进行分析,在针对截图进行的对比实验中,未被识别出的钓鱼网站具有以下特点:网页内

表 4 分类对比实验结果

| 模型             | 分类目标 | 召回率  | 精确率  | F1   |
|----------------|------|------|------|------|
| EMD            | 截图   | 62 % | 76 % | 0.68 |
| Alexnet        | 截图   | 82 % | 89 % | 0.85 |
| Alexnet        | logo | 95 % | 97 % | 0.95 |
| EfficientNetV2 | 截图   | 85 % | 91 % | 0.87 |
| EfficientNetV2 | logo | 97 % | 98 % | 0.97 |

容是动态变化的,导致对同一网页的多次截图不一致;钓鱼网站模仿的是品牌网站更新前的网页;钓鱼网站只在部分内容(如logo、favicon、少量文字和样式)上模仿了品牌网站,导致分类精度较低。相比之下,在使用logo进行分类时,对于页面相似的网站,其分类效果与使用截图相近甚至更好,而在面对差异较大的钓鱼网站时,由于只需要针对网页的logo部分,因此不会受到页面中其他内容的干扰。因此,针对logo进行分类实验的效果更佳。

### 3.7 评价模块效果

在图像识别模块中,为了实现更高的召回率和精确率,本文放弃使用截图对比,而是使用logo作为分类依据。这导致一些合法网站由于引用了其他保护品牌的logo,被错误地判定为钓鱼网站。目前,本文发现了48个此类假阳性报告,随机选取48个钓鱼网站和48个假阳性网站进行评价,结果如表5所示。

表 5 评价模块实验结果

| 评价依据     | 信用度特征       |           |         | 窃取功能特征     |                |                  | 评价模块判断为钓鱼网站 |
|----------|-------------|-----------|---------|------------|----------------|------------------|-------------|
|          | PageRank 值低 | Alexa 排名低 | 域名存在时间短 | 含有<form>标签 | 含有<password>标签 | 含有<submission>标签 |             |
| 假阳性样本    | 0           | 0         | 0       | 11         | 6              | 6                | 0           |
| 钓鱼网站数量/个 | 47          | 47        | 47      | 38         | 33             | 33               | 48          |

针对上述实验结果进行分析,在对比评价过程中发现,合法网站的域名存在时间通常远长于钓鱼网站的平均存活时间。此外,部分合法网站的HTML中不含有<form>、<password>、<submission>等标签,因此可以判定为不存在偷窃特性。对于网站的Page Rank值、Alexa排名等信用度评价,假阳性网站的信用度也远高于钓鱼网站。

## 4 结束语

本文提出一种通过已有的钓鱼网站信息寻找更多可疑网站的扩线方法,并通过目标切割模块,对扩线后的网页截图进行切割,以获取logo,再利用这些logo进行图像分类,以此实现网络钓鱼检测。通过真实场景和数据的实验验证了此方法的有效性,并通过增加评价模块降低了假阳率。但本文方法也存在一些不足,如每天扩线后的网页数量较多,需要花费大量时间进行截图工作。接下来,计划改进评价体系,并将其置于截图工作之前,增加预筛选模块,以减少所需截图网页的数量,节省时间。此外,对于通过二维码或二次跳转到新页面的钓鱼网站,本文所提方法无法进行识别和检测。在未来的工作中,将针对这些不足进行改进,并努力解决二次跳转问题。

### 参考文献:

- [1] ITU. Measuring Digital Development: Facts and Figures 2022[EB/OL]. (2023-08-01)[2023-08-14]. <https://www.itu.int/en/ITU-D/Statistics/Pages/facts/default.aspx>.
- [2] HULME G V, GOODCHILD J. Social Engineering: The Basics[EB/OL]. (2010-01-14)[2023-08-14]. <https://www2.cso.com.au/article/print/625643/social-engineering-basics/>.
- [3] CATAL C, GIRAY G, TEKINERDOGAN B, et al. Applications of Deep Learning for Phishing Detection: A Systematic Literature Review[J]. Knowledge and Information Systems, 2022, 64(6): 1457-1500.
- [4] Google. Google Safe Browsing[EB/OL]. (2021-11-14)[2023-08-14]. <https://safebrowsing.google.com>.
- [5] VERMA R, DYER K. On the Character of Phishing URLs: Accurate and Robust Statistical Learning Classifiers[C]//ACM. Proceedings of the 5th ACM Conference on Data and Application Security and Privacy(CODASPY'15). New York: ACM, 2015: 111-122.
- [6] KARIM A, SHAHROZ M, MUSTOFA K, et al. Phishing Detection System through Hybrid Machine Learning Based on URL[J]. IEEE Access, 2023(11): 36805-36822.
- [7] ALJOFEY A, JIANG Qingshan, QU Qiang, et al. An Effective Phishing Detection Model Based on Character Level Convolutional Neural Network from URL[EB/OL]. (2020-09-15)[2023-08-14]. <https://www.mdpi.com/2079-9292/9/9/1514>.
- [8] SAFI A, SINGH S. A Systematic Literature Review on Phishing Website Detection Techniques[J]. Journal of King Saud University-Computer and Information Sciences, 2023, 35(2): 590-611.
- [9] FU A Y, LIU Wenyin, DENG Xiaotie. Detecting Phishing Web Pages with Visual Similarity Assessment Based on Earth Mover's Distance (EMD)[J]. IEEE Transactions on Dependable and Secure Computing, 2006, 3(4): 301-311.
- [10] ABDELNABI S, KROMBHOLZ K, FRITZ M. VisualPhishNet: Zero-Day Phishing Website Detection by Visual Similarity[C]//ACM.



Proceedings of the 2020 ACM SIGSAC Conference on Computer and Communications Security (CCS'20). New York: ACM, 2020: 1681–1698.

[11] AFROZ S, GREENSTADT R. Phishzoo: Detecting Phishing Websites by Looking at Them[C]//IEEE. 2011 IEEE Fifth International Conference on Semantic Computing. New York: IEEE, 2011: 368–375.

[12] BOZKIR A S, AYDOS M. LogoSENSE: A Companion HOG Based Logo Detection Scheme for Phishing Web Page and E-Mail Brand Recognition[EB/OL]. (2020-08-01)[2023-08-14]. <https://www.sciencedirect.com/science/article/abs/pii/S0167404820301279>.

[13] CHANG E H, CHIEW K L, TIONG W K, et al. Phishing Detection via Identification of Website Identity[C]//IEEE. 2013 International Conference on IT Convergence and Security (ICITCS). New York: IEEE, 2013: 1–4.

[14] WANG Ge, LIU He, BECERRA S, et al. Verilogo: Proactive Phishing Detection via Logo Recognition[EB/OL]. (2011-01-01)[2023-08-14]. <https://hovav.net/ucsd/dist/verilogo.pdf>.

[15] YUN Lei, LI Dan, WANG Huanhuan. Summary of Research on Phishing Website Detection Technology[J]. Reliability and Environmental Testing of Electronic Products, 2021, 39(5): 114–119.

云雷, 李丹, 王欢欢. 钓鱼网站检测技术研究综述[J]. 电子产品可靠性与环境试验, 2021, 39(5): 114–119.

[16] ASIRI S, XIAO Yang, ALZHRANI S, et al. A Survey of Intelligent Detection Designs of HTML URL Phishing Attacks[J]. IEEE Access, 2023(11): 6421–6443.

[17] BU S J, CHO S B. Deep Character-Level Anomaly Detection Based on a Convolutional Autoencoder for Zero-Day Phishing URL Detection[EB/OL]. (2021-06-17)[2023-08-14]. <https://doi.org/10.3390/electronics10121492>.

[18] SOMESHA M, PAIS A R, RAO R S, et al. Efficient Deep Learning

Techniques for the Detection of Phishing Websites[EB/OL]. (2020-06-17)[2023-08-14]. <https://link.springer.com/article/10.1007/s12046-020-01392-4>.

[19] ADEBOWALE M A, LWIN K T, HOSSAIN M A. Intelligent Phishing Detection Scheme Using Deep Learning Algorithms[EB/OL]. (2020-06-04)[2023-08-14]. <https://www.emerald.com/insight/content/doi/10.1108/JEIM-01-2020-0036/full/html>.

[20] HE Kaiming, ZHANG Xiangyu, REN Shaoqing, et al. Deep Residual Learning for Image Recognition[C]//IEEE. Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. New York: IEEE, 2016: 770–778.

[21] REN Shaoqing, HE Kaiming, GIRSHICK R, et al. Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2015, 39(6): 1137–1149.

[22] ABDILLAH R, SHUKUR Z, MOHD M, et al. Performance Evaluation of Phishing Classification Techniques on Various Data Sources and Schemes[J]. IEEE Access, 2023(11): 38721–38738.

[23] TAN Mingxing, LE Q V. Efficientnetv2: Smaller Models and Faster Training[C]//PMLR. International Conference on Machine Learning. New York: PMLR, 2021: 10096–10106.

[24] DAVOUDI M R, YARI A R. Improving the Feature Section Method Based on Genetic Algorithm to Increase the Efficiency of Detecting Phishing Websites[J]. Automatic Control and Computer Sciences, 2023, 57(3): 213–221.

[25] LIN Yun, LIU Ruofan, DIVAKARAN D M, et al. Phishpedia: A Hybrid Deep Learning Based Approach to Visually Identify Phishing Webpages[C]//USENIX. In 30th USENIX Security Symposium. Berlin: USENIX, 2021: 3793–3810.

[26] WANG Jing, MIN Weiqing, HOU Sujuan, et al. Logo-2K+: A Large-Scale Logo Dataset for Scalable Logo Classification[C]//IEEE. Proceedings of the AAAI Conference on Artificial Intelligence. New York: IEEE, 2020, 34(4): 6194–6201.

## 电子科技大学许春香教授团队在 IEEE TIFS 发表研究成果

近日, 电子科技大学计算机科学与工程学院(网络空间安全学院)许春香教授团队的硕士研究生宋雅晴以第一作者在网络与信息安全领域期刊IEEE Trans. on Information Forensics and Security (CCF A类)上发表题为《Hardening Password-Based Credential Databases》的论文, 许春香教授、张源研究员为论文通讯作者, 电子科技大学为该论文唯一作者单位。

口令认证是目前较便捷高效的认证方法之一, 被广泛应用于诸多场景中。然而, 口令认证凭证数据库泄露导致的用户口令泄露是企业数据库最常发生的安全事件。例如, 2012年, LinkedIn遭到黑客攻击, 导致全球近650万用户的口令被泄露; 2020年, 有超过50万个Zoom用户的口令被黑客窃取并出售。该论文调研了现有3种口令凭证数据库保护机制。由于口令本身固有的低熵特性, 基于哈希函数的口令凭证保护机制无法抵御针对口令的离线字典猜测攻击。简单的“加盐”保护机制一定程度上可以增强对口令凭证的保护, 然而一旦“盐值”泄露, 敌手依然可以通过字典猜测攻击或者彩虹表攻击获得用户口令。

针对上述问题, 该论文提出了一个针对口令凭证数据库的保护机制, 该机制仅部署在服务器端, 由一组密钥服务器以门限方式分布式存储盐值, 以增强对口令认证凭证的保护, 使得敌手无法通过窃取口令认证凭证数据库, 进而恢复用户口令。该机制对用户完全透明且不改变用户端的交互模式, 能够兼容现有口令认证系统。同时, 该论文采取了两种前设防御机制, 使得服务器在保持用户透明性的同时更新盐值或更换密钥服务器, 进一步增强了系统安全保障。

许春香, 电子科技大学教授, 博士生导师, 中国密码学会理事, 中国保密协会隐私保护专委会常务委员, 电子科技大学第八届、第九届学术委员会委员。其研究方向包括应用密码学、隐私保护、云计算安全、物联网安全、区块链技术和人工智能安全等。近10年来, 主持国家自然科学基金面上项目3项, 国家重点研发计划子课题2项; 在IEEE TIFS、IEEE TDSC、IEEE TMC、IEEE TCC和IEEE TSC等重要国际学术期刊和会议上发表论文100多篇; 荣获IEEE Trans. Cloud Computing期刊年度最佳论文奖。(来源: 电子科技大学, <https://news.uestc.edu.cn/?n=UestcNews.Front.DocumentV2.ArticlePage&Id=90829>)