

Fishing for Smishing: Understanding SMS Phishing Infrastructure and Strategies by Mining Public User Reports

Sharad Agarwal
University College London
London, United Kingdom
sharad.agarwal@ucl.ac.uk

Guillermo Suarez-Tangil
IMDEA Networks
Madrid, Spain
guillermo.suarez-tangil@imdea.org

Antonis Papasavva
University College London
London, United Kingdom
antonis.papasavva@ucl.ac.uk

Marie Vasek
University College London
London, United Kingdom
m.vasek@ucl.ac.uk

ABSTRACT

Recently, there has been a worldwide surge in SMS phishing, aka smishing. However, the lack of open-access updated datasets makes it challenging for researchers to study this global issue. Mobile network operators and government agencies provide users special SMS spam reporting services. Though, these services are regional and users are largely unaware. So, users often turn to public forums such as Twitter or Reddit to report and discuss smishing. This paper presents a novel methodological approach to collect an updated smishing dataset and measure the infrastructure, targets, and strategies employed by attackers to lure victims. We programmatically collect users' smishing reports from five public forums, collating over 64.5k smishing image attachments and reports, which include 28.6k sender IDs and 25.9k URLs criminals abuse to conduct smishing campaigns across 66 languages. We unveil the exploited infrastructure ranging from mobile network operators to domains. We categorize smishing texts into seven scam types and explain lures criminals use to deceive victims into providing sensitive/financial information. Through a case study using real time measurements on a random sample of Twitter posts, we showcase how to uncover Android malware spread via smishing. We suggest effective mitigation approaches to curb this widespread cybercrime.

CCS CONCEPTS

• **Security and privacy** → **Usability in security and privacy**; **Phishing**.

KEYWORDS

smishing; cybercrime; sms scam; online financial fraud

ACM Reference Format:

Sharad Agarwal, Antonis Papasavva, Guillermo Suarez-Tangil, and Marie Vasek. 2025. Fishing for Smishing: Understanding SMS Phishing Infrastructure and Strategies by Mining Public User Reports. In *Proceedings of the 2025 ACM Internet Measurement Conference (IMC '25)*, October 28–31, 2025, Madison, WI, USA. ACM, New York, NY, USA, 19 pages. <https://doi.org/10.1145/3730567.3764431>

1 INTRODUCTION

Smishing or SMS phishing has been on the rise for the last few years with 300k to 400k smishing texts being sent daily [99]. While smishing is similar to email-based phishing, it is more challenging to detect and differentiate a smish from a benign text. Unlike phishing, it is difficult for mobile network operators to block smishing as the only metadata in an SMS is the sender ID and timestamp. Furthermore, the small URL bar of mobile browsers covers the complete URL, making it complex for a user to differentiate a phishing website from a legitimate one [74].

Adversaries continue to update their attack vectors and adapt to users' trends. As users prefer mobile phones over computers due to their convenience, accessibility, and functionalities, threat actors also shift to a new communication medium to target users — SMS/RCS/iMessage. As of January 2025, more than 90% of the world population has smartphones [34]. With the significant increase in smartphone usage and entities like financial institutions providing app-only solutions, spreading mobile malware via SMS has become the latest attack vector in banking defrauding customers [14]. With the high open and response rates of SMS, in contrast to emails [90], fraudsters leverage this to target users via smishing. Consumers in the US lost \$330 million in 2022 to text scams, more than double the losses reported in the previous year and nearly five times the amount lost in 2019 [43].

Smishing was forecasted as one of the top eight cybersecurity threats to organizations globally [93]. Criminals exploit every available opportunity to target users through smishing campaigns, including tax season and holiday periods. Reportedly, criminals even abuse global crises such as international conflicts, by impersonating charities raising funds [115]. Seven in ten people (71%) in the UK reported receiving a suspicious text, and almost a million people in just three months reported following scammers' instructions in text or call [82]. Phishing/smishing attacks are the most reported cybercrime in the US, with victims losing over \$129k from just toll-themed smishing texts [39].

Mobile network operators in certain countries, like the US, UK, and Australia, run a special reporting service to collect spam and scam texts from individual users [3, 42, 84]. Alas, this data is unavailable for research due to privacy concerns and the challenges of accurately distinguishing between spam, scams, and benign messages. Users may report a benign message as spam, which would reveal private information in the SMS message.

Research Gap. Despite smishing being a prominent global issue, the academic community lacks insights into this widespread threat due to the rapidly evolving nature of these scams and the unavailability of an updated data feed. To evade detection from mobile network operators’s filters, fraudsters have started moving from traditional SMS to encrypted online communication channels — RCS and iMessage [38]. This shift makes it difficult to access smishing texts sent over encrypted channels.

While some researchers published spam datasets more than a decade ago [21, 111], these contain limited smishing texts. URLs used in smishing tend to have a short lifespan, ranging from a few minutes to a maximum of a few days [66], and researchers devising smishing detection have been using over 10-year-old spam/ham datasets [87] where the information about the URLs is inaccessible. As smishing attack trends continuously change, research conducted on outdated data fails to produce relevant or actionable findings, e.g. by failing to note the rise of conversational smishing scams like ‘Hi mum and dad’ scams [5] or ‘Wrong number’ scams [25]. The latest research in the area has either published very small regional datasets [113], considered mixed spam, OTP, and scam texts [81], or used proprietary regional data that is unavailable to other researchers [6]. Smishing has not been researched enough to suggest/develop effective mitigation techniques preventing users from falling prey.

Contribution. Central to the lack of recent data, we see a methodological data collection gap that we address in this paper as the basis of the largest and most comprehensive measurement of smishing to date. A large majority (76%) of mobile users in the UK have never heard of the suspicious SMS reporting service – 7726 [83]. Other European and Asian countries do not have any official service to report suspicious messages.¹ As a result, users turn to online forums and post smishing messages to report them to authorities/brands or spread awareness to other online users. To this end, we leverage five online forums to collect reported smishing texts. We build an updated novel dataset and perform measurements (e.g., HLR lookup) that provide a holistic view of smishing and help suggest countermeasures.

Based on our data, we measure key aspects of smishing to better understand this networked phenomenon, including the underlying infrastructure, targeted victims, and strategies employed by active campaigns. Our measurements are designed to address the following research questions:

RQ1 What infrastructure do cybercriminals exploit to conduct smishing?

RQ2 How do cybercriminals lure victims into smishing?

This paper provides the following contributions:

- We provide a novel methodology to collect an updated SMS phishing dataset. By leveraging smishing reports collected from a distributed network of contributors in social media, we address a critical gap in data collection for understanding real-world SMS phishing activity. The pseudo-anonymized dataset is available at <https://github.com/reportsmishing/Smish>

ing-Dataset-IMC25 (described in Appendix C) enabling future work.

- Our measurement approach leverages both passive and active measurements. Our passive measurement presents insights into the infrastructure that criminals abuse to conduct smishing using mobile networks and trend analysis (§4). Our active measurement identifies, through a case study, malware that scammers spread through smishing texts, attempting to perform drive-by download attacks on Android (§6).
- To offer situational awareness, we systematically measure lure principles and characterize scams. We discover that the lures scammers use for conversation scams differ from those in other smishing texts (§5.5). We distribute the collected smishing texts into known SMS scam categories [6], with banking being the most popular one (§5.2).

2 MOTIVATION AND RELATED WORK

Smishing has become cybercriminals’ preferred medium as users trust mobile communications and have significantly higher URL click rates in mobile messaging compared to emails [99]. However, unlike phishing, an updated, comprehensive smishing dataset is unavailable, hindering researchers from studying this problem.

SMS Spam. Past researchers have published SMS spam and ham datasets consisting of limited scam texts, most of which were collected more than a decade ago [21, 30, 109, 111]. As adversaries keep changing their tactics, updated data is required to successfully detect, understand, and mitigate scams. Srinivasan et al. study SMS abuse campaigns without differentiating spam vs scam texts [106]. A recent study also curated a new spam dataset using Twitter posts [110]. Unfortunately, it lacks clear labeling between generic spam (unsolicited marketing, forgotten newsletter subscriptions, etc.) and smishing (scam texts, impersonation texts with malicious URLs). The inherent imbalance in such datasets, with smishing messages significantly underrepresented compared to generic spam, complicates analysis and limits the strength of any conclusions drawn. The patterns and trends derived from such data would disproportionately reflect the dominant category — spam, making it difficult to identify features unique to smishing.

Phishing aggregators like the AWPGE eCrime Exchange [50], OpenPhish [86], and Phishtank [22] provide lists of malicious domains. Still, these are primarily collated from emails, and they otherwise do not differentiate between collection mechanisms, be it SMS, email, or otherwise. While spam/ham datasets are available, there is an urgent need to study smishing. Our work addresses this gap by providing an updated labeled smishing dataset.

Smishing Data Collection Methods. While there has been an uptick in smishing, limited prior work has focused on collecting smishing data. A small historical collection of smishing screenshots has been curated and made publicly available on Pinterest [33]. Recently, two studies have collected a smishing dataset comprising 1,090 and 518 texts, collected via crowdsourcing through an online website [113] and a mobile application [92], respectively. While this method enables collecting data in real time, both datasets are small and geographically limited, with users reporting only from the US and Pakistan, which is insufficient to understand the global nature of smishing. Crowdsourcing scam messages is complex,

¹Recently, countries like India and Singapore have started to collect suspicious communication reports from users via their platforms. Specific mobile network operators in Germany have integrated Apple’s one-click reporting.

as the public lacks the trust to report these messages to a non-governmental, unofficial platform.

Others have collaborated with mobile network operators or security vendors to access blocked smishing texts [5, 6, 66]. Even though these datasets are large, the data here also remains regional and the data is publicly unavailable. It is possible to directly collect data here via honeypots [15]. While this approach is innovative and does not depend on third parties, it is challenging to effectively seed the mobile numbers while allowing for low cost scaling.

Public Online SMS Gateways. Public online SMS gateways (PSGs) provide disposable virtual numbers to receive text messages, where users do not want to provide their mobile numbers. Moreno et al. [78] collected 70m SMS messages and found only 41 URLs (125 SMS texts) considered harmful by Google Safe Browsing. Others also found limited malicious URLs in the text messages collected from PSGs [95, 96], i.e., plausible smishing messages. On the contrary, using 4 additional PSGs, Nahapetyan et al. found over 67k smishing messages [81], significantly more than prior studies. One of the reasons is the identifiers they consider for a message to be a smish – one-time passcodes (OTPs). Users primarily utilize PSGs to receive OTPs to install applications or access services. Differentiating smishing texts with OTPs from benign ones is challenging.

While researchers have studied PSGs to identify malicious URLs or smishing, they do not represent the various smishing texts users usually receive. This is because the mobile numbers are (1) publicly available, (2) not directly associated with users, and (3) get recycled frequently. Adversaries avoid sending malicious URLs to such numbers where their tactics would get exposed before reaching potential victims.

Smishing Detection. It is essential to efficiently detect and block smishing texts. Prior work has proposed broad rule-based approaches to filter smishing SMS [58, 59, 123], utilizing limited smishing texts from an old spam/ham dataset [111] or PhishTank [22]. While these initial efforts provide the essential groundwork, the derived rules were based on relatively small, dated samples, limiting their generalizability. Rule-based systems become ineffective in real-world scenarios. The current defense mechanism is to evolve detection mechanisms while threat actors evolve their methods in a continuous game of whack-a-mole.

Some researchers suggest using a Naive Bayesian classifier [61, 76] with additional checks like the presence of a URL providing an APK file [61, 75] or checking URL and phone number in blocklists [75] to detect smishing. Amrutkar et al. use web pages' static features to distinguish between benign and malicious mobile websites [12]. The underlying problem of data unavailability remains [87], preventing the training and evaluation of any proposed machine-learning models. Our work provides an updated novel smishing dataset collected through various online forums. The insights from our paper could help improve these detection methods.

3 METHODOLOGY

To address the problem of data availability, we collect smishing texts that users and analysts post online. This section explains the novel data curation and analysis techniques we conduct to understand this growing threat.

3.1 Data Collection

We identify five online forums where users voluntarily post smishing screenshots or report smishing texts that they receive, along with the sender IDs. We provide an overview of the collected data in Table 1.

3.1.1 Twitter. Security-conscious users on Twitter, now known as X, post screenshots of smishing texts to report them to companies or spread awareness to other users on the platform. To this end, we manually search multiple keywords and find that 'smishing,' 'phishing sms,' 'sms scam,' and 'sms fraud' return the best results for users' tweets reporting smishing. Fig. 4 in Appendix E shows users from various countries reporting smishing texts on Twitter.

We use the Twitter Academic API² to collect tweets and their image attachments in real-time from November 30, 2022, until the Twitter academic API shutdown on June 23, 2023. Where available, we also collect the original tweet if the keyword was in the reply to a tweet and its image attachment. Additionally, we query our keywords on Twitter to collect past tweets and their image attachments between January 1, 2017, and November 30, 2022. Table 15 in the Appendix shows the yearly distributions of tweets and image attachments.

3.1.2 Reddit. We find that users on Reddit also post about smishing text messages and discuss smishing. Even though specific subreddits such as *r/Scams* exist, we discover that users post about smishing in multiple subreddits. To this end, we use the Reddit API between January 1, 2017, and September 30, 2023, to search for the exact four identified keywords we use with Twitter. In total, we collect 1,707 image attachments from 1,771 unique submissions that users post on Reddit over 911 subreddits. While the majority (121) of submissions were posted on *r/Scams*, followed by 48 posts on *r/cybersecurity* and 42 on *r/ledgerwallet*, they are broadly distributed – 582 subreddits containing one post each.

3.1.3 Smishing.eu. Smishing.eu was a website where users from any country were able to report smishing texts that they received, focused towards European users. This online platform allowed users to fill in a form to submit a screenshot of the smishing text or the user's country, sender ID, impersonating brand, and text of the smishing message. We built a custom scraper to collect the report date, sender ID, impersonating brand, and the smishing text message from smishing.eu once a week (every Monday) between November 28, 2022, and October 16, 2023. We also collected all old users' posts until November 21, 2021, resulting in 121 smishing user reports. Smishing.eu seized operations on October 16, 2023, and is no longer available.

3.1.4 Pastebin. Individuals use online clipboards to store and share data with others. We investigate one of the most widely used online clipboards – Pastebin, where threat intelligence analysts create public pastes to share data. We find one user who creates pastes to store individual smishing text messages. The same data was reported to the IP abuse reporting platform abuseipdb.com. Fig. 5 in Appendix E shows an example paste. We collect 118 pastes with smishing texts and parse them to collect the sender's mobile

²<https://web.archive.org/web/20230707150805/https://developer.twitter.com/en/products/twitter-api/academic-research>

Table 1: Overview of our smishing dataset collected using posts ($n = 220,585$) and image attachments ($n = 64,284$).

Online Forum	Timeline	Posts	Image Attachments	SMS Messages		Sender IDs		URLs	
				Unique	Total	Unique	Total	Unique	Total
Twitter	2017 - 2023	215,842	60,209	25,517 (92.1%)	31,234	17,162 (88.9%)	26,185	18,306 (91.3%)	23,757
Reddit	2017 - 2023	2,136	1,707	309 (1.1%)	433	202 (1.0%)	326	178 (0.9%)	288
Smishtank	2022 - 2024	2,368	2,368	1,667 (6.0%)	1,963	1,722 (8.9%)	1,871	1,418 (7.1%)	1,682
Smishing.eu	2021 - 2023	121	-	117 (0.4%)	121	115 (0.6%)	121	64 (0.3%)	68
Pastebin	2021 - 2022	118	-	108 (0.4%)	118	113 (0.6%)	114	94 (0.5%)	108
Total		220,585	64,284	27,718	33,869	19,314	28,617	20,060	25,903

number, the timestamp of when the paste was created, and the text of the smishing message, including the URL, where available.

3.1.5 Smishtank. Timko and Rahman run the crowdsourcing website smishtank.com where individuals can report smishing texts as screenshots or text of the smishing message [113]. Their dataset from 2024 contains approximately 1k smishing texts. Additionally, we programmatically collect the updated list of 1,278 smishing user reports from the website between March 31, 2022, and April 8, 2024. These reports include the submission timestamp, sender ID, text of the smishing message, URL, and screenshot of the smishing text where available.

3.2 Smishing Data Curation

We collect over 200k smishing reports and 64k image attachments from five online forums (§3.1). As there is no verification or validation of users' reports or posts on these forums, users may submit images that are not screenshots of smishing texts. Users and organizations also use Reddit and Twitter to raise awareness about smishing instead of reporting or seeking advice; therefore, some posts and images may not be smishing reports.

Much of our data is in the form of screenshots of SMS texts, and we investigate options for extracting the text and associated metadata. Initially, we use Pytesseract to perform object character recognition (OCR) [52]. We find that it fails to work on all images and cannot differentiate between text messages, emails, or other kinds of images. OCR fails to extract text from multiple mobile messaging apps with custom background colors and designs that are available to users.

Threat actors also use various evasion squatting techniques to create domain names that imitate a legitimate domain using similar-looking characters [112]. For example, OCR fails to differentiate between 'l' and 'I'. To overcome this limitation, we consider the method employed by a recent study that extracts spam texts from screenshots posted on Twitter [110] using the Google Vision API. Even though the Google Vision API performs better than traditional OCR in recognizing individual characters, it often fails to preserve the correct reading order, resulting in incoherent text output. It also does not extract the complete URL from smishing images. SMS text consists of multiple lines where the URL spreads across more than one line (Fig. 4 in Appendix E). Incorrect ordering can fail to extract the complete URL.

As a result, we turn to OpenAI's Vision API. Unlike emails, SMSs have limited metadata; we can only extract the time the SMS was received and the sender ID. To that end, we develop

a custom script that utilizes OpenAI's Vision API to extract not only the message text and URL but also the timestamp and sender ID from images, where available. We write and test our prompt (found in Appendix D.1) before we provide it to OpenAI's Vision API. Altogether, we extract the following four variables from the smishing message that we further use for analysis:

Smishing Text. In addition to reporting smishing screenshots, users also post awareness posters and sometimes irrelevant screenshots with the keywords (§3.1) we query to collect posts from online forums. To this end, if the image does not represent an SMS image screenshot, we instruct OpenAI's Vision API to dismiss the image. If the image is indeed an SMS screenshot, then we instruct the API to also return the translated English version if the SMS text is not in English. To this end, we successfully extract the text from all the collected SMS-resembling images.

Timestamp. There is often a delay between when a user receives a smishing SMS and when they report it. Threat actors commonly broadcast the smishing text when they set up the malicious URL, impersonating the targeted brand. To capture a more accurate timestamp in our dataset, we extract the timestamp from the SMS screenshot that the user reports to the online forum, where available. We then use the Python library *dateparser* to parse date/time in various formats depending on the messaging application.

Sender ID. In most cases, users submit the full screenshot of their screen, including the SMS Sender ID. Depending on the smishing campaign, this could be a mobile number, email address, or an alphanumeric code. However, in some cases, users redact the sender ID before posting it on public online forums, likely due to privacy issues. If a sender ID is visible in the screenshot, OpenAI's Vision API extracts it successfully from the image.

URL. We instruct OpenAI's API to extract the URL in the SMS, where available. In some cases, users redact the URL or the short-code of the shortened URL to protect other users from opening the malicious URLs.

Overall, we collect a dataset with 27.7k smishing messages, 19.3k sender IDs, and 20k URLs. We present the breakdown of the variables from each online forum in Table 1.

3.3 Measurement Methods

We analyze the smishing data collected from five online forums to answer our research questions (§1). To this end, we enrich the collected variables (§3.2) to understand the scammer strategies and the underlying infrastructure they abuse. Fig. 1 provides an overview of the enrichment methods and measurements we perform to analyze

the collected data. Table 2 indicates the data sources we use towards our analysis methods, as we explain in the relevant subsections.

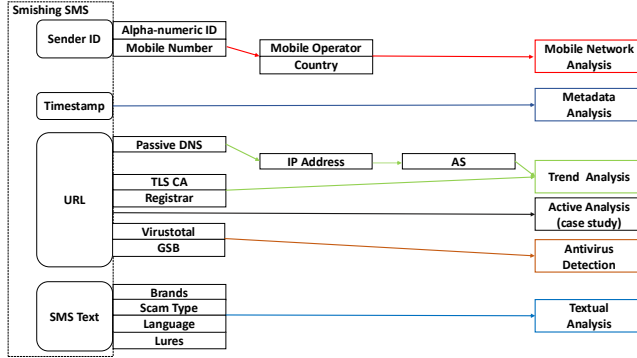


Figure 1: Overview of our analysis methodology.

Table 2: Data sources used in analysis methods (Fig. 1).

Analysis Methods	Data sources
Mobile network analysis	Twitter, Reddit, Smishing.eu, Pastebin & Smishtank
Metadata analysis	Twitter, Reddit & Smishtank
Trend analysis	Twitter, Reddit, Smishing.eu, Pastebin & Smishtank
Active analysis (case study)	Twitter
Antivirus detection	Twitter, Reddit, Smishing.eu, Pastebin & Smishtank
Textual analysis	Twitter, Reddit, Smishing.eu, Pastebin & Smishtank

3.3.1 Mobile Network Analysis. Scammers abuse mobile numbers, email addresses, or alphanumeric sender IDs to send smishing texts to potential victims. Apple allows users to send encrypted texts via email addresses (registered with an iCloud account) over the internet. Fraudsters manipulate SMS aggregators [71] to spoof alphanumeric sender IDs and send smishing texts [114]. To this end, we create regular expressions to differentiate between mobile numbers, email addresses, and alphanumeric sender IDs. As alphanumeric shortcodes are not standardized globally and open-source tools are unavailable, we cannot investigate where emails or alphanumeric shortcode sender IDs originate from or which aggregator service they abuse to run campaigns.

We use the Home Location Register (HLR) lookup [79] to investigate identified mobile numbers. HLR lookup provides detailed information about a mobile number’s current status (live/inactive/dead) and its original and current mobile network operator. We perform a one-time HLR lookup on all the mobile numbers in our dataset. We only focus on the original mobile network operator as we perform the lookup once the complete dataset is collected. As mobile numbers can be re-issued/recycled, the current mobile network operator and status do not represent when they were abused.

3.3.2 Metadata Analysis. We investigate the time of the day per week to understand when scammers prefer to target users. To this end, we analyze the extracted timestamps from the smishing text (§3.2). We exclude the smishing reports from *smishing.eu* and *pastebin* towards this, as the reports on these forums include the date when the smishing text was received, missing the time of the

day. Additionally, in case of smishing screenshots collected from the other forums — Reddit, Twitter, and Smishtank, the timestamp in a few images does not contain the date, probably because the screenshot of the text was taken within the week when it was received. As the date when the user posts/reports the smishing text does not represent the date it was received, we cannot find the day of the week for these smishing texts. Hence, we exclude only these cases from this analysis.

3.3.3 Trend Analysis. We query the extracted URLs to understand what infrastructure criminals abuse to run smishing campaigns, identifying the following:

URL Shortener. Services use URL shorteners to shorten a regular URL that might contain personal information in the HTTP header. For example, this can be a mobile or parcel tracking number in a legitimate text. Scammers have started to abuse these in smishing to evade detection by hiding the redirected phishing URL [125]. To this end, we manually search online for various shortening services and create a list of 33 URL shorteners. We then compare it against the URLs extracted from the smishing texts to identify the various URL shorteners criminals abuse towards smishing campaigns.

Top-level Domain (TLD). Top-level domains are selected by entities based on various factors such as the business being country-specific, global, or an organization. We use the `tld` package available in Python to extract the TLDs of smishing domains [16]. IANA classifies TLDs into six groups: generic, country-code, generic-restricted, sponsored, infrastructure, and test [55]. We use this to identify the various groups of TLDs that criminals abuse for smishing.

Registrar. Domain name registrars manage and sell domain names for a website. The `WHOIS` command returns the registrar details of the queried domain [27]. However, `WHOIS` restricts users to automate the query programmatically for privacy reasons. To this end, we collaborate with `WhoisXMLAPI`’s `WHOIS` API to query all the collected domains to identify the registrars scammers abuse to conduct smishing [122].

TLS Certificate. A TLS certificate is set up to encrypt website interactions and verify the server’s identity. We query `crt.sh`, an open-source website that monitors all publicly issued TLS certificates, through their API [89]. This allows us to collect details of all TLS certificates issued to the smishing domains we collect in our dataset. This service provides the authority that issued the certificate, its date of issue, and the expiration date.

Autonomous System (AS). Autonomous systems (ASes) are groups of networks that follow a unified routing policy. Each AS controls a specific set of IP addresses. Identifying abused ASes and the ones that collude with criminals can help stakeholders, such as law enforcement, take required actions. To this end, we collaborate with Spamhaus to access their passive DNS API that returns all IP addresses each domain in our dataset resolved to in the past year [105]. We then map these IP addresses to their corresponding ASes and countries using the IP to AS Number (ASN) and IP to country database provided by `ipinfo.io` [57].

3.3.4 Antivirus Detection. We query the collected URLs on the two most popular antivirus (AV) detection services — VirusTotal [118] and Google Safe Browsing (GSB) [47] using their public APIs. We query VirusTotal as it lists the public detection results from over 70 AV vendors, providing the flags — malicious or suspicious, from all scanners [117]. As GSB also provides a web platform that offers more details than its API, we also query the domains on the GSB transparency website [46]. However, the their transparency report website did not allow us to programmatically check all URLs using our scripts, stopping us from querying 9,948 URLs. This limits us to look at only the 50% data points for the Transparency Report website’s results.

3.3.5 Active Analysis. As a case study, we manually analyze 200 smishing reports posted by users during the real-time smishing data collection from Twitter (§3.1.1). While we identify the URL shortening services in our dataset to understand the abuse of URL shorteners (§3.3.3), it is not possible to retrieve the redirected URL once the shortened URL stops working (either taken down by the service or scammers). To this end, from the 200 reports, we manually investigate all 145 URLs.

As we manually open the shortened URLs, we find that 18 redirect to webpages trying to download an APK file. We save these 18 APK files and query their hash against the set of hashes provided by AndroZoo [9]. AndroZoo provides researchers with the analysis from tens of different AV products on over 25m Android applications. We do not find any of our hashes in AndroZoo. We then submit the APK samples for analysis to VirusTotal [117]. VirusTotal provides results for all AV scanners that use their naming conventions, but they often mislabel samples [119]. To overcome this challenge, we use Euphony [53] which parses malware labels from VirusTotal reports and returns a single malware family per file.

3.3.6 Textual Analysis. We annotate our dataset on four properties (scam type, language, brands, and lure principles) to shed light on the different tactics scammers use in smishing. To this end, we use OpenAI’s GPT-4o, as it was the best-performing model at the time of this analysis. Our prompt is available in Appendix D.2.

Scam Type. We investigate smishing texts to identify the types of scams. Towards this end, we use OpenAI to categorize the texts into six known SMS scams — Hey mum/dad, Delivery, Banking, Government, Telecom, Wrong number, and Others, along with Spam [6]. This distinction allows us to understand SMS scams (which cause financial harm) apart from spam (which is annoying, but not directly disruptive).

Language. About a third of the messages in our dataset were not written in English. We thus annotate/translate the language of the original text using OpenAI and note the original language used.

Impersonated Brand/Organization. Named-entity recognition (NER) using NLP libraries such as *SpaCy* often fails to detect entities from smishing texts [81]. This is likely due to (1) scammers using special characters and combinations of alphanumeric to evade detection from mobile network operators. E.g., N3tfl!x cannot be detected as Netflix from off-the-shelf models, and (2) NER models not trained to detect entities globally. To this end, we use OpenAI to extract the entity that criminals impersonate.

Scam Lure. We aim to understand how scammers deceive potential victims into smishing. We adopt the seven lure techniques from Stajano and Wilson [107] and annotate the smishing texts with the identified lures.

3.4 OpenAI Evaluation

We extract 150 random messages from our dataset, and two authors label the scam category, impersonated brand, and lures used by the scammer in each smishing text. We use this subset to calculate the inter-rater reliability (IRR) between the two annotators, and then use it as a ground truth to evaluate OpenAI’s model annotation and fine-tune the prompt. We omit non-English smishing texts for the IRR calculation, as English is the only common language between annotators. OpenAI’s model 4o performs relatively well for translation tasks [60] and we expect little variation on short texts.

We use Cohen’s κ [24], the standard metric for IRR, between the two authors. There is near-perfect agreement across all three properties: impersonated brands ($\kappa = 0.82$), scam types ($\kappa = 0.94$), and lure principle ($\kappa = 0.85$). After discussing the disagreements, we develop a consensus ground truth annotated set of 150 texts. We develop a prompt for Open AI to label these properties using this, performing multiple iterations before finalizing the prompt (Appendix D.2). We then analyze the performance of OpenAI against humans. GPT-4o achieves near-perfect agreement for the identified brands ($\kappa = 0.85$) and scam types ($\kappa = 0.93$), with substantial agreement for lure principle ($\kappa = 0.7$).

4 SMISHING INFRASTRUCTURE

Criminals exploit communication channels and web infrastructure to send smishing texts and host phishing websites. In this section, we set out to answer **RQ1** using the passive analysis methods described in §3.3.1, §3.3.3 and §3.3.4.

4.1 Sender-related Information

We collect 692 (3.7%) unique email addresses, 12,299 (65.6%) unique phone numbers, and 5,762 (30.7%) unique alphanumeric shortcodes. Prior work that crowdsourced smishing texts from users only in the US found email addresses (23.9%) being abused more than shortcodes (1%) as sender IDs [113]. On the contrary, we find that scammers abuse alphanumeric shortcodes more than email addresses to send smishing texts. This is likely due to our dataset consisting of texts from across the world, and not just the US, where email-to-text (like iMessage) is popular.

Phone Numbers. Our HLR lookups yield the types of phone numbers exploited in smishing messages (Table 3). There are a wide variety of numbers found which are not able to actually send texts (and thus likely spoofed) and would be easy fodder to block. These hide scammers’ original sender ID and evade detection [17]. For instance, we find landline numbers, random sender IDs with more digits than the maximum in a valid number in any country, and voicemail-only numbers. A collective group of law enforcement recently took down one such service that allowed mobile number spoofing to call and send messages [37]. Scammers have also started recently abusing SMS blasters, i.e., fake base stations to send messages that allow sender ID spoofing [23, 28].

Table 3: Types of phone numbers abused as sender IDs to conduct smishing ($n = 12,299$).

Type	Phone Numbers
Valid Numbers	
Mobile	8,209 (66.7%)
Mobile or Landline	283 (2.3%)
VOIP	249 (2.0%)
Toll Free	73 (0.6%)
Pager	8 (0.1%)
Universal Access Number	5 (0.0%)
Personal number	2 (0.0%)
Others	11 (0.1%)
Invalid/Suspicious Numbers	
Bad Format	2,991 (24.3%)
Landline	466 (3.8%)
Voicemail Only	2 (0.0%)

Mobile Network Operators. HLR lookup on unique mobile numbers provides us with their original mobile network operators (§3.3.1). We find that Vodafone is the most abused mobile network operator with scammers using their network to send smishing texts from 18 countries (Table 4). While Airtel is primarily abused for banking and telecom scams, we find that scammers prefer Vodafone for banking and delivery scams.³ This indicates that scammers prefer different mobile network operators to conduct various scams, likely based on their target countries.

Table 4: Top 10 mobile network operators abused to send smishing messages.

MNOs	Mobile #s	Countries
Vodafone	1,166 (13.3%)	ESP, IND, GBR, NLD, AUS, CZE, DEU, GHA, HUN, IRL, ITA, NZL, PRT, QAT, ROU, TUR, UKR, ZAF
AirTel	953 (10.9%)	IND, COD, KEN, LKA, MWI, NGA
BSNL Mobile	676 (7.7%)	IND
Reliance Jio	493 (5.6%)	IND
O2	429 (4.9%)	GBR, DEU, IRL
T-Mobile	396 (4.5%)	USA, NLD, CZE
Lycamobile	262 (3.0%)	NLD, BEL, ESP, FRA, AUS, DEU, IRL
SFR	192 (2.2%)	FRA, GLP
KPN Mobile	190 (2.2%)	NLD
EE Limited	184 (2.1%)	GBR

Takeaway. This subsection addresses **RQ1** by indicating that scammers abuse Vodafone in multiple countries to send smishing texts, followed by Airtel. They abuse multiple mobile network operators, but strategically choose different ones depending on the scam campaign. We also find that scammers spoof sender IDs and abuse email addresses to send smishing via online encrypted messaging such as RCS and iMessage [7].

³Our dataset includes the original mobile network operator for every scam text (Appendix C).

4.2 URL Shorteners

Scammers abuse URL shorteners in smishing texts to evade detection from mobile network operators and threat intelligence companies [81, 113, 125], similar as for phishing [64]. URL shorteners can also make it challenging for users to differentiate between legitimate and smishing texts. We find 27 abused URL shortening services to hide the redirected phishing websites. `bit.ly` is preferred for all scam types. While `is.gd` is the second most preferred shortener for banking scams, scammers prefer `cutt.ly` for delivery and government impersonation scams (Table 5).

We also identify 205 `wa.me` URLs (likely) from from scammers asking users to initiate a conversation over WhatsApp. This helps them evade detection from mobile network operators. ‘Hey mum/dad’ scams use this approach [5].

Table 5: Top 10 URL shorteners abused per scam type (B: Banking, D: Delivery, G: Government, T: Telecom, W: Wrong Number and H: Hey mum/dad).

URL Shorteners	URLs	Scam Types					
		B	D	G	T	W	H
<code>bit.ly</code>	1,830 (30.6%)	1,140	112	176	104	6	-
<code>is.gd</code>	1,023 (17.2%)	970	19	7	7	-	-
<code>cutt.ly</code>	516 (8.7%)	310	86	44	11	-	-
<code>tinyurl.com</code>	443 (7.4%)	326	37	32	9	-	-
<code>bit.do</code>	404 (6.8%)	254	40	31	25	-	-
<code>shrtco.de</code>	271 (4.5%)	269	-	-	-	-	-
<code>rb.gy</code>	230 (3.9%)	199	12	11	-	-	-
<code>t.ly</code>	172 (2.9%)	112	20	23	2	-	-
<code>bitly.ws</code>	161 (2.7%)	153	1	3	-	-	-
<code>t.co</code>	157 (2.6%)	94	32	12	2	1	-

Takeaway. Scammers abuse third-party URL shortening services, particularly `bit.ly`, to conduct smishing, answering **RQ1**. We also highlight that scammers prefer different URL shortening services depending on the type of scam. For example, the second most preferred for banking is `is.gd`, but not for other scam types.

4.3 Top-level Domains (TLDs)

We find over 280 top-level domains (TLDs) that scammers abuse to conduct smishing. The most abused TLD is `.com`, followed by `.info`, consistent with previous findings [8, 113].

The majority 7,539 (72.33%) of the URLs abuse gTLDs, followed by 2,829 (27.14%) abusing ccTLDs (Table 16 in Appendix F). Criminals likely select gTLDs based on the brands and sectors they target. For example, using `.online` for brands that impersonate online technical companies such as Facebook (`fb.user-page[.online]`). Prior work has investigated specific ccTLDs for phishing abuse [80]; we find that scammers abuse 130 ccTLDs towards smishing.

Some of this is from scammer registered websites; others are from free website building services such as Google’s Firebase, ngrok, Heroku which allow them to more easily deploy smishing websites [101]. There are a few reasons behind their popularity here.

Table 6: Top 10 TLDs abused to set up unique smishing URLs ($n = 10,423$).

TLDs	Smishing URLs	TLDs	Shortened URLs
com	4,951	ly	2,482
info	574	com	383
in	404	gd	352
me	291	do	311
net	286	gy	233
co	234	de	170
top	225	co	137
us	202	ws	122
online	201	cc	68
xyz	159	fr	39

First, these services are free of cost, saving scammers costly domains and hosting infrastructure. Second, services like these provide the advantage of quickly spinning up a web application (often using phishing kits [18] to impersonate a brand). We identify 303 web.app domains, 186 ngrok.io domains and 184 domains with five other TLDs – app, firebase.app, vercel.app, herokuapp.com and netlify.app. As users cannot identify domains with unusual TLDs to conduct phishing [94], scammers misuse various TLDs to register domains. The insights on TLDs abused towards smishing will help stakeholders update their policy frameworks to prevent this threat.

Takeaway. We identify that scammers prefer to abuse the .com TLD to register smishing domains, followed by .info TLD, addressing **RQ1**.

4.4 Registrars

Scammers register new domains to host phishing websites that they abuse for smishing. The most abused registrar in our dataset is GoDaddy, followed by NameCheap. We provide the top 10 registrars that criminals abuse to purchase smishing domains in Table 17 in Appendix F. Previous work that crowdsourced limited user reports found that scammers abuse NameCheap the most towards smishing [113]. Another study on phishing domains identified GoDaddy as the most abused registrar [112]. While banking, delivery, and telecom scams exploit GoDaddy the most, scammers prefer to abuse Gname over other registrars for government impersonation scams. These insights help inform stakeholders such as ICANN towards targeted intervention and refining registrar-specific policies.

Takeaway. This subsection answers **RQ1** by highlighting that scammers abuse GoDaddy the most to register smishing domains, followed by Namecheap. These are well-known registrars whose affordability and ease of registration make them convenient entry points for malicious activity.

4.5 TLS Certificates

We find 263,318 TLS certificates issued to 6,766 domains by 357 different TLS issuer IDs corresponding to over 100 issuing organizations. Cybercriminals sometimes use multiple TLS certificates for smishing URLs, similar to phishing [18]. We identify TLS certificates with between 1 and 4,681 per URL (mean: 39, median: 4). The

most abused certificate authority (CA) is Let’s Encrypt, followed by DigiCert and cPanel (Table 7). While Let’s Encrypt and cPanel provide free certificates, DigiCert charges money even for the basic TLS certificate. We see that Let’s Encrypt is the most abused for both the number of certificates and the domains they are issued to, while Sectigo is the second most abused in terms of the number of domains with relatively fewer certificates issued. This is likely due to Sectigo charging fees to provide features like a much longer validity period and multi-domain TLS.

The preference for Let’s Encrypt is unsurprising: it issues TLS certificates at no cost, multiple hosting platforms use its services to provide TLS certificates, and most of the Internet uses its services. Previous work investigating maliciously registered domains supports this finding [62]. Certificates issued by Let’s Encrypt are only valid for 90 days, which likely inflates their numbers (Table 7). With TLS certificate validity lengths being reduced to 47 days, the number of certificates issued by a domain will increase further [26].

Table 7: Top 10 TLS certificate authorities abused to run smishing campaigns.

Certificate Authority	Certificates	Domains
Let’s Encrypt	141,878	4,773
DigiCert	19,340	736
cPanel	17,619	915
Google Trust Services	16,712	957
Globalsign	15,341	144
Comodo	14,128	250
Amazon	7,746	273
Entrust	6,599	73
Sectigo	6,477	1,372
Cloudflare	4,075	683

Takeaway. Scammers primarily abuse Let’s Encrypt for TLS certificates to set up smishing websites, followed by DigiCert and cPanel, addressing **RQ1**. Table 7 indicates that the abuse is unevenly distributed across certificate authorities. In particular, scammers exploit free, automated, and widely accessible services as a low-barrier option to obtain TLS certificates for malicious websites quickly.

4.6 Autonomous Systems (ASes)

We find 466 domains that resolve to 1266 IP addresses found by our passive DNS queries. Cloudflare controls 487 of these IP addresses, corresponding to 88 domains. Criminals abuse Cloudflare as it provides a free proxy service to hide their actual IP addresses. Amazon, Akamai, and Google are the next most prominent ASes— three more proxies/cloud providers (Table 8). Prior work found Amazon as the top hosting provider [81]. There is a likely bulletproof hosting provider in our top 10 – Frantech Solutions [67]. Bulletproof hosting providers (BHP) offer hosting services to criminals for conducting illicit activities and evade enforcement by ignoring or delaying legal requests, or are based in unreachable jurisdictions [10, 11, 68]. We also have IP addresses that belong to other known BHPs, like Proton66 OOO (AS198953 with IPs in RU) [8, 20] and Stark Industries (AS44477 with IPs in NL) [32]. Passive DNS on

identified IPs that belong to BHP’s AS could help unveil additional maliciously registered domains.

Table 8: Top 10 ASes abused to host smishing web pages along with the host countries

AS Name	IPs	ASNs	Countries
Amazon	188	AS16509, AS14618	US, JP, IE, IN, MA
Akamai	147	AS63949	US, IN
Google	59	AS15169, AS396982	US
Multacom	49	AS35916	US
SEDO GmbH	31	AS47846	DE
Alibaba	16	AS45102, AS37963	HK, US, CN
Tencent	15	AS132203	US, DE
FranTech Solutions	11	AS53667	US, LU
HKBN Enterprise	11	AS17444	HK
The Constant Company	11	AS20473	US

Takeaway. This subsection answers **RQ1** by indicating that 18.8% of domains that resolve to an IP address abuse Cloudflare to evade detection. While scammers prefer to abuse traditional ASes such as Amazon the most, certain threat actors also use bulletproof hosting providers to resist takedowns and evade law enforcement actions.

4.7 Antivirus detection

Over 8,911 (44.9%) smishing URLs are not marked as malicious or suspicious by any antivirus (AV) scanners on VirusTotal (Table 9). While more than 9.8k (49.6%) URLs are marked malicious by at least one AV vendor on VirusTotal, only 56 (0.3%) URLs are marked malicious by more than 15. Additionally, more than one AV vendor marks 3,574 (18%) as suspicious. This confirms previous evidence indicating that different providers build their blocklists in different ways [40]. Importantly, our work focuses on URLs abused in smishing; few AV scanners are popular in the mobile ecosystem.

Table 9: VirusTotal detection results for all smishing URLs ($n = 19,864$).

VirusTotal Results	URLs
Malicious = 0 and Suspicious = 0	8,911 (44.9%)
Malicious ≥ 1	9,851 (49.6%)
Malicious ≥ 3	5,136 (25.9%)
Malicious ≥ 5	3,236 (16.3%)
Malicious ≥ 10	727 (3.7%)
Malicious ≥ 15	56 (0.3%)
Suspicious ≥ 1	3,574 (18.0%)
Suspicious ≥ 3	31 (0.2%)
Suspicious ≥ 5	0 (0%)

While Google Safe Browsing (GSB) is listed as a scanner on VirusTotal, prior work notes that there are inconsistencies between VirusTotal and vendor’s own scanners [8, 91]. To this end, we find that GSB’s public API detects only 191 (1%) URLs (Table 18 in Appendix F), while it marks 319 (1.6%) URLs malicious on VirusTotal. While GSB continuously updates its API, VirusTotal submissions are less frequent. Contrarily, GSB’s transparency report website returns 802 (8.1%) URLs as unsafe and 440 (4.4%) as partially unsafe. The partially unsafe likely indicates their priority to avoid

blocklisting entire domain when only a single page or subdomain contains malicious content. GSB cannot detect 5,883 (59.3%) URLs, but returns ‘no available data’ for 2,823 (28.5%) URLs.

Takeaway. This subsection addresses **RQ1** by supporting the abused infrastructure detected by AV vendors. It also seconds previous work showing AV vendors build their blocklists in different ways [8, 40]. While the majority of URLs (49.6%) are marked malicious by at least one vendor on VirusTotal, only 0.3% are marked by more than 15.

5 SCAMMER STRATEGIES

Criminals employ various tactics to target users via smishing. In this section, we analyze the timestamp (§3.3.2), mobile network operators’ originating countries (§3.3.1) and the text of the smishing SMS (§3.3.6) to understand how users are lured into taking an action, answering **RQ2**.

5.1 Timestamps

We extract the complete timestamp from 9,432 smishing reports. To understand when scammers prefer to send smishing texts, we plot the time of the day per week when scammers send a smishing text in Fig. 2.

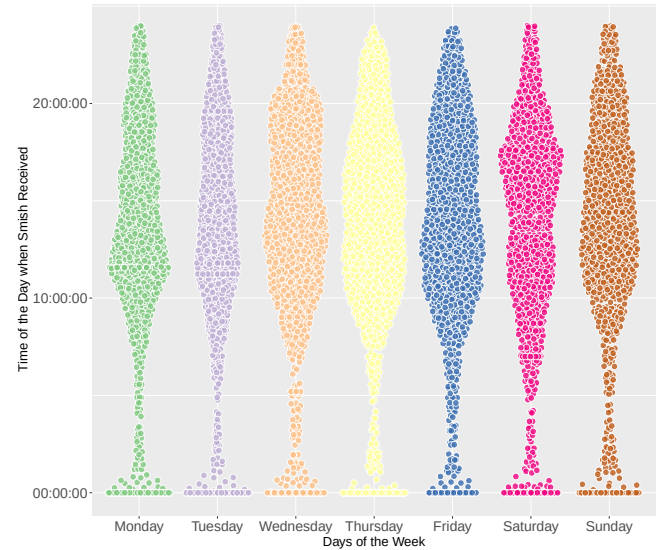


Figure 2: Time of the day per week when scammers send a smishing text ($n = 8,580$). We remove the 2021 campaign.

We find that most scammers are active between 09:00 and 20:00 on weekdays (with medians: Mon – 12:38:00, Tues – 12:26:00, Wed – 14:36:30, Thurs – 14:24:30, Fri – 13:17:00, Sat – 14:38:00, Sun – 13:19:00). This shows that scammers send messages throughout the day when a victim would be busy at work and might make a rushed decision. This aligns with prior work that interacted with ‘Hey mum/dad’ scammers during weekdays, found scammers actively engage between 10:00-15:00 [5]. To understand if the scammers send texts in a similar pattern everyday, we perform a two-sample Kolmogorov-Smirnov test. We identify that the distribution

of sending smishing texts on Monday, Tuesday, Wednesday, and Saturday is different, with a significant p-value for these combinations ($p < 0.05$).

We identify a popular smishing campaign from 2021 that targeted users in India at 11:34 a.m. on Tue, Aug 3rd. Over 850 messages from Tue at 11:34 a.m. belong to this campaign from our Twitter dataset. The smishing text impersonates a popular financial institution (SBI) in India and provides a malicious URL aiming to steal the users' banking login credentials. To avoid the distribution on Tuesday from being skewed, we remove this campaign from Fig. 2.

Takeaway. We identify that most smishing texts are sent between 09:00 and 20:00 on weekdays. This highlights that scammers strategically exploit times when users are busy and prone to taking irrational decisions, answering RQ2.

5.2 Scam Categories

Our smishing texts primarily target the banking sector (45.12%), followed by delivery/parcel companies and government agencies. This distribution is shown in Table 10. While we strategically select keywords to collect user reports (§3.1), we still discover over 1.7k spam texts. This indicates that data collected via user reports needs to be checked for spam before being marked as scam. Note that we also find over 6.9k messages that belong to the 'Others' category. While differentiating this is future work, we manually investigate a random subset and find job-related conversation scams, investment-related conversation scams, cryptocurrency scams, OTPs (potentially a call-back scam), and tech company-impersonation scams for companies like Telegram, Facebook, and Netflix.

Table 10: Distribution of collected smishing messages ($n = 33,869$) into eight categories, including spam.

Scam Category	Messages	Top 4 Languages
Banking	15,277 (45.1%)	en, es, nl, it
Delivery	3,810 (11.3%)	en, es, de, fr
Government	3,248 (9.6%)	en, fr, es, nl
Telecom	2,226 (6.6%)	en, fr, es, nl
Wrong number	332 (0.9%)	en, ja, id, es
Hey mum/dad	263 (0.8%)	en, de, es, nl
Others	6,944 (20.6%)	en, es, fr, nl
Spam	1,710 (5.0%)	en, es, id, tl

Prior work that analyzed data concentrating on a single UK mobile network operator found that delivery-themed smishing texts were the most prominent scam type [6]. While this might be true in the UK, our data is not limited to a particular region and focuses on smishing texts that reach the end user, bypassing mobile network operators' detection systems, where implemented. Scammers impersonate banking institutions across the globe and lure users into providing their confidential details to steal their funds.

Delivery-themed smishing texts impersonate popular international and regional postal entities such as 'USPS' and 'Correos' and deceive users into providing their credit card details along with personal information that is either sold online or abused towards card-not-present fraud [19]. Similarly, for government or telecom

scams, scammers lure users into providing their private or financial details, which could be used for identity theft. The last two categories — 'Wrong number' and 'Hey mum/dad' scams — are conversational scams that lure users into replying or initiating a conversation, gain trust, and deceive users into investing in fake cryptocurrency schemes or requesting funds [5, 97].

Takeaway. This subsection answers RQ2 by indicating that scammers use banking impersonation the most to conduct smishing, followed by delivery and government impersonation scams. This suggests that scammers exploit users' trust in essential legitimate services such as banks, increasing their likelihood of becoming victims of smishing.

5.3 Text Languages

Most smishing messages we collect are written in English (65.3%), followed by Spanish (13.7%) (Table 11). While we detect 66 languages, only 13 have over 100 messages. This aligns with prior work [81], which primarily found English texts targeting the US and UK.

Table 11: Top 10 languages used in smishing messages ($n = 33,869$) and the most spoken languages [35].

Language	Code	Messages	Language	Population (m)	Countries
English	en	22,078 (65.2%)	English	1,500	46
Spanish	es	4,639 (13.7%)	Mandarin Chinese	1,200	2
Dutch	nl	1,945 (5.7%)	Hindi	609	2
French	fr	1,163 (3.4%)	Spanish	558	21
German	de	810 (2.4%)	Arabic	335	24
Italian	it	669 (1.9%)	French	312	29
Indonesian	id	347 (1.0%)	Bengali	284	2
Portuguese	pt	280 (0.8%)	Portuguese	267	9
Japanese	ja	257 (0.8%)	Russian	253	4
Hindi	hi	175 (0.5%)	Indonesian	252	2

This language distribution does not reflect the global population. Mandarin Chinese is the second most spoken language worldwide, yet it only accounts for 46 messages in our dataset — less than 0.2%. Whereas Dutch appears in nearly 2k reported messages, despite not ranking among the world's most spoken languages. This mismatch could be due to the nature of the platforms from which we collect data (e.g., Reddit, Twitter), which tend to contain posts mostly in English [31, 108]. We find many smishing campaigns that target non-native English speakers yet use English in the smishing text [54]. E.g., we only find 176 smishes in Hindi yet SBI is our top impersonated brand (§5.4). This is likely due to global organizations increasingly using English for their communications [102] and scammers adopting that norm.

We investigate the intersection of scam category and language (Table 10). Other than English, Spanish is a popular language with banking and delivery scams and French with government and telecom scams. We also notice some 'wrong number' scams in Indonesian, Japanese, and Chinese.

These findings suggest that scammers may adapt scams that work better for audiences differentiated by language. For instance, 'Hey mum/dad' scams predominantly target English, German, Spanish, and Dutch users [5], which may reflect Western family dynamics or communication norms exploited in such attacks. This adds weight to the work of Simoiu et al. [103], who conclude that attacks

often focus, among other things, on individuals' risk level, including age, locality, and even prior security incidents. More work by social scientists is required to understand scammer behavior.

Takeaway. We find that scammers send smishing texts primarily in English (65.2%), followed by Spanish. This showcases that scammers continue to adapt as global organizations increasingly send texts to users in English, addressing **RQ2**.

5.4 Targeted Brands

The most frequently impersonated brands in our dataset are financial institutions based in India — SBI, PayTM, and HDFC (Table 12). This finding feeds into our broader discussion on the global use of the English language by organizations (§5.3). The majority of texts impersonating these organizations are written in English; it is one of India's official languages.

Smishing texts impersonating Santander are in Spanish, followed by English and Portuguese, while those targeting Amazon are in English, followed by Spanish and Japanese. The ones that impersonate IRS are in English, followed by Spanish. Messages that impersonate Netflix are written in English, followed by French and Spanish. Smishing targeting regional entities uses native languages, e.g.: texts targeting Rabobank are in Dutch, BBVA in Spanish, and CaixaBank are in Spanish and Portuguese.

The overwhelming majority of organizations being impersonated by scammers feature financial institutions, in line with §5.2. Even though delivery scams are second most popular, we do not find any service in the top 10 due to the diversity of various impersonated brands in the sector, such as Correos, DHL, and USPS.

Table 12: Top 10 brands that scammers impersonate to lure victims via smishing ($n = 33,869$).

Brand Name	Category	Messages
State Bank of India (SBI)	Banking	3,925 (11.6%)
PayTM	Banking	1,001 (3.0%)
Housing Development Finance Corporation (HDFC)	Banking	974 (2.9%)
Santander (BNC, SAN)	Banking	519 (1.5%)
Amazon (AMZ)	Others	460 (1.4%)
Internal Revenue Service (IRS)	Government	418 (1.2%)
Rabobank	Banking	382 (1.1%)
Banco Bilbao Vizcaya Argentaria (BBVA)	Banking	363 (1.1%)
Netflix (NFLX)	Others	361 (1.1%)
CaixaBank	Banking	326 (1.0%)

Takeaway. In line with our findings of scam categories (§5.2), we find four banks as the most targeted brands, followed by Amazon and the IRS, addressing **RQ2**.

5.5 Scam Lures

We define scam lures using Stajano and Wilson's typology [107] in Table 13. While previous work has identified the lures scammers use to deceive victims of cryptocurrency fraud [4, 104], we present the various lures criminals use to create smishing texts. We find that conversation scams like 'Hey mum/dad' and 'Wrong Number' lure victims into replying to their initial scam text by applying distraction and kindness lures. The 'Hey mum/dad' scam also abuses time and urgency, supporting prior work [5].

Smishing campaigns that impersonate a delivery company, mobile network operators, government agency, or bank primarily employ the authority principle as they pretend to be a trusted entity. These scams create a false sense of legitimacy, often urging the user into taking a hasty decision by showing urgency (Time/Urgency lure). Some campaigns in these categories also appeal to a victim's interests or greed, such as those offering tax refunds, leveraging the need and greed lure. Unsurprisingly, the lure least used by scammers in our dataset is dishonesty (0.5%) as it relies on the victim's complicity, which is less plausible in unsolicited SMS campaigns. Victims are unlikely to engage if they recognize the action as fraudulent from the outset. Similarly, for the herd lure, we find only 393 messages (1.2%) as smishing does not tend to convince victims to take risks others supposedly have taken (unlike spam).

Takeaway. This subsection addresses **RQ2** by highlighting that scammers use time/urgency lures in all smishing texts, except 'Wrong Number' scams. While they use distraction and kindness lures for 'Hey mum/dad' and 'Wrong number' scams, they prefer to deceive users into banking, delivery, government, and telecom scams using authority and need & greed lures.

5.6 Sender ID Originating Countries

Most smishing texts were sent from mobile numbers belonging to Indian mobile network operators followed by US mobile network operators (Table 14). While this directly shows what countries send smishes, it potentially shows which countries receive smishes. Further work needs to be done to validate this hypothesis, particularly given the different SMS filters for different countries. India does not have a reporting mechanism for SMS phishing, likely contributing to their prominence in our dataset. However, our dataset also contains mobile network operators from countries like the US, UK, and Australia, where special SMS reporting services like 7726 exist [3, 42, 84]. This indicates that users still report these messages to online discussion forums, suggesting lack of awareness of official reporting channels [83].

We identify that numbers originating from mobile network operators in India are primarily abused towards banking scams (see Fig. 3). Meanwhile, mobile numbers from the US are abused to send scam texts belonging to the 'Others' category, followed by banking and delivery scams. The 'Others' category includes texts impersonating various services and tech companies such as Netflix, Amazon, and Facebook, cryptocurrency-related scams, and conversation scams. We observe a similar trend in Indonesia, where the 'Others' category dominates. These include impersonation of services such as WhatsApp and Telegram, as well as conversation scams like fake recruitment and investment opportunities.

Takeaway. We discover that most smishing texts originate from mobile numbers belonging to Indian mobile network operators and are primarily abused to send banking scams. Whereas scammers prefer to abuse mobile numbers from the US mobile network operators to impersonate platforms such as Netflix, Amazon, and Facebook and conduct other scams, answering **RQ2**.

6 CASE STUDY: MALWARE VIA SMISH

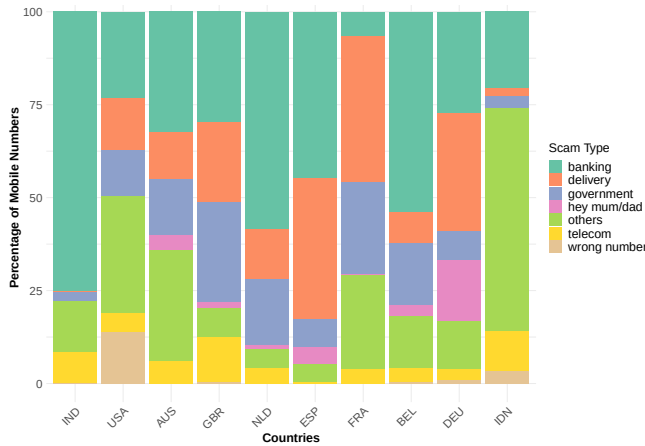
Scammers target victims through smishing by providing a malicious URL in the SMS text. While this is primarily known to

Table 13: Description of lures scammers use to deceive victims by scam category (B: Banking, D: Delivery, G: Government, T: Telecom, W: Wrong Number and H: Hey mum/dad).

Lure	Definition	Scam Types					
		B	D	G	T	W	H
Authority	Scammers refer to trusted third parties to convince users to comply	✓	✓	✓	✓		
Dishonesty	Scammers invite users willingly and knowingly into taking fraudulent action						
Distraction	Scammers provide unrelated details to distract the user		✓			✓	✓
Need & Greed	Scammers leverage users' greed and offer attractive benefits	✓	✓	✓	✓		
Herd	Scammers convince that others have won taking the same risk						
Kindness	Scammers leverage the willingness of people to help others					✓	✓
Time & Urgency	Scammers put time pressure on users so they make an irrational decision	✓	✓	✓	✓		✓

Table 14: Top 10 countries by sender ID mobile numbers.

Originating Countries	MNOs	Mobile Numbers	
		All	Live
India (IND)	10	2,722	396
United States of America (USA)	72	1,369	281
Netherlands (NLD)	7	801	229
United Kingdom (GBR)	15	767	138
Spain (ESP)	9	494	361
Australia (AUS)	6	392	151
France (FRA)	14	387	202
Belgium (BEL)	6	271	85
Indonesia (IDN)	6	216	28
Germany (DEU)	5	187	70

**Figure 3: Distribution of scam types for the top 10 mobile network operator originating countries.**

redirect to a phishing website, criminals also spread Android malware through smishing [44, 120, 121]. We find 18 malicious APK files while manually inspecting 145 URLs from 200 smishing texts (§3.3.5). We notice that these URLs are designed to redirect depending on the user's device and operating system. For example, `shrtco[.]de/2Rq2La`, when opened on a desktop browser, redirects to `sa-krs[.]web[.]app/`, which displays a smishing webpage impersonating a bank. However, if opened using an Android

device, it redirects to `sa-krs[.]web[.]app/?d=s1` and automatically downloads an APK file named `s1.apk`. 24 antivirus scanners on VirusTotal mark this APK file malicious.⁴

Our case study suggests that 'SMSspy' [70, 124] is the most common malware that scammers use to target users through smishing (see Table 19 in Appendix G). Criminals target victims' mobile phones using these malicious applications to steal SMS for one-time passcodes (OTPs). Additionally, we discover 89 further URLs ending with '.apk' in our dataset. For example, `download[.]china-telecom[.]cn/internet.apk` and `ceskaposta[.]online/PostaOnlineTracking.apk`.

Takeaway. We highlight that scammers can also spread malicious APK files through smishing texts, an emerging threat, which addresses RQ2. This indicates one way in which criminals target victims' mobile phones to steal SMS-based one-time passcodes – through malware such as 'SMSspy.'

7 DISCUSSION AND CONCLUSION

Smishing has led to a significant financial loss to users globally. The unavailability of updated data restricts researchers from studying this cybercrime. While previous work has crowdsourced a small amount of smishing reports [92, 113] or studied online public gateways [81], this paper uses a novel methodology to collect smishing texts from five public online forums and contributes a pseudo-anonymized, updated, and labeled smishing dataset. We identify the infrastructure criminals abuse to conduct smishing campaigns – mobile network and web hosting ecosystem (§4) and the tactics scammers employ to lure victims (§5). While previously URL shortening service APIs allowed retrieving redirected phishing websites from taken-down shortened URLs [49, 51, 64], they have restricted their APIs. Actively measuring smishing URLs could help identify malicious APKs [100] (§6) and capture phishing kits scammers use to set up phishing websites quickly [18]. Measuring smishing data from online forums could help stakeholders, such as regulators, inform policies and take-down companies, prioritize their actions for efficiently fighting against this threat.

7.1 Limitations

As with any measurement study, our methodology has limitations. Like other social media studies, we cannot collect all entries that

⁴<https://www.virustotal.com/gui/file/34ae95c0a19e3c72f199c812f64dc8f38bbc7f0f5746efe0bd756737163ed8ec/detection>

report smishing. Users can delete content before we collect their historical data, though not before we collect our real-time data. Not all relevant posts use our keywords. This would result in us underestimating the volume of smishing. Particularly, we only use English keywords, biasing our data towards users who speak this language.

As with any data collected from online forums, we also encounter certain biases. Our dataset is biased towards countries where mobile network operators do not have scam/spam text blocking mechanisms and users are more active in reporting smishing on these forums.⁵ For example, the 2021 smishing campaign from India. Twitter and Reddit also made their APIs inaccessible during the time period of our study. There could be other community-specific forums where users from various countries prefer to report smishing texts. As our dataset contains smishing messages primarily until 2023, it may not reflect more recent trends, particularly in the modern LLM-enhanced threat landscape. Despite these limitations, we collect a large updated smishing dataset and, by overcoming several measurement challenges, we draw insights into this ecosystem.

7.2 Recommendations and Mitigations

Based on our dataset and measurements, we suggest countermeasures to potentially mitigate smishing.

Technical Measures. Researchers could use our labeled dataset with new features such as scam typologies to develop multi-class detection models, as prior work predominantly relies on decade-old spam/ham datasets to build binary classifiers (§2). Mobile network operators should implement checks for shortened URLs in texts for redirection to abused domains in their XDR filtering solutions to identify and block scam texts [69] to circumvent SMS fraud abusing this technology. Mobile network operators should also deploy XDR filtering widely. Official reporting services such as 7726, and one-click reporting [45], currently in limited countries such as the UK [84, 85], should be enabled broadly, since we demonstrate that users actively report smishing texts. Online forums like Twitter should have automated algorithms to identify and share user-reported smishing texts with stakeholders.

Registrars such as GoDaddy and NameCheap (§4.4) should check for prior abuse of domains and restrict domains that could be abused to impersonate popular brands before (re)issuing domains [8]. URL shortening services such as bit.ly and is.gd (§4.2) should use threat intelligence to check for domain abuse before providing services. CAs such as Let's Encrypt in the past have been known to use Google Safe Browsing results before issuing TLS certificates [1, 2]. Updating this approach by incorporating more relevant data sources, particularly maliciously registered domains, will result in better detection. Involving TLS CAs as a stakeholder in the ecosystem to prevent scams has been a hot topic of debate. The abuse of their services (§4.5) suggests that CAs are an important stakeholder and should work collectively with others to potentially mitigate cybercrime.

Government. While countries like the UK and Spain have been actively working towards tackling sender ID spoofing (§4.1) [72, 85], telecom regulators worldwide should create sender ID registries to

detect and stop shortcode abuse [77]. Criminals purchase services from underground forums (Fig. 6 in Appendix H), including Telegram [98] to send bulk smishing texts using illicit devices such as SIM Farms/Boxes.⁶ Similar to laws in the UK, regulators in other countries should ensure that the sale of these devices is restricted to legitimate use, like call centers [116]. Regulators should work with mobile network operators to implement effective know-your-customer (KYC) checks that will significantly decrease the abuse of shortcodes and mobile numbers to conduct SMS scams [36, 41]. International cooperation among law enforcement agencies is required to take down bulletproof hosting providers [63] who provide hosting services to criminals (§4.6).

Educational. Scammers create well-crafted smishing texts to deceive victims into taking action. Educating users about the lures we identify (§5.5) can help potential victims avoid falling prey to such scams. Academic researchers could use our labeled dataset and the insights from this paper to build educational awareness tools that could be adopted by industry stakeholders. Text messaging platforms like Google and Apple can contextualize their *potential scam warnings* shown to users by explaining the lure used [48].

ACKNOWLEDGMENTS

We thank the reviewers and the anonymous shepherd for their feedback. Additionally, we thank Spamhaus, HLRLookup.com, IPInfo, WhoisXMLAPI, and Google for providing researcher access to their services. This project is supported by the UCL Computer Science Strategic Research Grant. Dr. Antonis Papasavva was funded by the Dawes Trust. Guillermo Suarez-Tangil is a 2020 RyC fellow RYC2020-029401-I funded by MICIU/AEI/10.13039/501100011033 and the ESF Investing in your future. His work was also supported by the project PID2022-143304OB-I00 funded by MICIU/AEI/10.13039/501100011033/ and by the ERDF, EU.

REFERENCES

- [1] Josh Aas. 2015. The CA's Role in Fighting Phishing and Malware. <https://letsencrypt.org/2015/10/29/phishing-and-malware.html>.
- [2] Josh Aas. 2019. Let's Encrypt No Longer Checking Google Safe Browsing. <https://community.letsencrypt.org/t/let-s-encrypt-no-longer-checking-google-safe-browsing/82168>.
- [3] ACCC. 2010. Spam SMS: Report it! <https://www.scamwatch.gov.au/about-us/news-and-alerts/spam-sms-report-it>.
- [4] Sharad Agarwal, Gilberto Atondo-Siu, Marilyne Ordekian, Alice Hutchings, Enrico Mariconti, and Marie Vasek. 2023. Short Paper: DeFi Deception—Uncovering the Prevalence of Rugpulls in Cryptocurrency Projects. In *International Conference on Financial Cryptography and Data Security*. Springer Nature Switzerland, Cham, 363–372.
- [5] Sharad Agarwal, Emma Harvey, Enrico Mariconti, Guillermo Suarez-Tangil, Marie Vasek, et al. 2025. 'Hey mum, I dropped my phone down the toilet': Investigating Hi Mum and Dad SMS Scams in the United Kingdom. In *34th USENIX Security Symposium (USENIX Security 25)*. USENIX Association, Seattle, WA, USA.
- [6] Sharad Agarwal, Emma Harvey, and Marie Vasek. 2024. Poster: A comprehensive categorization of sms scams. In *Proceedings of the 2024 ACM on Internet Measurement Conference (IMC '24)*. Association for Computing Machinery, New York, NY, USA, 755–756.
- [7] Sharad Agarwal, Guillermo Suarez-Tangil, and Marie Vasek. 2025. An Overview of 7726 User Reports: Uncovering SMS Scams and Scammer Strategies. <https://arxiv.org/abs/2508.05276>. arXiv:2508.05276 [cs.CR]
- [8] Sharad Agarwal and Marie Vasek. 2025. Examining Newly Registered Phishing Domains at Scale. In *In 24th Workshop on the Economics of Information Security (WEIS)*. WEIS, Tokyo, Japan.

⁵This might have changed since our data collection concluded.

⁶Law enforcement arrests show use of SIM boxes abused to send smishing campaigns. [29, 56]

- [9] Kevin Allix, Tegawendé F. Bissyandé, Jacques Klein, and Yves Le Traon. 2016. AndroZoo: Collecting Millions of Android Apps for the Research Community. In *Proceedings of the 13th International Conference on Mining Software Repositories (MSR '16)*. IEEE, Austin, TX, USA, 468–471.
- [10] Sumayah Alrwais, Xiaojing Liao, Xianghang Mi, Peng Wang, Xiaofeng Wang, Feng Qian, Raheem Beyah, and Damon McCoy. 2017. Under the shadow of sunshine: Understanding and detecting bulletproof hosting on legitimate service provider networks. In *2017 IEEE Symposium on Security and Privacy (SP)*. IEEE, San Jose, CA, US, 805–823.
- [11] Sumayah Alrwais, Xiaojing Liao, Xianghang Mi, Peng Wang, Xiaofeng Wang, Feng Qian, Raheem Beyah, and Damon McCoy. 2017. Under the Shadow of Sunshine: Understanding and Detecting Bulletproof Hosting on Legitimate Service Provider Networks. In *2017 IEEE Symposium on Security and Privacy (SP)*. IEEE, San Jose, CA, USA, 805–823.
- [12] Chaitrali Amrutkar, Young Seuk Kim, and Patrick Traynor. 2017. Detecting Mobile Malicious Webpages in Real Time. *IEEE Transactions on Mobile Computing* 16, 8 (2017), 2184–2197.
- [13] Icy Fresno Anabo, Iciar Elexpuru-Albizuri, and Lourdes Villardón-Gallego. 2019. Revisiting the Belmont Report's ethical principles in internet-mediated research: Perspectives from disciplinary associations in the social sciences. *Ethics and Information Technology* 21, 2 (2019), 137–149.
- [14] Ross Anderson, Chris Barton, Rainer Böhm, Richard Clayton, Carlos Gañán, Tom Grasso, Michael Levi, Tyler Moore, and Marie Vasek. 2019. Measuring the changing cost of cybercrime. In *In 18th Workshop on the Economics of Information Security (WEIS)*. WEIS, Boston, MA, USA.
- [15] Marco Balduzzi, Payas Gupta, Lion Gu, Debin Gao, and Mustaqe Ahamad. 2016. MobiPot: Understanding Mobile Telephony Threats with Honeycards. In *Proceedings of the 11th ACM on Asia Conference on Computer and Communications Security (ASIA CCS '16)*. Association for Computing Machinery, New York, NY, USA, 723–734.
- [16] Artur Barseghyan. 2013. tld 0.13. <https://pypi.org/project/tld>.
- [17] Bradley Barth. 2016. Snack attack: A crimeware-as-a-service menu for wannabe hackers. <https://www.scworld.com/news/snack-attack-a-crimeware-as-a-service-menu-for-wannabe-hackers>.
- [18] Hugo Bijmans, Tim Booi, Anneke Schwedersky, Aria Nedgabat, and Rolf van Wegberg. 2021. Catching Phishers By Their Bait: Investigating the Dutch Phishing Landscape through Phishing Kit Detection. In *30th USENIX Security Symposium (USENIX Security 21)*. USENIX Association, Virtual, USA, 3757–3774.
- [19] Amanda Bodker, Phil Connolly, Oliver Sing, Benjamin Hutchins, Michael Townsley, and Jacqueline Drew. 2022. Card-not-present fraud: Using crime scripts to inform crime prevention initiatives. *Security Journal* 36 (2022), 693–711.
- [20] Broadcom. 2024. Russian bulletproof hosting services exploited for malicious activities, SocGholish malware campaigns. <https://www.broadcom.com/support/security-center/protection-bulletin/russian-bulletproof-hosting-services-exploited-for-malicious-activities-socgholish-malware-campaigns>.
- [21] Tao Chen and Min-Yen Kan. 2012. Creating a Live, Public Short Message Service Corpus: The NUS SMS Corpus. *Language Resources and Evaluation* 47 (Aug. 2012), 299–335.
- [22] Cisco Talos Intelligence Group. 2024. PhishTank. <https://www.phishtank.com>.
- [23] City of London Police. 2024. Two people arrested in connection with investigation into homemade mobile antenna used to send thousands of smishing text messages to the public. <https://www.cityoflondon.police.uk/news/city-of-london/news/2024/june/two-people-arrested-in-connection-with-investigation-into-homemade-mobile-antenna-used-to-send-thousands-of-smishing-text-messages-to-the-public/>.
- [24] Jacob Cohen. 1960. A coefficient of agreement for nominal scales. *Educational and psychological measurement* 20, 1 (1960), 37–46.
- [25] Kevin Collier. 2021. Odd text from a wrong number? It's probably a scam. <https://www.nbcnews.com/tech/security/wrong-number-text-scam-rcna39793>.
- [26] Casey Crane. 2025. Industry Leaders Approve the Move to a 47-Day SSL Certificate Validity Period. <https://sectigostore.com/blog/47-day-ssl-certificate-validity/>.
- [27] Leslie Daigle. 2004. WHOIS protocol specification. <https://www.rfc-editor.org/rfc/rfc3912>.
- [28] Gert Van de Ven. 2023. IMSI-catcher used in massive phishing campaign leads to arrests in France. <https://www.conquer-your-risk.com/2023/02/24/imsi-catcher-used-in-massive-phishing-campaign-leads-to-arrests-in-france/>.
- [29] Dedicated Card and Payment Crime Unit (DCPCU). 2023. This week, with support from partners, the DCPCU executed four search warrants across England. <https://bit.ly/3NuqjwD>.
- [30] Sarah Jane Delany, Mark Buckley, and Derek Greene. 2012. SMS spam filtering: Methods and data. *Expert Systems with Applications* 39, 10 (2012), 9899–9908.
- [31] Fabio Duarte. 2025. Reddit User Age, Gender, & Demographics (2025). <https://explodingtopics.com/blog/reddit-users>.
- [32] Gerry Eaton. 2024. Bulletproof Hosting Havens for FIN7 and Russian Cyber Threat Groups. <https://threatshare.ai/research/bulletproof-hosting-havens-for-fin7-and-russian-cyber-threat-groups/>.
- [33] Sec Edu. 2021. SMiShing dataset. <https://in.pinterest.com/seceduau/smishing-dataset/>.
- [34] Ericsson Mobility. 2025. Ericsson Mobility Visualizer. <https://www.ericsson.com/en/reports-and-papers/mobility-report/mobility-visualizer?f=1&ft=3&r=1&t=8&s=1&u=1&y=2014,2025&c=1>.
- [35] Ethnologue. 2025. What are the top 200 most spoken languages? <https://www.ethnologue.com/insights/ethnologue200/>.
- [36] European Union. 2006. DIRECTIVE 2006/24/EC OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL. <https://eur-lex.europa.eu/eli/dir/2006/24/oj>.
- [37] EUROPOL. 2023. Action against criminal website that offered 'spoofing' services to fraudsters: 142 arrests. <https://www.europol.europa.eu/media-press/newsroom/news/action-against-criminal-website-offered-ãÿspoofingãÿ-services-to-fraudsters-142-arrests>.
- [38] Harry Everett. 2024. Out of the shadows – 'darcula' iMessage and RCS smishing attacks target USPS and global postal services. <https://www.netcraft.com/blog/darcula-smishing-attacks-target-usps-and-global-postal-services/>.
- [39] FBI. 2025. https://www.ic3.gov/AnnualReport/Reports/2024_IC3Report.pdf.
- [40] Álvaro Feal, Pelayo Vallina, Julien Gamba, Sergio Pastrana, Antonio Nappa, Oliver Hohlfeld, Narseo Vallina-Rodriguez, and Juan Tapiador. 2021. Blocklist babel: On the transparency and dynamics of open source blocklisting. *IEEE Transactions on Network and Service Management* 18, 2 (2021), 1334–1349.
- [41] Federal Network Agency. 2025. Monitoring measures and information. <https://www.bundesnetzagentur.de/863984>.
- [42] Federal Trade Commission (FTC). 2022. How to Recognize and Report Spam Text Messages. <https://consumer.ftc.gov/articles/how-recognize-and-report-spam-text-messages>.
- [43] Emma Fletcher. 2023. IYKYK: The top text scams of 2022. <https://www.ftc.gov/news-events/data-visualizations/data-spotlight/2023/06/iyykyk-top-text-scams-2022>.
- [44] Artur Geers, Aaron Ding, Carlos Hernandez Gañán, and Simon Parkin. 2023. Lessons in Prevention and Cure: A User Study of Recovery from Flubot Smartphone Malware. In *Proceedings of the 2023 European Symposium on Usable Security (EuroUSEC '23)*. Association for Computing Machinery, New York, NY, USA, 126–142.
- [45] Google. 2023. Report spam in Google Messages. <https://support.google.com/messages/answer/9061432>.
- [46] Google. 2024. Google Safe Browsing - Google Transparency Report. <https://transparencyreport.google.com/safe-browsing/search>.
- [47] Google. 2024. Safe Browsing APIs (v4). <https://developers.google.com/safe-browsing/v4/>.
- [48] Google. 2025. New AI-Powered Scam Detection Features to Help Protect You on Android. <https://security.googleblog.com/2025/03/new-ai-powered-scam-detection-features.html>.
- [49] Chris Grier, Kurt Thomas, Vern Paxson, and Michael Zhang. 2010. @spam: the underground on 140 characters or less. In *Proceedings of the 17th ACM Conference on Computer and Communications Security (CCS '10)*. Association for Computing Machinery, New York, NY, USA, 27–37.
- [50] Anti-Phishing Working Group. 2004. The APWG eCrime Exchange (eCX). <https://apwg.org/ecx/>.
- [51] Neha Gupta, Anupama Aggarwal, and Ponnurangam Kumaraguru. 2014. bit.ly/malicious: Deep dive into short URL based e-crime detection. In *2014 APWG Symposium on Electronic Crime Research (eCrime)*. IEEE, Birmingham, AL, US, 14–24.
- [52] Samuel Hoffstaetter. 2014. pyteseract 0.3.13. <https://pypi.org/project/pyteseract/>.
- [53] Médéric Hurier, Guillermo Suarez-Tangil, Santanu Kumar Dash, Tegawendé F Bissyandé, Yves Le Traon, Jacques Klein, and Lorenzo Cavallaro. 2017. Euphony: Harmonious unification of cacophonous anti-virus vendor labels for android malware. In *2017 IEEE/ACM 14th International Conference on Mining Software Repositories (MSR)*. IEEE, Buenos Aires, Argentina, 425–435.
- [54] Kiarah Reyshylle Ibañez. 2024. Scammer Strategies and Social Actions in Online Filipino Transactions. *Southeastern Philippines Journal of Research and Development* 29, 1 (2024), 43–75.
- [55] Internet Assigned Numbers Authority (IANA). 2024. Root Zone Database. <https://www.iana.org/domains/root/db>.
- [56] Interpol. 2025. More than 300 arrests as African countries clamp down on cyber threats. <https://www.interpol.int/en/News-and-Events/News/2025/More-than-300-arrests-as-African-countries-clamp-down-on-cyber-threats>.
- [57] IPinfo. 2024. The Trusted Source For IP Address Data. <https://ipinfo.io/>.
- [58] Ankit Kumar Jain and B. B. Gupta. 2018. Rule-Based Framework for Detection of Smishing Messages in Mobile Environment. *Procedia Computer Science* 125 (jan 2018), 617–623.
- [59] Ankit Kumar Jain and Brij B Gupta. 2019. Feature based approach for detection of smishing messages in the mobile environment. *Journal of Information Technology Research (JITR)* 12, 2 (2019), 17–35.
- [60] Wenxiang Jiao, Wenxuan Wang, Jen tse Huang, Xing Wang, Shuming Shi, and Zhaopeng Tu. 2023. Is ChatGPT A Good Translator? Yes With GPT-4 As The

- Engine. <https://arxiv.org/abs/2301.08745>. arXiv:2301.08745 [cs.CL]
- [61] Jae Woong Joo, Seo Yeon Moon, Saurabh Singh, and Jong Hyuk Park. 2017. S-Detector: an enhanced security model for detecting Smishing attack for mobile computing. *Telecommunication Systems* 66, 1 (sep 2017), 29–38.
 - [62] Doowon Kim, Haehyun Cho, Yonghwi Kwon, Adam Doupe, Soel Son, Gail-Joon Ahn, and Tudor Dumitras. 2021. Security Analysis on Practices of Certificate Authorities in the HTTPS Phishing Ecosystem. In *Proceedings of the 2021 ACM Asia Conference on Computer and Communications Security (ASIA CCS '21)*. Association for Computing Machinery, New York, NY, USA, 407–420.
 - [63] Maria Konte, Roberto Perdisci, and Nick Feamster. 2015. Aswatch: An as reputation system to expose bulletproof hosting ases. In *Proceedings of the 2015 ACM Conference on Special Interest Group on Data Communication*. Association for Computing Machinery, New York, NY, USA, 625–638.
 - [64] Sophie Le Page, Guy-Vincent Jourdan, Gregor V. Bochmann, Jason Flood, and Isif-Violel Onut. 2018. Using URL shorteners to compare phishing and malware attacks. In *2018 APWG Symposium on Electronic Crime Research (eCrime)*. IEEE, San Diego, CA, US, 1–13.
 - [65] Kevin Lee and Arvind Narayanan. 2021. Security and Privacy Risks of Number Recycling at Mobile Carriers in the United States. In *2021 APWG Symposium on Electronic Crime Research (eCrime)*. IEEE, Boston, MA, USA, 1–17.
 - [66] Mingxuan Liu, Yiming Zhang, Baojun Liu, Zhou Li, Haixin Duan, and Donghong Sun. 2021. Detecting and Characterizing SMS Spearphishing Attacks. In *Annual Computer Security Applications Conference*. Association for Computing Machinery, New York, NY, USA, 930–943.
 - [67] Larry Loeb. 2019. Servers Discovered With Multiple Malware Families, Staged & Ready to Launch. <https://www.darkreading.com/application-security/servers-discovered-with-multiple-malware-families-staged-ready-to-launch>.
 - [68] Dhia Mahjoub. 2017. Behaviors and patterns of bulletproof and anonymous hosting providers. In *Usenix Enigma Conference*. USENIX Association, OAKLAND, CA, USA.
 - [69] Mavenir. 2025. SpamShield Messagin Fraud. <https://www.mavenir.com/portfolio/mavapps/fraud-security/spamshield-messaging-fraud/>.
 - [70] Tomas Meskauskas. 2023. How to remove SMSSpy malware from your Android device. <https://www.pcrisk.com/removal-guides/23541-smsspy-malware-android>.
 - [71] Mara Miller. 2025. What an SMS Aggregator Is and How to Choose One. <https://www.vibes.com/blog/what-is-an-sms-aggregator>.
 - [72] Ministry for Digital Transformation and Public Service. 2025. BOE-A-2025-2870 Orden TDF/149/2025, de 12 de febrero, por la que se establecen medidas para combatir las estafas de suplantación de identidad a través de llamadas telefónicas y mensajes de texto fraudulentos y para garantizar la identificación de la numeración utilizada para la prestación de servicios de atención al cliente y realización de llamadas comerciales no solicitadas. <https://www.boe.es/buscar/act.php?id=BOE-A-2025-2870>.
 - [73] Vickie A Miracle. 2016. The Belmont Report: The triple crown of research ethics. *Dimensions of critical care nursing* 35, 4 (2016), 223–228.
 - [74] Sandhya Mishra and Devpriya Soni. 2019. SMS Phishing and Mitigation Approaches. In *2019 Twelfth International Conference on Contemporary Computing (IC3)*. IEEE, Noida, India, 1–5.
 - [75] Sandhya Mishra and Devpriya Soni. 2020. Smishing Detector: A security model to detect smishing through SMS content analysis and URL behavior analysis. *Future Generation Computer Systems* 108 (jul 2020), 803–815.
 - [76] Sandhya Mishra and Devpriya Soni. 2021. DSmsishSMS-A System to Detect Smishing SMS. *Neural Computing and Applications* 35 (jul 2021), 4975–4992.
 - [77] Mobile Ecosystem Forum (MEF). 2020. SMS SenderID Protection Registry. <https://mobileecosystemforum.com/sms-senderid-protection-registry/>.
 - [78] José Miguel Moreno, Srdjan Matic, Narseo Vallina-Rodriguez, and Juan Tapiador. 2023. Your Code is 0000: An Analysis of the Disposable Phone Numbers Ecosystem. In *2023 7th Network Traffic Measurement and Analysis Conference (TMA)*. IEEE, Naples, Italy, 1–10.
 - [79] David Morris. 2021. What is a HLR Lookup? <https://www.hlrlookup.com/what-is-a-hlr-lookup/>.
 - [80] Giovane C. M. Moura, Thomas Daniels, Maarten Bosteels, Sebastian Castro, Moritz Müller, Thymen Wabeke, Thijs van den Hout, Maciej Korczyński, and Georgios Smaragdakis. 2024. Characterizing and Mitigating Phishing Attacks at cTLD Scale. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security (CCS '24)*. Association for Computing Machinery, New York, NY, USA, 2147–2161.
 - [81] Aleksandr Nahapetyan, Sathvik Prasad, Kevin Childs, Adam Oest, Yeganeh Ladwig, Alexandros Kapravelos, and Bradley Reaves. 2024. On sms phishing tactics and infrastructure. In *2024 IEEE Symposium on Security and Privacy (SP)*. IEEE, San Francisco, CA, USA, 169–169.
 - [82] Ofcom. 2023. 45 million people targeted by scam calls and texts this summer. <https://www.ofcom.org.uk/phones-and-broadband/scam-calls-and-messages/45-million-people-targeted-by-scams>.
 - [83] Ofcom. 2024. Experiences of suspicious calls, texts and app messages. <https://www.ofcom.org.uk/siteassets/resources/documents/research-and-data/telecoms-research/scams-research/experiences-of-suspicious-calls-texts-and-app-messages-research-2024.pdf>.
 - [84] Ofcom. 2024. How to report scam texts and mobile calls to 7726. <https://www.ofcom.org.uk/phones-and-broadband/scam-calls-and-messages/7726-reporting-scam-texts-and-calls/>.
 - [85] Ofcom. 2024. Reducing mobile messaging scams. <https://www.ofcom.org.uk/siteassets/resources/documents/consultations/category-1-10-weeks/call-for-input-reducing-mobile-messaging-scams/main-documents/cfi-reducing-mobile-messaging-scams.pdf>.
 - [86] OpenPhish. 2024. OpenPhish - Phishing Intelligence. <https://openphish.com/index.html>.
 - [87] Antonis Papasavva, Samantha Lundrigan, Ed Lowther, Shane Johnson, Enrico Mariconti, Anna Markovska, and Nilufer Tuptuk. 2025. Applications of AI-Based Models for Online Fraud Detection and Analysis. *Crime Science* 14, 1 (2025), 7.
 - [88] Craig Partridge and Mark Allman. 2016. Ethical Considerations in Network Measurement Papers. *Commun. ACM* 59, 10 (sep 2016), 58–64.
 - [89] A. Paul. 2021. crtsh 0.3.1. <https://pypi.org/project/crtsh/>.
 - [90] Chris Pemberton. 2016. Tap Into The Marketing Power of SMS. <https://www.gartner.com/en/marketing/insights/articles/tap-into-the-marketing-power-of-sms>.
 - [91] Peng Peng, Limin Yang, Linhai Song, and Gang Wang. 2019. Opening the Blackbox of VirusTotal: Analyzing Online Phishing Scan Engines. In *Proceedings of the Internet Measurement Conference (IMC '19)*. Association for Computing Machinery, New York, NY, USA, 478–485.
 - [92] Fahad Pervaz, Rai Shah Nawaz, Muhammad Umer Ramzan, Maryem Zafar Usmani, Shirang Mare, Kurtis Heimerl, Faisal Kamiran, Richard Anderson, and Lubna Razaq. 2019. An assessment of SMS fraud in Pakistan. In *Proceedings of the 2nd ACM SIGCAS Conference on Computing and Sustainable Societies*. Association for Computing Machinery, New York, NY, USA, 195–205.
 - [93] Proofpoint. 2022. Cybersecurity: The 2022 Board Perspective. <https://www.proofpoint.com/sites/default/files/white-papers/pfpt-us-wp-board-perspective-report.pdf>.
 - [94] Florian Quinkert, Martin Degeling, Jim Blythe, and Thorsten Holz. 2020. Be the Phisher – Understanding Users' Perception of Malicious Domains. In *Proceedings of the 15th ACM Asia Conference on Computer and Communications Security (ASIA CCS '20)*. Association for Computing Machinery, New York, NY, USA, 263–276.
 - [95] Bradley Reaves, Nolen Scaife, Dave Tian, Logan Blue, Patrick Traynor, and Kevin R. B. Butler. 2016. Sending Out an SMS: Characterizing the Security of the SMS Ecosystem with Public Gateways. In *2016 IEEE Symposium on Security and Privacy (SP)*. IEEE, San Jose, CA, USA, 339–356.
 - [96] Bradley Reaves, Luis Vargas, Nolen Scaife, Dave Tian, Logan Blue, Patrick Traynor, and Kevin R. B. Butler. 2018. Characterizing the Security of the SMS Ecosystem with Public Gateways. *ACM Trans. Priv. Secur.* 22, 1 (dec 2018), 31 pages.
 - [97] Robokiller. 2023. Wrong-number text scams exposed: What you need to know. <https://www.robokiller.com/blog/wrong-number-text-scams>.
 - [98] Sayak Saha Roy, Elham Pourabbas Vafa, Kobra Khanmohammadi, and Shirin Nilizadeh. 2025. DarkGram: A Large-Scale Analysis of Cybercriminal Activity Channels on Telegram. In *34th USENIX Security Symposium (USENIX Security 25)*. USENIX Association, Seattle, WA, USA.
 - [99] RTE. 2024. MWC hears SMS fraud a headache for telecom operators. <https://www.rte.ie/news/business/2024/02/28/1434908-mwc-hears-sms-fraud-a-headache-for-telecom-operators/>.
 - [100] Ryu Saeki, Leo Kitayama, Jun Koga, Makoto Shimizu, and Kazumasa Oida. 2022. Smishing Strategy Dynamics and Evolving Botnet Activities in Japan. *IEEE Access* 10 (2022), 114869–114884.
 - [101] Sayak Saha Roy, Unique Karanjit, and Shirin Nilizadeh. 2023. Phishing in the Free Waters: A Study of Phishing Attacks Created using Free Website Building Services. In *Proceedings of the 2023 ACM on Internet Measurement Conference (IMC '23)*. Association for Computing Machinery, New York, NY, USA, 268–281.
 - [102] Shagufta Khurram Siddiqui, Farhana Yasmeen Qadri, and Zulfiqar Ali Chachar. 2023. Developing proficiency in english language and embarking on the journey of global entrepreneurship: a pathway to achieving success. *Journal of Entrepreneurship, Management, and Innovation* 5, 5 (2023), 821–837.
 - [103] Camelia Simoiu, Ali Zand, Kurt Thomas, and Elie Bursztin. 2020. Who is targeted by email-based phishing and malware? Measuring factors that differentiate risk. In *Proceedings of the ACM Internet Measurement Conference (IMC '20)*. Association for Computing Machinery, New York, NY, USA, 567–576.
 - [104] Gilberto Atondo Siu, Alice Hutchings, Marie Vasek, and Tyler Moore. 2022. "Invest in crypto!": An analysis of investment scam advertisements found in Bitcointalk. In *2022 APWG Symposium on Electronic Crime Research (eCrime)*. IEEE, Boston, MA, USA, 1–12.
 - [105] Spamhaus. 2024. Passive DNS. <https://www.spamhaus.com/product/passive-dns/>.
 - [106] Bharat Srinivasan, Payas Gupta, Manos Antonakakis, and Mustaque Ahamad. 2016. Understanding cross-channel abuse with sms-spam support infrastructure attribution. In *Computer Security—ESORICS 2016: 21st European Symposium on*

- Research in Computer Security*. Springer International Publishing, Cham, 3–26.
- [107] Frank Stajano and Paul Wilson. 2011. Understanding scam victims: seven principles for systems security. *Commun. ACM* 54, 3 (2011), 70–75.
 - [108] Statista. 2024. Leading countries based on number of X (formerly Twitter) users as of April 2024. <https://www.statista.com/statistics/242606/number-of-active-twitter-users-in-selected-countries/>.
 - [109] Caroline Tagg. 2009. *A corpus linguistics study of SMS text messaging*. Ph. D. Dissertation. University of Birmingham.
 - [110] Siyuan Tang, Xianghang Mi, Ying Li, XiaoFeng Wang, and Kai Chen. 2022. Clues in Tweets: Twitter-Guided Discovery and Analysis of SMS Spam. In *Proceedings of the 2022 ACM SIGSAC Conference on Computer and Communications Security (CCS '22)*. Association for Computing Machinery, New York, NY, USA, 2751–2764.
 - [111] A. Tiago, J. Hidalgo, and Y. Akebo. 2011. Contributions to the Study of SMS Spam Filtering: New Collection and Results. In *Proceedings of the 11th ACM Symposium on Document Engineering*. Association for Computing Machinery, New York, NY, USA, 259–262.
 - [112] Ke Tian, Steve T. K. Jan, Hang Hu, Danfeng Yao, and Gang Wang. 2018. Needle in a Haystack: Tracking Down Elite Phishing Domains in the Wild. In *Proceedings of the Internet Measurement Conference 2018 (IMC '18)*. Association for Computing Machinery, New York, NY, USA, 429–442.
 - [113] Daniel Timko and Muhammad Lutfur Rahman. 2024. Smishing Dataset I: Phishing SMS Dataset from Smishtank. com. In *Proceedings of the Fourteenth ACM Conference on Data and Application Security and Privacy*. Association for Computing Machinery, New York, NY, USA, 289–294.
 - [114] Bill Toulas. 2022. Revolut hack exposes data of 50,000 users, fuels new phishing wave. <https://www.bleepingcomputer.com/news/security/revolut-hack-exposes-data-of-50-000-users-fuels-new-phishing-wave/>.
 - [115] Trend Micro. 2022. Ukraine Charitable Donation Scams Are Misusing the Name of a Legitimate Charity. <https://news.trendmicro.com/2022/04/12/ukraine-charity-scammers-impersonating-legitimate-charity/>.
 - [116] UK Home Office. 2025. Major step for fraud prevention with landmark ban on SIM farms. <https://www.gov.uk/government/news/major-step-for-fraud-prevention-with-landmark-ban-on-sim-farms>.
 - [117] VirusTotal. 2024. VirusTotal. <https://docs.virustotal.com/docs/how-it-works>.
 - [118] VirusTotal. 2024. What is the difference between the public API and the private API? <https://docs.virustotal.com/docs/difference-public-private>.
 - [119] Jingjing Wang, Liu Wang, Feng Dong, and Haoyu Wang. 2023. Re-measuring the Label Dynamics of Online Anti-Malware Engines from Millions of Samples. In *Proceedings of the 2023 ACM on Internet Measurement Conference (IMC '23)*. Association for Computing Machinery, New York, NY, USA, 253–267.
 - [120] Liu Wang, Ren He, Haoyu Wang, Pengcheng Xia, Yuanchun Li, Lei Wu, Yajin Zhou, Xiapu Luo, Yulei Sui, Yao Guo, et al. 2021. Beyond the virus: a first look at coronavirus-themed Android malware. *Empirical Software Engineering* 26, 4 (2021), 82.
 - [121] Liu Wang, Haoyu Wang, Ren He, Ran Tao, Guozhu Meng, Xiapu Luo, and Xuanzhe Liu. 2022. MalRadar: Demystifying Android Malware in the New Era. *Proc. ACM Meas. Anal. Comput. Syst.* 6, 2, Article 40 (June 2022), 27 pages.
 - [122] WhoisXMLAPI. 2024. WHOIS API offers unified & consistent data. <https://whois.whoisxmlapi.com>.
 - [123] Longfei Wu, Xiaojian Du, and Jie Wu. 2014. MobiFish: A lightweight anti-phishing scheme for mobile phones. In *2014 23rd International Conference on Computer Communication and Networks (ICCCN)*. IEEE, Shanghai, China, 1–8.
 - [124] Angelin Yeoh. 2022. Maybank warns of malicious SMSSpy campaign targeting Android users in Malaysia. <https://www.thestar.com.my/tech/tech-news/2022/06/02/maybank-warns-of-malicious-smsspy-campaign-targeting-android-users-in-malaysia>.
 - [125] Zimperium. 2025. The Evolution of Mobile-Specific Phishing Attacks. <https://lp.zimperium.com/hubfs/MTD/REPORT/GENERIC/SmishingThreatReport26Q1.pdf>.

A APPENDIX: ETHICS

Our research has some ethical concerns. We collect image attachments of smishing messages reported by users on online forums. These attachments might include metadata containing personally identifiable information (PII). We do not manually open the images and delete them immediately after extracting the text programmatically. Direct consent is not possible in the case of internet measurement research [88]. Instead, we can view our research using the beneficence principle and make a risk-benefit assessment. Following the Belmont report [13, 73], we determine that the risks

to any stakeholder are negligible, and our work has broader societal benefits.

We extract the sender IDs, including the mobile numbers abused by scammers to send smishing messages for HLR lookup. We also perform measurements on the extracted URLs, such as antivirus detection. However, as mobile numbers are recycled [65] and parameters or path of the URLs may contain PII, we do not share the raw mobile numbers or complete URLs in our published dataset. Additionally, our research provides an understanding of smishing that helps suggest countermeasures and allow stakeholders to tackle this cybercrime. We perform data protection impact assessments to minimize risks. After a thorough review, the university’s research ethics committee approved our study.

B APPENDIX: ARTIFACT AVAILABILITY

We have released our labeled dataset and the code used to generate the plots presented in the paper towards open science, available at <https://github.com/reportsmishing/Smishing-Dataset-IMC25>. This helps ensure the availability and reproducibility of our work. A detailed description of our dataset is provided in Appendix C, and the prompts employed in our research have already been made available in Appendix D.

C APPENDIX: DATASET DESCRIPTION

Our paper provides an updated novel smishing dataset. This contains the following fields:

- **Sender ID:** In line with the ethical principles, we can not share the actual mobile number, email address, or the alphanumeric sender ID as it falls under personally identifiable information (PII). Therefore, we provide anonymized sender ID, i.e., ‘phone number,’ ‘email,’ or ‘alphanumeric.’
- **Sender ID Type:** The phone number type returned from the HLR lookup, where the sender ID is a phone number.
- **Sender ID Original Mobile Network Operator:** The original mobile network operator returned from the HLR lookup, where the sender ID is a validated phone number.
- **Sender ID’s Origin Country:** The origin country returned from the HLR lookup, where the sender ID is a validated phone number.
- **Text Message:** The SMS text received by the users after removing the PII information, i.e., names, URLs, and phone numbers, if any.
- **Translated Text Message:** The translation of the SMS text in English, where the text is in a different language.
- **URL Shortener:** The name of the URL shortener that is abused in the smishing text.
- **Brand Impersonated:** Extracted name of the entity/brand from the smishing text that scammers impersonate to lure a victim.
- **Scam Category:** Identified scam type based on known scam categories [6].
- **Lure Principles:** Lures scammers use in smishing texts to deceive the victim into taking an action.
- **Language:** Language of the original smishing text.

D APPENDIX: OPENAI API PROMPTS

D.1 OpenAI Vision API prompt

You will receive a json object with an 'image'. The 'image' is reported by a user as a phishing SMS. This should most likely be a screenshot of the text message received on a user's mobile phone. Based on the instructions below, process the message and return a json object. Instructions: Do not extract the details if it is not a screenshot of the SMS message and return the below parameters empty. If it is a mobile message screenshot, you need to extract the following and return a JSON response consisting of the following: 'timestamp': This should be the date and time in the screenshot when the SMS message was received. If the timestamp is not there, leave it empty. 'text': This should be the text in the SMS message. If unavailable in the screenshot, leave it empty. 'url': If the SMS contains a URL, extract it; otherwise, leave it empty. 'sender-id': This should be the sender ID (mobile number, alphanumeric sender ID, or email address) that sent the SMS message. If it is not available, leave it empty.

D.2 OpenAI API annotation prompt

You will receive a json object with an 'id' and a 'message'. The 'id' is the id of the message and the 'message' is text extracted from a phishing SMS screenshot. Based on the instructions below, process the message and return a json object. Instructions: 1. Translate the text to English, ONLY if it is not in English ("translation" key in the json object. This key should ONLY exist in the JSON if the text is not in the English language). 2. Identify the brand, organization, or any other named entity that the message is trying to impersonate in the text. If unspecified or unknown, then leave it empty. ("named_entity" key in the json object. This key should always be returned in the json.). 3. Classify the type of smishing message ("scam_type" key in the json. This key should always be returned in the json.) The scam_types can be: a) Hey mum/dad - text addressed to mom/mum or dad and asking to text back potentially giving a reason about phone being broken or using a different mobile number. b) Delivery/Parcel - text impersonating a parcel/delivery company asking to click on a link, text back or call on a number. c) Banking - text impersonating a bank or a financial institution asking to click on a link, text back or call on a number. d) Government - text impersonating a government organization asking to click on a link, text back or call on a number. e) Telecom - text impersonating a mobile network operator asking to click on a link, text back or call on a number. f) Wrong number - text addressed to an individual that looks like a normal greeting or asking about someone and/or to reply back. g) Spam - illicit marketing message including casino, betting, random draws, etc. h) Others - If it does not fit as of the above category. 4. Provide which lure principles apply for each text message ("lure_principles" key should be a list and always be provided in the json object. If you cannot detect any lure principles, leave the list empty.) Lure principles are: a) Distraction Principle - providing various reasons to distract the user. b) Authority Principle - providing trust to the user to not question authority. could be done by making references to legitimate entities. c) Herd Principle - encouraging a user to not miss out on opportunities by relating to the popularity of a scheme. convincing how others have won things or take the same risk. d) Dishonesty Principle - inviting users

willingly and knowingly participating into a fraudulent scheme. e) Kindness Principle - Fraudsters leverage the willingness of people to help others. for eg. hi mum/dad texts or cases where someone asks for help. f) Need and Greed Principle - leveraging users' greed and offering attractive (monetary) benefits so user would take an action asked in the text. g) Time/Urgency Principle - putting time pressure on users so they make a rushed decision. 5. Every json object should include the "id" of the message being classified. 6. Return the language code of the text ("language" key in the json object. This key should always be returned in the json.)

E APPENDIX: SMISHING TEXTS

Table 15: Annual distribution of tweets reporting smishing texts and their image attachments we collect from Twitter.

Year	Tweets	Image Attachments
01/2017 - 12/2017	6,345 (2.9%)	1,747 (2.9%)
01/2018 - 12/2018	9,957 (4.6%)	2,717 (4.5%)
01/2019 - 12/2019	16,403 (7.6%)	6,537 (10.9%)
01/2020 - 12/2020	34,265 (15.9%)	8,750 (14.5%)
01/2021 - 12/2021	45,486 (21.1%)	11,717 (19.5%)
01/2022 - 11/2022	51,690 (23.9%)	10,315 (17.1%)
12/2022 - 06/2023	51,696 (23.9%)	18,426 (30.6%)
Total	215,842	60,209

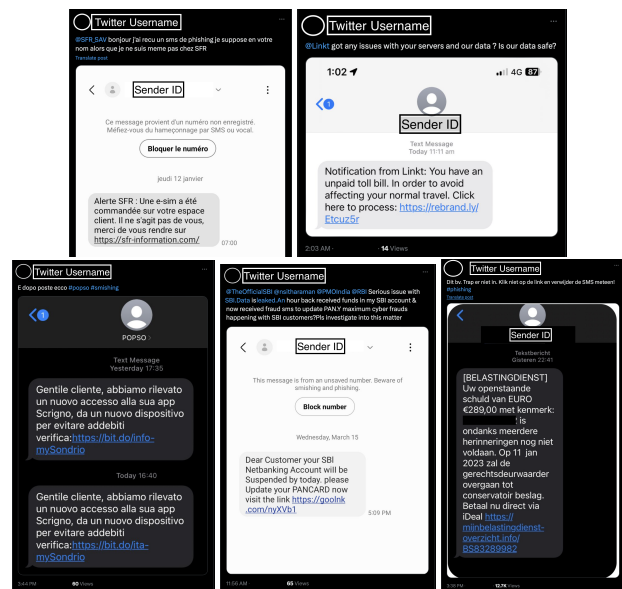


Figure 4: Examples of users reporting smishing texts globally on Twitter.

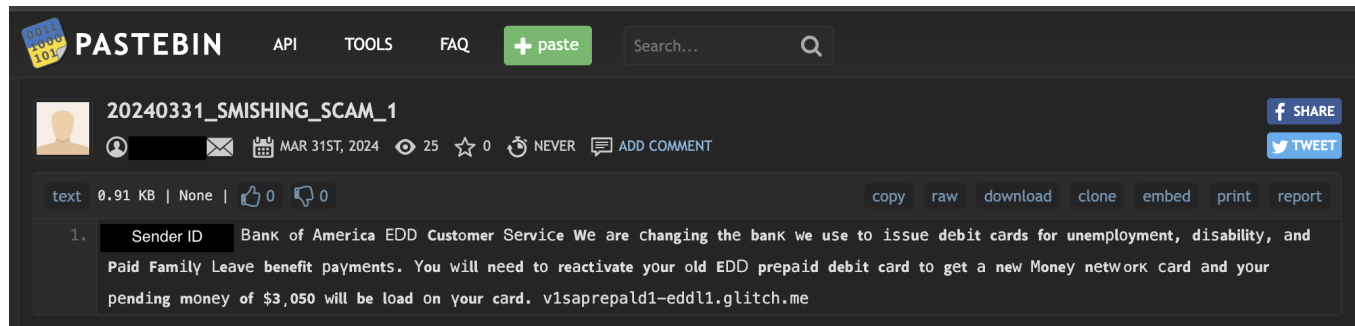


Figure 5: Example of a paste containing a smishing text on Pastebin.

F APPENDIX: SMISHING INFRASTRUCTURE

Table 16: Distribution of unique smishing URLs' TLDs based on IANA classification.

Type	URLs	TLDs
Generic (gTLD)	7,539 (72.3%)	146
Country-Code (ccTLD)	2,829 (27.1%)	130
Generic-restricted (grTLD)	30 (0.3%)	3
Sponsored (sTLD)	25 (0.2%)	8
Infra (iTLD)	-	-

Table 17: Top 10 registrars scammers abuse to register smishing domains.

Registrars	Domains
GoDaddy	464
NameCheap	153
Gname	98
Dynadot	79
Tucows	74
PublicDomainRegistry	71
NameSilo	64
Key-Systems	60
MarkMonitor	53
Gandi	52

G APPENDIX: SCAMMER STRATEGIES

H APPENDIX: BULK SMS SERVICES

Figure 6: Post from an underground forum where a threat actor is offering bulk SMS services with custom sender IDs.

Table 18: Google Safe Browsing (GSB) antivirus detection results for smishing URLs ($n = 19,864$) through GSB's API, transparency report website, and on VirusTotal.

Google Safe Browsing	Unsafe URLs	Partially Unsafe URLs	URLs Undetected	No Available Data	Not Queried (§3.3.4)
API	191 (1.0%)	-	19,673 (99.0%)	-	-
Transparency Report	802 (4.0%)	440 (2.2%)	5,883 (29.6%)	2,823 (14.2%)	9,948 (50.1%)
on VirusTotal	319 (1.6%)	-	19,545 (98.4%)	-	-

Table 19: Case study: A distribution of 18 APK malware actively identified from smishing messages ($n = 200$).

Indicator of Compromise (IoC)	Malware Family
5dceeb810142f65e692cddbe6fd1b1123f0f606575b6d7c6d666e0e65f62de2f	SMSSpy
1ef6913e78da66294e8738b414c6ff06b59f7c9fd808af4e54586833e4019341	SMSSpy
99422143d1c7c82af73f8fd8bf5a0ce4ff32f899014241be5616a804d2104ebf	HQWar
b9481cdb24105cba4b8f4c067798ea8deed8715e0c57f3570f860afaa23e8027	SMSSpy
c66c801ab7b4373bb0c461c763b22b43c96fa9cea5f5ead8abbc99bd73d19c10	SMSSpy
34ae95c0a19e3c72f199c812f64dc8f38bbc7f0f5746efe0bd756737163ed8ec	SMSSpy
94b7d6c376871d154ade8518c4770bfc86571f58e212a632c1703fd806e1ee5c	SMSSpy
c79b0aabc24bb2169e62e17cf36a476bf9629797dbd1822dec515b9f916b4be0	SMSSpy
512ba356c6a7e79435e1178b61289b506fdc3432e3ede91b8a6fba1e4e41f89f	SMSSpy
28e826bd811c250a5bd10f0f07975a6f61ee5d79f8f5fc352bcd50b8318b2f34	Rewardsteal
aa785b6d68ff7760e192381755f764b31e19ad64a788afe6c86e30be4a7e9cd2	SMSSpy
ea6d6efec35d09e63edbc790bc7cfef8acf6bef2eeb5eaddf919083f87cf9ab	SMSSpy
a5e9b5a296bc557d747d43de4f2c86c090cf44ec202af739780f28b0c72dc470	SMSSpy
c9d5c28cefdd3c2234c844b334ebe2871ea6f9b9b55f1cccd5678dcaec883ada	SMSSpy
c7f5152aa924a03883f8f6a17dae79955216f42f3d9fbf9113df80981b8da030	Artemis
91231094ebfec2833f7f606baa2987322a36000edc2f51f4dce2b860b0b1b3d2	SMSSpy
77c5ef34f044f53c4bbb703ec6dbbaeeae61c9d16b6b56c16af1f84b0652f388	SMSSpy
ccd0375d69902236b880ece65182e5eeb393b95578e869b5b054b6d5df6dc976	SMSSpy