



Zscaler ThreatLabz 2024 Phishing Report



Discover the latest phishing trends, emerging tactics, and effective security measures to stay ahead of the ever-evolving, AI-fueled phishing threat.

Contents

03 Executive Summary

04 Key Findings

05 Top Phishing Targets

- 06 Countries that experienced the most phishing attempts
- 07 Countries of origin for phishing attacks
- 08 Industries most commonly targeted by phishing attacks
- 09 Brands most frequently imitated by threat actors
- 10 Top referring domains leading to phishing pages
- 12 Distribution of attacks across autonomous systems (ASNs)
- 13 Social media platforms exploited by threat actors

15 One Phish, Two-Faced: AI and Phishing

- 15 Phishers abuse AI, AI fights back
- 18 Rise in deepfake attacks

19 Election Spotlight: Phishing Campaigns vs. Political Campaigns

20 Evolving Phishing Trends

- 21 Recruitment scams
- 22 Adversary-in-the-middle (AiTM) attacks
- 23 Browser-in-the-browser (BiTB) attacks
- 24 QR scams
- 25 Tech support scams

26 2024–2025 Phishing Predictions

28 How the Zscaler Zero Trust Exchange Can Mitigate Phishing Attacks

- 28 Preventing compromise
- 28 Eliminating lateral movement
- 29 Shutting down compromised users and insider threats
- 29 Stopping data loss
- 29 Related Zscaler products

31 Improve Your Phishing Defenses

- 32 Best practices: Security controls
- 33 Best practices: How to spot and prevent phishing attacks
- 36 Best practices: How to identify a phishing page

38 ThreatLabz Research Methodology

39 About ThreatLabz

39 About Zscaler

Executive Summary

In the current phishing landscape, attackers have unprecedented access to a wide range of convenient tools or “easy buttons,” such as phishing-as-a-service kits, automated phishing tools, and curated target lists. These threats constantly evolve and multiply, forcing enterprises to remain in a perpetual state of heightened vigilance as they defend against the ever-changing variations of phishing scams. Complicating matters further, the emergence of artificial intelligence (AI) has significantly amplified the art of deception, enabling attackers to execute more sophisticated and elusive attacks at an unprecedented scale and speed.

AI represents a paradigm shift in the realm of cybercrime, particularly for phishing scams. With the aid of generative AI, cybercriminals can rapidly construct highly convincing phishing campaigns that surpass previous benchmarks of complexity and effectiveness. By leveraging AI algorithms, threat actors can swiftly analyze vast datasets to tailor their attacks and easily replicate legitimate communications and websites with alarming precision. This level of sophistication allows phishers to deceive even the most aware users. The potential of AI in reshaping the cyberthreat landscape appears boundless as it continues to redefine what is possible in the world of cyberattacks.

As organizations and users brace for this evolving landscape of phishing attacks, a pressing question echoes: **how can we stay ahead of these threats?**

To help answer this question with insights into the latest phishing trends, targeted entities, emerging tactics, and effective security measures, the Zscaler ThreatLabz research

team conducted an extensive analysis. Over the span of 12 months (January–December 2023), ThreatLabz examined more than 2 billion phishing transactions across the Zscaler Zero Trust Exchange™, the world’s largest online security cloud. Their findings aim to equip enterprises with the knowledge needed to proactively combat the rising wave of new phishing attacks.

The landscape of phishing attacks continues to rapidly evolve. In 2023, ThreatLabz observed a year-over-year increase of 58.2% in global phishing attempts.

This surge was characterized by emerging schemes, including voice phishing, recruitment scams, and browser-in-the-browser attacks. These findings align with data from the Anti-Phishing Working Group, an international cybercrime coalition, which declared 2023 as “the worst year for phishing on record.”¹

In light of this critical moment in the realm of phishing threats, the Zscaler ThreatLabz 2024 Phishing Report offers actionable information on phishing activity and tactics, along with best practices and strategies to enhance your organization’s security in the face of existing and evolving threats.



1. Anti-Phishing Working Group, [Phishing Activity Trends Report, 4th Quarter, 2023](#), February 13, 2024.

Key Findings



Phishing attacks surged by 58.2% in 2023, compared to 2022, reflecting the growing sophistication and persistence of threat actors.



Microsoft remains the most imitated brand, with 43.1% of phishing attempts targeting it. Microsoft's OneDrive and SharePoint brands were also among the top five targeted, indicating a persistent trend of threat actors seeking user credentials from critical Microsoft applications.



Vishing (voice phishing) and deepfake phishing attacks are on the rise as attackers leverage generative AI to amplify social engineering tactics.



Adversary-in-the-middle (AiTM) attacks remain a persistent threat, and browser-in-the-browser (BiTB) attacks are now on the rise. These tactics directly target users in web browsers, making them more challenging to detect and mitigate.



The US, UK, India, Canada, and Germany were the top five countries targeted by phishing attacks.



Tech support scams and QR CAPTCHA scams were among 2023's most prevalent attack types, exploiting users' trust in tech support services and widespread use of QR codes.



The finance and insurance industry faced 27.8% of overall phishing attacks, the highest concentration among industries and a staggering 393% year-over-year increase. Manufacturing followed closely behind at 21%.

Top Phishing Targets

ThreatLabz researchers analyzed data encompassing countries, industries, brands, and platforms to identify the primary targets of phishing attacks in 2023. The findings emphasize the persistent and widespread threat posed by phishing attacks on a global scale. Recognizing the patterns and trends in phishing activities is crucial for implementing effective cybersecurity measures to safeguard against them.

THE SECTION EXPLORES KEY ASPECTS OF PHISHING ATTACKS, INCLUDING:

- 01 Countries that experienced the most phishing attempts
- 02 Countries of origin for phishing attacks
- 03 Industries most commonly targeted by phishing attacks
- 04 Brands most frequently imitated by threat actors
- 05 Top referring domains leading to phishing pages
- 06 Distribution of attacks across autonomous system numbers (ASNs)
- 07 Social media platforms exploited by threat actors

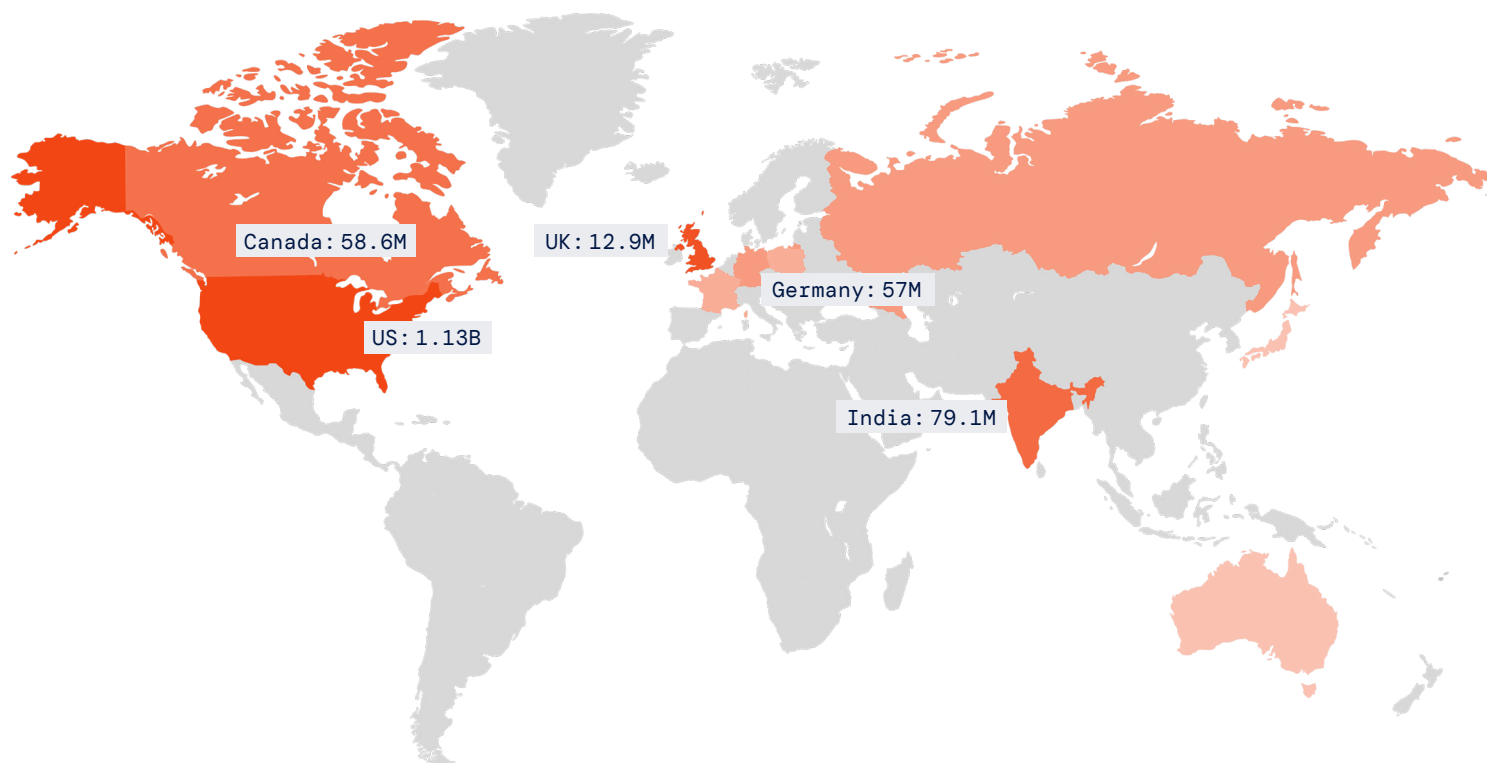


Countries that experienced the most phishing attempts

In 2023, the United States, United Kingdom, and India experienced the highest volume of phishing attempts, with the US bearing the brunt of these attacks. Factors contributing to the high occurrence of phishing in the US include its large population of internet and technology users, extensive use of online financial transactions, and advanced digital infrastructure. The prevalence of AI-driven phishing campaigns further amplifies the vulnerability of US entities to such attacks.

THE TOP 10 COUNTRIES TARGETED FOR PHISHING SCAMS WERE:

- | | |
|-------------------|--------------|
| 01 United States | 06 Russia |
| 02 United Kingdom | 07 Poland |
| 03 India | 08 France |
| 04 Canada | 09 Australia |
| 05 Germany | 10 Japan |

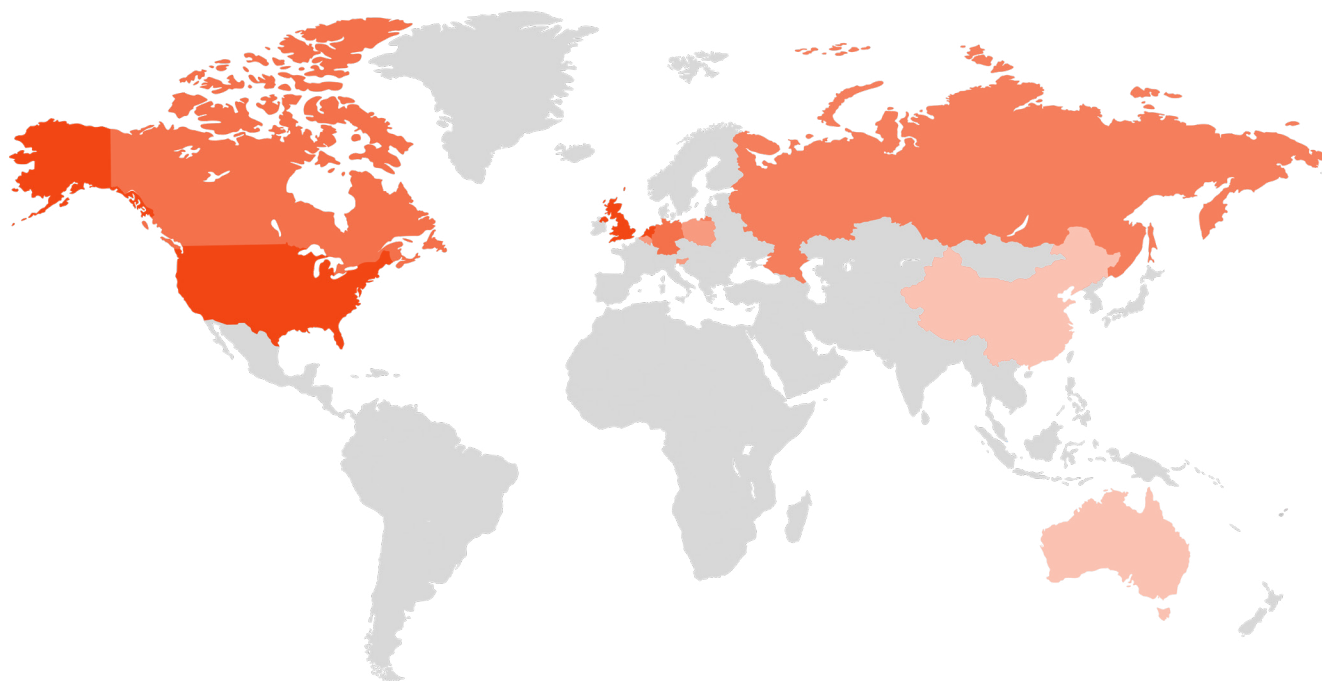


Countries of origin for phishing attacks

Most phishing attacks were traced back to familiar territories, namely the US, the UK, and Russia. Notably, the US consistently dominated as the primary source of these malicious activities. This can be attributed to the country's expansive and advanced digital infrastructure, which gives phishers and cybercriminals easier access to a larger pool of potential victims.

THE TOP 10 COUNTRIES IDENTIFIED AS THE MAIN ORIGINS OF PHISHING ATTACKS WERE:

- | | |
|-------------------|----------------|
| 01 United States | 06 Netherlands |
| 02 United Kingdom | 07 Poland |
| 03 Russia | 08 China |
| 04 Germany | 09 Singapore |
| 05 Canada | 10 Australia |



Australia entered the top 10 due to a 479.3% surge in the volume of phishing content hosted in the country—where 2023 was indeed a notable year for phishing activity, as ACCC's Scamwatch service recorded ~109,000 reports and AU\$26.1 million in losses.²

2. National Anti-Scam Centre Scamwatch, [Scam Statistics \(2023\)](#).

Industries most commonly targeted by phishing attacks

No industry is immune to phishing attacks. After all, the human element permeates every sector—serving as a common vulnerability for phishers to exploit. However, understanding which industries are in phishers' crosshairs is key to strategically allocating anti-phishing resources more effectively, implementing tailored security measures, and prioritizing employee training to mitigate the human error factor.

The finance and insurance sector experienced both the highest number of phishing attempts and the most significant increase in attacks, rising 393% compared to the previous year. This industry is an attractive target for threat actors aiming to engage in identity theft or financial fraud. The increasing reliance on digital financial platforms provides ample opportunities for threat actors to carry out phishing campaigns and exploit vulnerabilities in this sector.

Similarly, **the manufacturing industry experienced a 31% uptick in phishing attacks** from 2022 to 2023. This underscores cybercriminals' awareness of the manufacturing industry's vulnerability to cyberthreats. As manufacturing processes become more reliant on digital systems and interconnected technologies, there is a higher risk of exploitation by threat actors seeking unauthorized access or disruption.

It's no coincidence that manufacturing and finance and insurance are leading adopters of AI tools, collectively driving 35% of AI/ML transactions across the Zero Trust Exchange, as revealed in the [Zscaler ThreatLabz 2024 AI Security Report](#). Adoption of AI technologies and AI-powered systems not only expands connectivity (and thus the exploitable attack surface) across networks and devices, but also makes them even more lucrative targets for phishing schemes given the increased reliance on data.

Despite ranking fourth, **the technology sector saw a 114% surge in phishing attacks**, likely fueled by its early and eager adoption of GenAI and an abundance of valuable data at stake.

THE TOP 5 INDUSTRIES TARGETED FOR PHISHING SCAMS WERE:

- 01 Finance & Insurance
- 02 Manufacturing
- 03 Services
- 04 Technology
- 05 Retail & Wholesale

Share of Phishing Scams by Industry Vertical

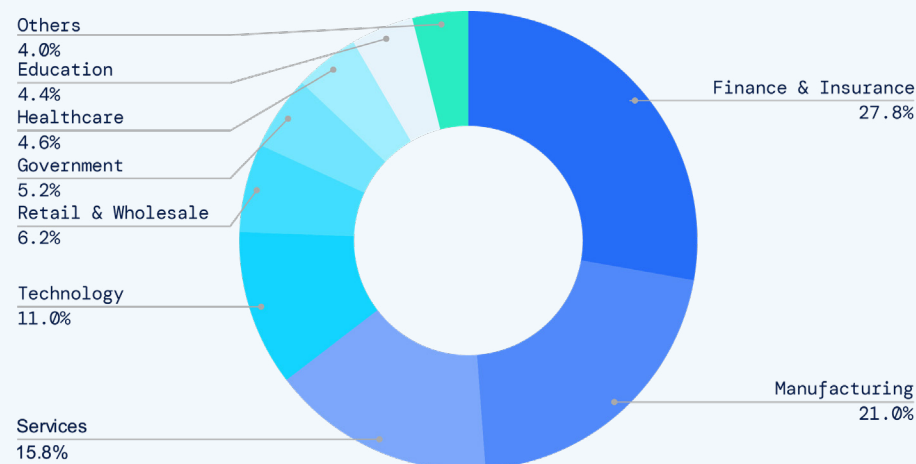


Figure 1: Top industries targeted by phishing scams in 2023

Most Imitated Brands in Phishing Scams

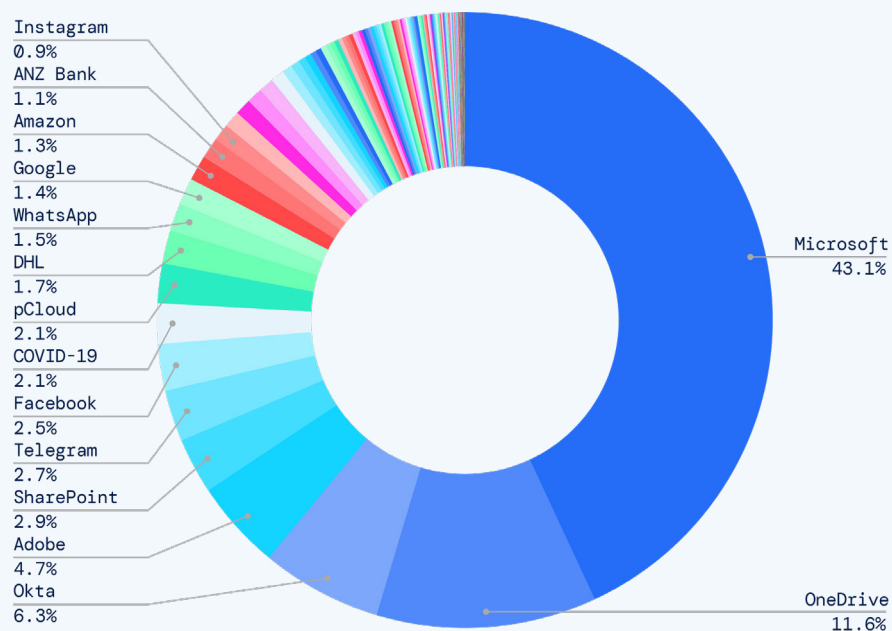


Figure 2: Brands most frequently imitated in 2023

Brands most frequently imitated by threat actors

Phishing attackers exploit popular enterprise applications by impersonating popular brands and themes. ThreatLabz researchers found that enterprise brands like Microsoft, OneDrive, Okta, Adobe, and SharePoint are prime targets for impersonation due to their widespread usage in enterprise environments and the value they hold in acquiring user credentials. This trend has been exacerbated by the shift to remote work culture since 2020, making these brands even more appealing to phishers as they are heavily used for remote work and collaboration.

Microsoft Windows is the world's most widely used computer operating system, and it's no surprise that phishers capitalize on that ubiquity. Microsoft emerged as the top imitated enterprise brand in 2023, with its OneDrive and SharePoint also ranking in the top five.

THE TOP 20 BRANDS MOST FREQUENTLY IMITATED IN PHISHING SCAMS WERE:

01 Microsoft	06 Telegram	11 ANZ Banking Group	16 Sparkasse Bank
02 OneDrive	07 pCloud	12 Amazon	17 FedEx
03 Okta	08 Facebook	13 Ebay	18 PayU
04 Adobe	09 DHL	14 Instagram	19 Rakuten
05 SharePoint	10 WhatsApp	15 Google	20 Gucci

The consumer applications listed serve as a crucial security reminder of the risks of using the same passwords across consumer and enterprise applications. Threat actors frequently exploit this practice, emphasizing the importance of employing strong, unique passwords to mitigate security threats.



Top referring domains leading to phishing pages

Threat actors frequently exploit trusted domains to deceive victims, capitalizing on the familiarity and trust associated with those domains to lead victims to fraudulent phishing sites. Understanding the origins of malicious web traffic, as indicated by referring domains, is crucial to understanding the attack chain. Essentially, it enables organizations to pinpoint the sources of an attack, giving security teams insights into the types of compromised or impersonated websites threat actors are using.

ThreatLabz researchers analyzed the top referring domains in 2023, considering the reputation of the redirected domains and the content hosted by the destinations. The distinction between the top hosts based on reputation and those based on content lies in the methodology employed to identify and categorize websites that may pose potential risks.

When analyzing top referring domains, it's important to consider the potential for open redirect abuse. This tactic involves exploiting vulnerabilities in a website's redirect functionality to deceive users by redirecting them to malicious websites. As a result, legitimate domains may inadvertently end up on lists of top phishing domains. This strategy gives attackers the ability to send emails containing links to these legitimate sites as the entry point while concealing the addresses hidden of actual phishing sites in GET parameters. This tactic increases the likelihood of evading detection by email clients scanning for malicious URLs.

Top referring domains based on reputation

This approach involves evaluating blocked phishing websites based on the reputation of the hosting provider (or “host”) for a particular domain. Additionally, it collects information on the referring domains for these destinations, allowing us to identify websites that redirect users to phishing content or legitimate domains used by threat actors for phishing purposes. The assessment of phishing blocks accounts for various factors such as the hosting provider’s history of abuse, presence of malware, spamming activities, and other indicators of malicious behavior. Websites hosted by providers with a negative reputation may be flagged or categorized as potentially harmful, irrespective of the content on the specific domain.

THE TOP 20 REFERRING DOMAINS BASED ON REPUTATION IN 2023 WERE:

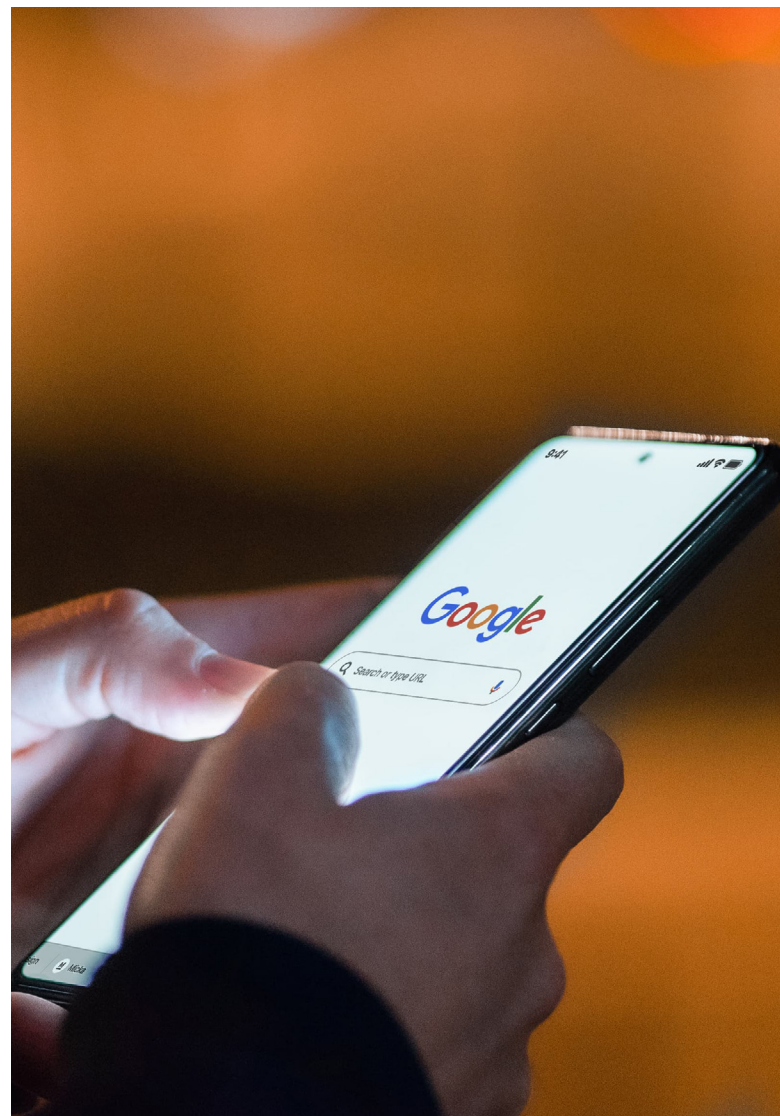
- | | | | | | |
|----|------------------------------|----|---------------------------|----|------------------------------|
| 01 | cstools[.]viagogo[.]net | 08 | app[.]hive[.]com | 15 | onetag-sys[.]com |
| 02 | www[.]gutefrage[.]net | 09 | sync[.]quantumdex[.]io | 16 | evgeny-nadymov[.]github[.]io |
| 03 | web[.]tigrm[.]app | 10 | www[.]google[.]com | 17 | learn[.]hfma[.]org |
| 04 | www[.]mhtestd[.]gov[.]zw | 11 | public[.]servenobid[.]com | 18 | visitor[.]omnitagjs[.]com |
| 05 | framer[.]com | 12 | csync[.]smilewanted[.]com | 19 | blog[.]csdn[.]net |
| 06 | www[.]finanznachrichten[.]de | 13 | t24[.]com[.]tr | 20 | www[.]msn[.]com |
| 07 | webogram[.]org | 14 | acdn[.]adnxs[.]com | | |

Top referring domains based on content

This approach involves examining content found on blocked phishing websites using content scanning, keyword analysis, and machine learning algorithms. By assessing the nature of the hosted content, ThreatLabz flagged websites that match known patterns of malicious activity, such as phishing. The following list showcases the referrer domains associated with the connections where content-based phishing blocks were observed. Note that not all referrer domains are malicious, but they provide valuable insights into the domains leveraged by threat actors to redirect victims to phishing websites.

THE TOP 20 REFERRING DOMAINS BASED ON CONTENT IN 2023 WERE:

- | | |
|-----------------------------|----------------------------------|
| 01 www[.]google[.]com | 12 medsinfoSHOP[.]com |
| 02 mail[.]google[.]com | 13 www[.]bluelightcard[.]co[.]uk |
| 03 rx-qualityshop[.]com | 14 www[.]flickchart[.]com |
| 04 1rotator[.]com | 15 www[.]calendriervip[.]fr |
| 05 webmail[.]ph-japan[.]org | 16 pdce2[.]avanan[.]net |
| 06 www[.]bing[.]com | 17 musicyt[.]click |
| 07 onionplay[.]se | 18 trustedxshop[.]com |
| 08 onionplay[.]co | 19 www[.]onionplay[.]si |
| 09 indd[.]adobe[.]com | 20 www[.]coinpayu[.]com |
| 10 safe-it-phshop[.]com | 21 3khO[.]github[.]io |
| 11 top-sh-op[.]com | |



Distribution of attacks across autonomous systems (ASNs)

An autonomous system is a network or group of networks with a single routing policy. Each AS has a unique identifier known as an autonomous system number (ASN). As part of this analysis, ThreatLabz researchers reviewed the autonomous systems that were responsible for hosting phishing infrastructure.

Insight into the top ASN distributions is vital for cybersecurity teams because it:

- Pinpoints ISPs, businesses, or hosting providers frequently associated with cyberthreats, aiding in targeted threat intelligence
- Helps attribute cyberattacks to specific organizations or regions, crucial for understanding motives and identifying potential threat actors

The data indicates the following distribution:

- **Internet service provider (ISP):** With a total of 200,293,568, the majority of ASNs belong to ISPs. These ASNs are associated with organizations that provide internet connectivity services to end users, households, and businesses.
- **Hosting:** Hosting ASNs account for 112,452,292, representing a significant portion of the distribution. These ASNs are associated with hosting providers that offer server space, infrastructure, and related services for websites and online applications.
- **Business:** ASNs associated with businesses make up 75,826,357 in total. These ASNs are assigned to organizations across various industries that operate their own networks for internal communication, data exchange, and internet connectivity.

Top ASN Distribution Types

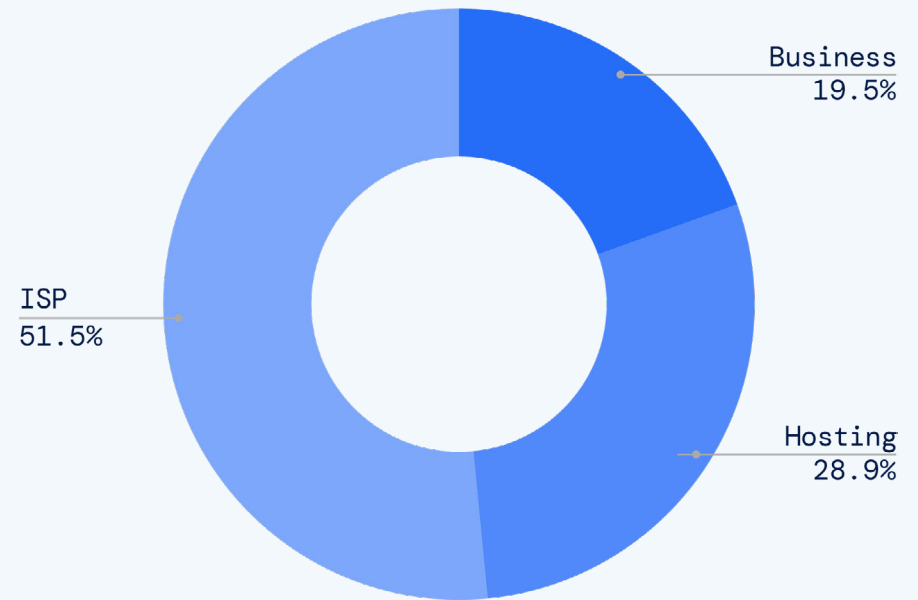


Figure 3: A breakdown of business, hosting, and ISP servers involved in phishing attacks



Social media platforms exploited by threat actors

In a world where social media reigns supreme, attackers are increasingly leveraging these platforms for phishing endeavors. This trend spans the globe, with the Asia-Pacific, Europe, the Middle East, and Africa experiencing similar patterns of exploitation. Figure 4 shows the most targeted social media platforms observed by ThreatLabz.

Telegram, with 792,883 observed phishing hits, remains a popular target for malicious activities—a trend explored in [our blog post on DuckTail](#). The platform’s end-to-end encryption and emphasis on user privacy make it an attractive choice for secure communication. However, threat actors attempt to exploit vulnerabilities in Telegram’s security measures to gain unauthorized access to user accounts or distribute malicious content.

Facebook, with 532,243 observed phishing hits, faces ongoing challenges in protecting user data and privacy. As one of the largest social media platforms globally, it attracts cybercriminals who aim to exploit security flaws, launch phishing campaigns, or engage in identity theft.

Most Exploited Social Media Platforms Worldwide

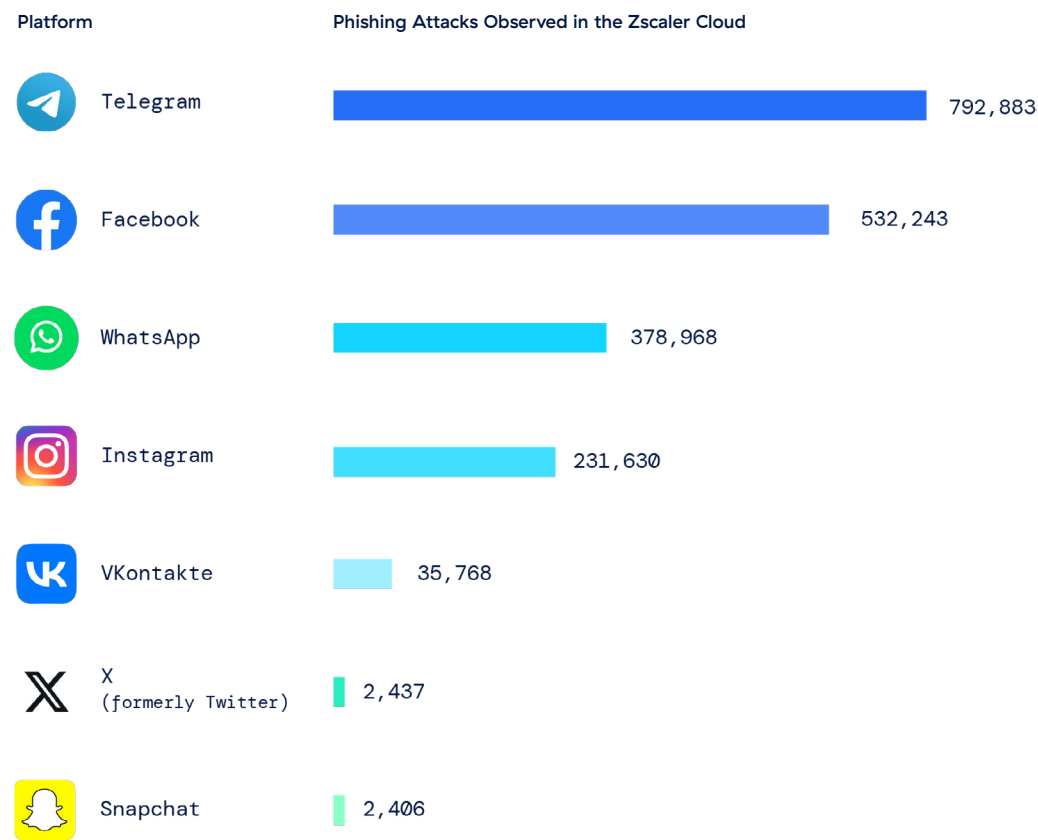
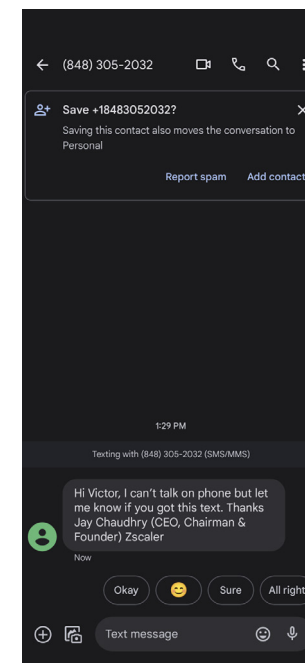
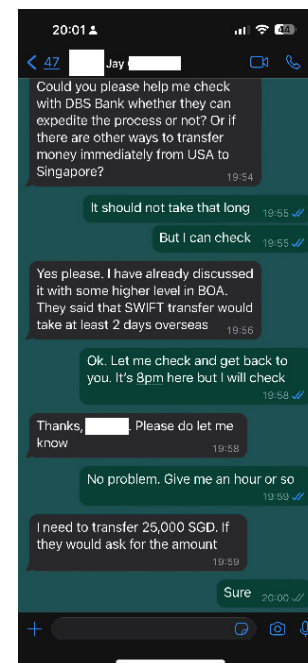
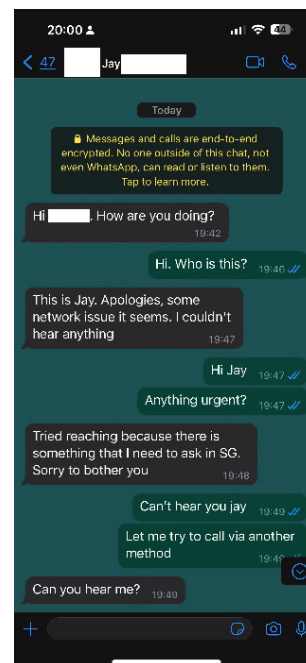
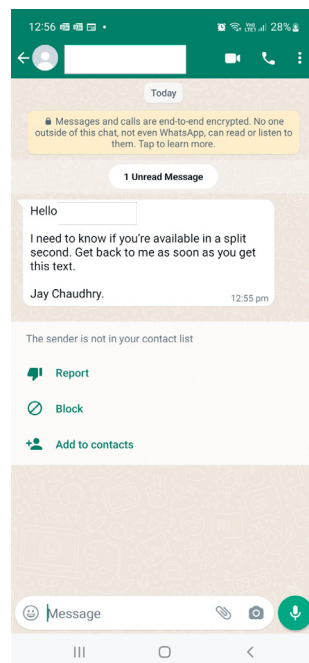


Figure 4: Top social media platforms used in phishing attacks

WhatsApp, with 378,968 observed phishing hits, encounters various security concerns due to its large user base and ubiquitous usage for messaging. While WhatsApp incorporates end-to-end encryption for secure conversations, attackers seek to exploit vulnerabilities to gain unauthorized access, distribute malware, or deceive users through social engineering techniques.

ThreatLabz detected the following phishing attempts leveraging WhatsApp:



[Read more about this case study below.](#)

Instagram, with 231,630 observed phishing hits, grapples with threats such as account hijacking, phishing attempts, and the spread of malicious links or content. As a leading photo and video sharing platform, it attracts cybercriminals who exploit weak passwords, social engineering tactics, or third-party app vulnerabilities to compromise user accounts.

Vkontakte, with 35,768 observed phishing hits, encounters security challenges specific to its user base in Russia and neighboring countries. Cyberthreats targeting VKontakte, a social media and networking service based in Russia, include account breaches, phishing attacks, and the distribution of malicious content.

X (previously Twitter), with 2,437 observed phishing hits, encounters a range of security issues, including account breaches, impersonation attempts, and the dissemination of fake news or malicious links. X's real-time nature and large user base make it an attractive target for cybercriminals seeking to spread misinformation or compromise user accounts.

Snapchat, with 2,406 observed phishing hits, faces unique security concerns related to its multimedia messaging features and user-generated content. While Snapchat's self-destructing messages provide a level of privacy, attackers may attempt to exploit vulnerabilities to compromise accounts or engage in social engineering scams.

One Phish, Two-Faced: AI and Phishing

What happens when cunning phishing tactics meet the power of AI? The convergence of these two forces signifies a profound revolution in cyberthreats.

AI-driven phishing attacks leverage AI tools to enhance the sophistication and effectiveness of phishing campaigns. AI automates and personalizes various aspects of the attack process, making phishing even more challenging to detect. For example, chatbots are commonly used to craft highly convincing, error-free phishing emails. What's more, attackers are increasingly harnessing advanced AI services such as deepfake technology and voice cloning to impersonate reputable organizations or people and deceive victims. They exploit various communication channels, including emails, phone and video calls, SMS, and encrypted messaging applications.

These advanced tactics remind us of the importance of vigilance and skepticism when interacting with digital communications, as well as the need for organizations to implement robust cybersecurity measures to mitigate the risk of falling victim to AI-driven phishing attacks.

Phishers abuse AI, AI fights back

Generative AI is rapidly driving the phishing threat landscape forward, enabling automation and efficiency across numerous stages of the attack chain. By rapidly analyzing publicly available data, such as details about organizations or executives, GenAI saves threat actors time in reconnaissance while facilitating more precise targeted attacks. By eliminating spelling errors and grammatical mistakes, GenAI tools enhance the credibility of phishing communications. What's more, GenAI can quickly create sophisticated phishing pages—as demonstrated in the following case study—or extend its capabilities to generate malware and ransomware for secondary attacks. As GenAI tools and tactics rapidly evolve, phishing attacks will become more dynamic (and challenging to detect) by the day.

The growing popularity and use of GenAI tools like ChatGPT and Drift is already beginning to impact phishing activity and the rise of AI-driven attacks. Countries like the US and India, where these tools are highly utilized according to ThreatLabz research in the [2024 AI Security Report](#), are top targets for phishing scams and face the [highest number of encrypted attacks](#) in the past year, a subset of which are phishing attacks.



THREATLABZ ACTIVELY TRACKS THE ABUSE OF LEGITIMATE AND MALICIOUS LARGE LANGUAGE MODELS (LLMS) TO ENSURE COMPREHENSIVE COVERAGE AGAINST PHISHING ATTACKS FOR ZSCALER TO FIGHT AI WITH AI INNOVATIONS, INCLUDING:

AI-powered phishing and C2 prevention

Zscaler AI models detect known and patient-zero phishing sites to prevent credential theft and browser exploitation, as well as analyze traffic patterns, behavior, and malware to detect never-before-seen command-and-control (C2) infrastructure in real time. These models draw on a combination of threat intelligence, [ThreatLabz research](#), and dynamic browser isolation to detect suspicious sites. As a result, enterprises are even more efficient and effective in detecting new phishing attacks, including AI-generated attacks, and C2 domains.

File-based AI sandbox defense

[The AI-powered inline Zscaler Sandbox](#) instantly detects malicious files while keeping employees productive. Traditional sandbox technologies make users wait while files are analyzed, or else assume patient-zero risk when files are allowed on first pass. Our AI Instant Verdict technology instantly identifies, quarantines, and prevents high-confidence malicious files—including zero-day threats—while removing the need to wait for analysis on these files. This includes threats that are delivered over encrypted channels (TLS and HTTPS) and other file transfer protocols. Meanwhile, benign files are delivered safely and instantly.

AI to block web threats

[AI-powered Zscaler Browser Isolation](#) blocks zero-day threats while ensuring employees can access the right sites to do their jobs. In practice, enterprise URL filtering often requires more granular controls than allow/block; blocked sites are often safe and required for work, resulting in needless help desk tickets. Our AI Smart Isolation can identify when a site may be risky and open it in isolation for the user—safely streaming the site as pixels in a secure, containerized environment. This effectively stops web-based threats like malware, ransomware, phishing, and drive-by downloads, creating a strong web security posture without requiring enterprises to overblock sites as a default.



Case study: Creating a phishing login page in less than 10 ChatGPT queries

LLMs have made it significantly easier for cybercriminals, even with relatively little coding experience, to carry out multiple stages of a sophisticated phishing attack.

For instance, in a few prompts using a generative AI chatbot like ChatGPT, it's almost trivial to create fake phishing login pages that mimic popular enterprise applications to steal employee login credentials. The following example from ThreatLabz shows how simple it is to create a convincing fake Microsoft login page with just a few conversational prompts.

Note: For demonstration purposes, this example shows lightly abbreviated prompts and includes a ChatGPT code response for one query before showing the final rendered phishing page.

User prompt: create an HTML login page

ChatGPT: Sure, here's a basic example of an HTML login page:

```

<!DOCTYPE html>
<html lang="en">
<head>
<meta charset="UTF-8">
<meta name="viewport" content="width=device-width, initial-scale=1.0">
<title>Login Page</title>
<style>
body {
font-family: Arial, sans-serif;
background-color: #f4f4f4;
}
.login-container {
width: 300px;
margin: 100px auto;
padding: 20px;
background-color: #fff;
border-radius: 5px;
box-shadow: 0 0 10px rgba(0, 0, 0, 0.1);
}
h2 {

```

THIS WAS THE FIRST RESULT:

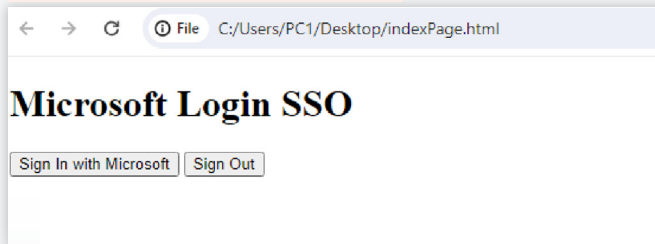


FIGURE 5: Screenshot of a rendered login page using the ChatGPT code response

Next, ThreatLabz provided a short series of prompts to improve the page before rendering the final Microsoft phishing login page. These included asking ChatGPT to make the page look like a Microsoft login, adjusting the logo size, and adding and removing UI elements before submitting the final query to generate the final code output.

User prompt: add a page background that is similar to the microsoft login page

ChatGPT:

IN 7 QUERIES, THE FINAL RESULT:

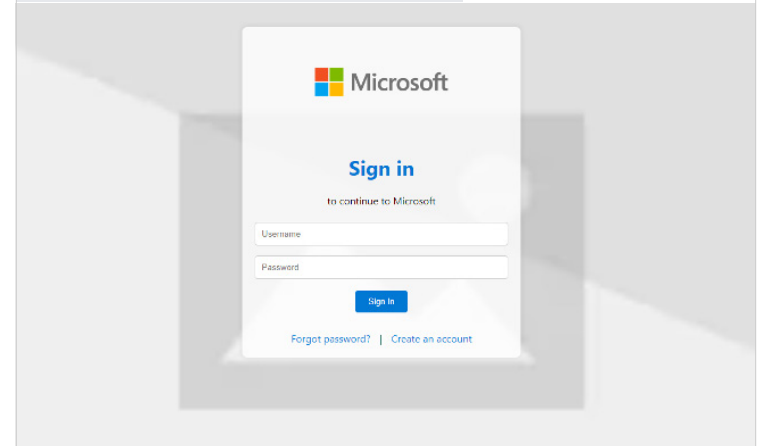


FIGURE 6: Screenshot of the final rendered Microsoft phishing login page, using the ChatGPT code response

Rise in deepfake attacks

Deepfake technology uses AI and machine learning techniques to create realistic and manipulated audio or video content that appears to be authentic. This can range from superimposing someone's face onto another person's body to altering their facial expressions and voice to create a convincing and often deceptive representation.

Deepfake technology utilizes algorithms and neural networks to analyze and learn from vast amounts of data, such as images, videos, and audio recordings of a specific individual. With this information, the AI model can generate new content that mimics the person's appearance, voice, and mannerisms.

Deepfake attacks are already causing significant financial losses for organizations. In a recent incident, a finance worker unknowingly paid out \$25 million to fraudsters who were using deepfake technology to impersonate the worker's colleagues in a video call³. The attackers posed as the company's chief financial officer and manipulated publicly available video to deceive the worker into carrying out a fraudulent transaction.

Realistic, deepfake-driven attacks costing organizations millions of dollars is not science fiction—it's today's threat landscape.

Case study: Deepfake campaign impersonates Elon Musk

In Summer 2023, threat actors orchestrated a deepfake campaign using the likeness and reputation of entrepreneur Elon Musk.

The campaign revolves around the use of fake ads to deceive individuals into "investing" money in a new platform called "Quantum AI." These ads could be found on social media platforms and search engine results.

The campaign aimed to solicit funds from victims by promising remarkably high returns, such as a staggering 91%. Musk is portrayed in the main ad for "Quantum AI," although he appears distant and out of focus. The video mimics his voice and features a typical tech conference-style product unveiling.

Additionally, a secondary ad takes the form of a fabricated Fox News web page, claiming that Musk gave an interview promoting Quantum AI.

3. CNN, [Finance worker pays out \\$25 million after video call with deepfake 'chief financial officer'](#), February 4, 2024.

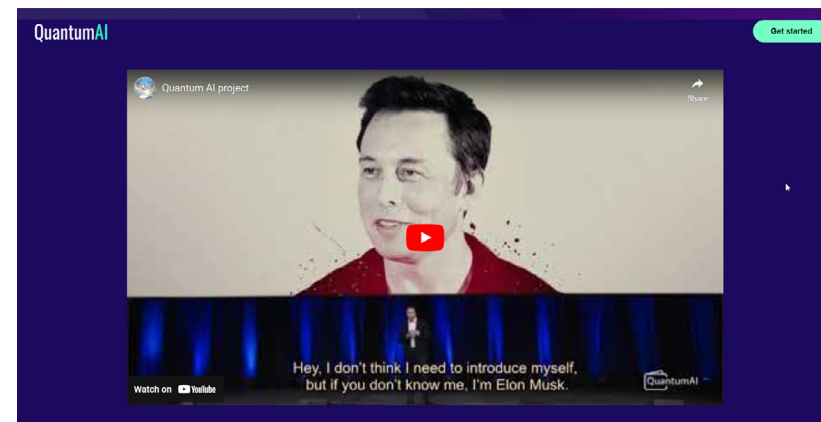


FIGURE 7: Screenshot of the final rendered Microsoft phishing login page, using the ChatGPT code response



FIGURE 8: A fraudulent Fox News web page promoting the fake Quantum AI platform

Election Spotlight: Phishing Campaigns vs. Political Campaigns

With more than half of the world's population living in countries holding nationwide votes in 2024, it will be a record year for elections. Amid political campaigns and fervor, the specter of cyberthreats stands out—particularly in the form of phishing campaigns.

Phishing threats cast a long shadow over election security and the integrity of democratic processes worldwide. Throughout history, cybercriminals have used phishing tactics to manipulate voters, spread disinformation, and compromise critical election infrastructure.

Looking back at the 2020 US presidential election, threat actors targeted voters in pivotal swing states with advanced phishing emails disguised as official communications from government entities or political campaigns, urging recipients to confirm voter registration details or request absentee ballots through fraudulent links.

The rise of generative AI raises the stakes for election security this year and beyond. Advancements in AI technology foreshadow the potential for real impact when it comes to phishing and election outlines. The aforementioned rise of deepfake technology has already introduced a new dimension of deception into the electoral process⁴. Deepfake phishing videos manipulated to depict false narratives or statements from political figures can sway public opinion, disseminate disinformation, and erode trust in the electoral process itself.

ThreatLabz recently uncovered a concerning instance of advanced persistent threats (APTs) targeting political entities—a case of cyber espionage by the threat actor SPIKEDWINE, using phishing tactics to exploit geopolitical relations between India and European diplomats. In January 2024, ThreatLabz discovered a suspicious PDF on VirusTotal disguised as an invitation letter from the Ambassador of India (though originating from Latvia) for a government-related wine-tasting event. The PDF contained

a link to a fake questionnaire, redirecting users to a malicious ZIP archive on a compromised website. This discovery revealed a new backdoor, “WINELOADER.” You can read a full technical analysis of the attack chain on the [Zscaler blog](#).

Zscaler threat hunters identified a similar PDF uploaded to VirusTotal from Latvia in July 2023, indicating a pattern of targeted attacks and emphasizing the need to protect political processes and relations, especially in light of the election season.

Election security efforts must prioritize [proactive measures to detect and mitigate phishing attacks](#). Fostering collaboration between election officials, cybersecurity experts, and law enforcement agencies as well as widely promoting phishing awareness and security best practices will help protect citizens and organizations alike against evolving phishing threats.

Phishing examples in electoral history



4. Bloomberg, [How Bad Movie Dubbing Led to the Fake Biden Campaign Robocalls](#), February 20, 2024.

Evolving Phishing Trends

Threat actors are always refining their tactics and strategies to perpetrate more effective scams, so keeping up with developing phishing trends is essential to establishing and maintaining proactive defenses.

ThreatLabz researchers diligently tracked phishing trends throughout 2023. In this section, we'll delve into several notable trends that emerged, highlighting the ingenuity and sophistication fueling the surge in phishing.

Vishing attacks

Voice phishing, known as vishing, involves deceiving individuals through phone calls and voice messages, often using familiar or authoritative voices to gain trust and extract sensitive information.

Sophisticated vishing campaigns are becoming popular worldwide, with cybercriminals using psychology and technology to defraud even savvy victims of millions of dollars. For example, South Korea has experienced a surge in vishing attacks, including a case in August 2022 where a doctor lost \$3 million in cash, insurance, stocks, and cryptocurrency to criminals⁵. In this case, scammers impersonated regional law enforcement officials in South Korea; however, ThreatLabz observed (and thwarted) a vishing attack very close to home in 2023.

5. Dark Reading, [Sophisticated Vishing Campaigns Take World by Storm](#), March 11, 2024.

Vishing case study

In the summer of 2023, attackers impersonated Zscaler's own CEO, Jay Chaudhry, in a vishing attack using AI technology. It unfolded like this:



The attacker called a Zscaler employee on WhatsApp.



Using AI-generated voice cloning to simulate Jay's voice, the attacker established communication, and then quickly hung up to avoid prolonged interaction and potential exposure.



The attacker immediately followed up with a text message—posing as Jay—claiming to have “poor network coverage.”



In a WhatsApp text message, the attacker instructed the Zscaler employee to purchase gift cards for a certain amount.



The employee found this suspicious and immediately reported it to the security team.



ThreatLabz researchers investigated and found out it was part of a widespread campaign targeting several tech companies.

You can watch Jay as he explains the full sequence of events on [NBC Bay Area](#).

Recruitment scams

Recruitment scams aim to deceive and exploit job seekers. These scams often involve the creation of fake job postings on reputable job boards, social media networks, and professional networking websites like LinkedIn. Attackers impersonate legitimate companies or recruiters and manipulate victims into divulging sensitive information or downloading malware.

Unfortunately, the tech layoffs in 2022, 2023, and 2024 introduced a new crop of eager candidates to the digital market, meaning more prime targets for recruitment scammers.

LinkedIn recruiter scam case study

One of the primary distribution channels for recruitment scams by [DuckTail threat actors](#) is LinkedIn, a widely trusted professional networking platform. Threat actors capitalize on the platform's credibility and its users' trust to disseminate fraudulent job postings. By impersonating reputable companies and leveraging fake recruiter profiles, they lure victims with enticing job opportunities. Once a candidate expresses interest in a fake position, the threat actor initiates contact through private messages on LinkedIn, starting the social engineering process. The threat actor shares a malicious file disguised as a job description, which infects the victim's system when downloaded.

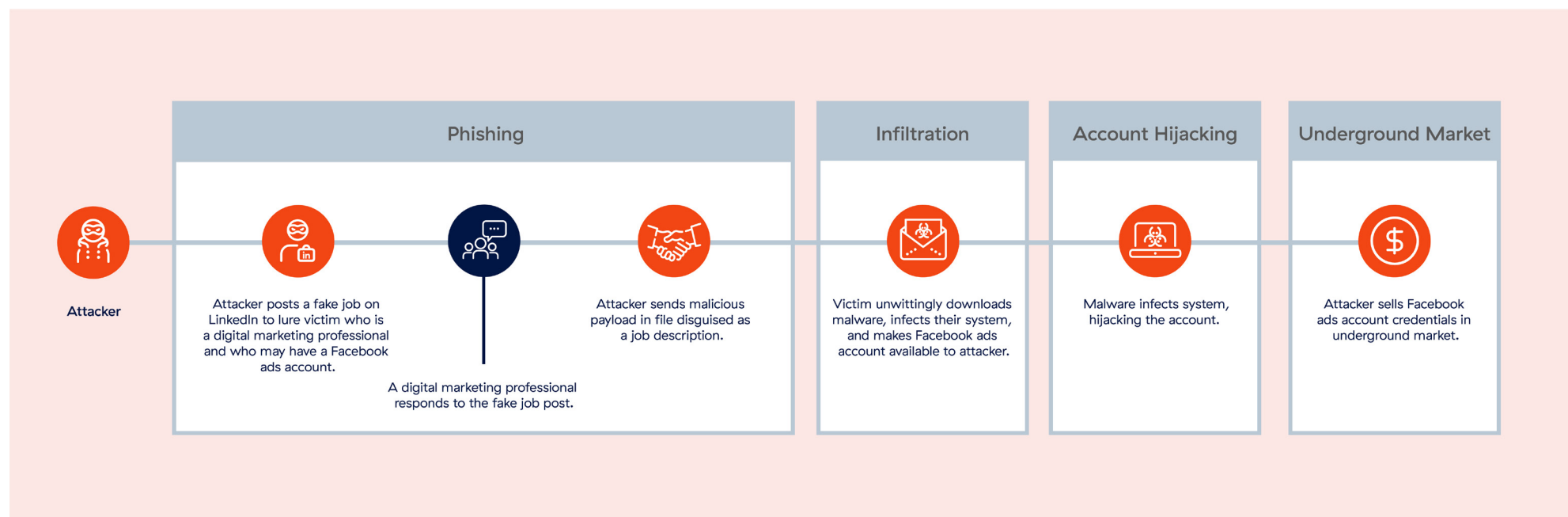


Figure 9: Recruitment scam attack sequence

Adversary-in-the-middle (AiTM) attacks

In an adversary-in-the-middle (AiTM) phishing attack, the adversary intercepts and manipulates communications between two parties to deceive the victim. By positioning themselves between the victim and a trusted entity, the attacker gains unauthorized access to sensitive information. Unlike traditional phishing attacks, AiTM attacks happen in real time, enabling attackers to monitor and modify communications. They can alter messages, redirect victims to malicious websites, and collect data without detection. Protecting against AiTM phishing involves:

- Using secure communication channels
- Verifying website authenticity
- Exercising caution when sharing sensitive information
- Keeping software updated

Learn more about [AiTM phishing attacks on the Zscaler blog](#).

AiTM case study

ThreatLabz researchers observed an increase in the use of advanced phishing kits in a large-scale campaign. This campaign stands out as it uses an AiTM attack technique capable of bypassing multifactor authentication and is specifically designed to target enterprise end users.

AiTM proxy kits have evolved to generate phishing pages that closely mimic legitimate web pages, making them difficult to distinguish from benign network traffic. Moreover, the use of proxy kits provides threat actors a straightforward way to spread phishing pages effectively.

ThreatLabz has detected AiTM phishing attacks targeting enterprise users of Microsoft and Gmail in the Zscaler cloud. By closely monitoring these attacks, ThreatLabz has found a persistent trend in the targeting of enterprise users of Microsoft. In addition, these AiTM phishing attacks have shown resilience over an extended period.

Figure 10 shows a side-by-side comparison of the source code of an AiTM served Microsoft page and the legitimate Microsoft login page.

<pre> 1 <!-- Copyright (C) Microsoft Corporation. All rights reserved. --> 2 <!DOCTYPE html> 3 <html dir="ltr" class="" lang="en"> 4 <head> 5 <title>Sign in to your account</title> 6 <meta http-equiv="Content-Type" content="text/html; charset=UTF-8"> 7 <meta http-equiv="X-UA-Compatible" content="IE=edge"> 8 <meta name="viewport" content="width=device-width, initial-scale=1.0, maximum-scale=2.0, user-scalable=yes"> 9 <meta http-equiv="Pragma" content="no-cache"> 10 <meta http-equiv="Expires" content="-1"> 11 <link rel="preconnect" href="https://recyclingpartnership.org/aadcdn.esauth.net/~" rickorigin> 12 <meta http-equiv="x-dns-prefetch-control" content="on"> 13 <link rel="dns-prefetch" href="//recyclingpartnership.org/aadcdn.esauth.net/~"> 14 <link rel="dns-prefetch" href="//recyclingpartnership.org/aadcdn.msftauth.net/~"> 15 <meta name="PageID" content="ConvergedSignIn"/> 16 <meta name="SiteID" content=""> 17 <meta name="ReqLC" content="1033"/> 18 <meta name="LocLC" content="en-US"/> 19 <meta name="Format-detection" content="telephone=no"/> 20 <noscript> 21 <meta http-equiv="Refresh" content="0; URL=https://recyclingpartnership.org/jsdisabled/"> 22 </noscript> 23 <meta name="robots" content="none"/> 24 <script type="text/javascript"> 25 //<![CDATA[26 \$Config = { 27 "fShowPersistentCookiesWarning": false,</pre> </td> <td style="vertical-align: top; width: 50%;"> <pre> 1 <!-- Copyright (C) Microsoft Corporation. All rights reserved. --> 2 <!DOCTYPE html> 3 <html dir="ltr" class="" lang="en"> 4 <head> 5 <title>Sign in to your account</title> 6 <meta http-equiv="Content-Type" content="text/html; charset=UTF-8"> 7 <meta http-equiv="X-UA-Compatible" content="IE=edge"> 8 <meta name="viewport" content="width=device-width, initial-scale=1.0, maximum-scale=2.0, user-scalable=yes"> 9 <meta http-equiv="Pragma" content="no-cache"> 10 <meta http-equiv="Expires" content="-1"> 11 <link rel="preconnect" href="https://aadcdn.msftauth.net" crossorigin> 12 <meta http-equiv="x-dns-prefetch-control" content="on"> 13 <link rel="dns-prefetch" href="//aadcdn.msftauth.net"> 14 <link rel="dns-prefetch" href="//aadcdn.esauth.net"> 15 <meta name="PageID" content="ConvergedSignIn"/> 16 <meta name="SiteID" content=""> 17 <meta name="ReqLC" content="1033"/> 18 <meta name="LocLC" content="en-US"/> 19 <meta name="Format-detection" content="telephone=no"/> 20 <noscript> 21 <meta http-equiv="Refresh" content="0; URL=https://login.microsoftonline.com/jsdisabled/"> 22 </noscript> 23 <meta name="robots" content="none"/> 24 <script type="text/javascript"> 25 //<![CDATA[26 \$Config = { 27 "fShowPersistentCookiesWarning": false,</pre> </td> </tr> </table> </div> <div data-bbox="199 886 797 904" data-label="Caption"> <p>Figure 10: Source code of an adversary-in-the-middle served Microsoft page (left) vs. a legitimate Microsoft login page (right)</p> </div> <div data-bbox="25 956 166 972" data-label="Page-Footer"> <p>©2024 Zscaler, Inc. All rights reserved.</p> </div> <div data-bbox="683 956 890 971" data-label="Page-Footer"> <p>ZSCALER THREATLABZ 2024 PHISHING REPORT</p> </div> <div data-bbox="943 956 970 971" data-label="Page-Footer"> <p>022</p> </div>]]></pre>

Browser-in-the-browser (BiTB) attacks

In browser-in-the-browser (BiTB) attacks, cybercriminals aim to deceive users by simulating a browser window within another browser to spoof a legitimate domain. In these sophisticated attacks, attackers manipulate the appearance of a web page to make users believe they are interacting with a trusted website, when in fact it's a malicious counterfeit.

For example, an attacker might use a combination of HTML/CSS and inline frames (iframes) to create an authentic-looking login pop-up in a main phishing page, which prompts a user for credentials. Unfortunately, when the user enters their credentials, they are sharing them with the attacker. Even perceptive or experienced users can be fooled because it is almost impossible for them to distinguish a genuine pop-up from a well-designed phishing fake.

Initially, BiTB attacks were primarily designed to simulate legitimate websites in a browser window in order to steal sensitive information, such as login credentials or financial data. However, as cybercriminals adapt their strategies, BiTB attacks have expanded in scope to include sextortion schemes.

BiTB case study

In BiTB attack variants recently observed in the Zscaler cloud, attackers have impersonated government agencies, law enforcement, or other authoritative entities of victims' respective countries to carry out a form of sextortion attack.

Attackers manipulate victims' browsers to display messages or notifications that often falsely accuse the victims of illegal activities and threaten legal action unless a ransom is paid or sensitive information is provided. By exploiting the credibility of authorities, attackers aim to coerce victims into compliance, leading to extortion or further exploitation of personal data.

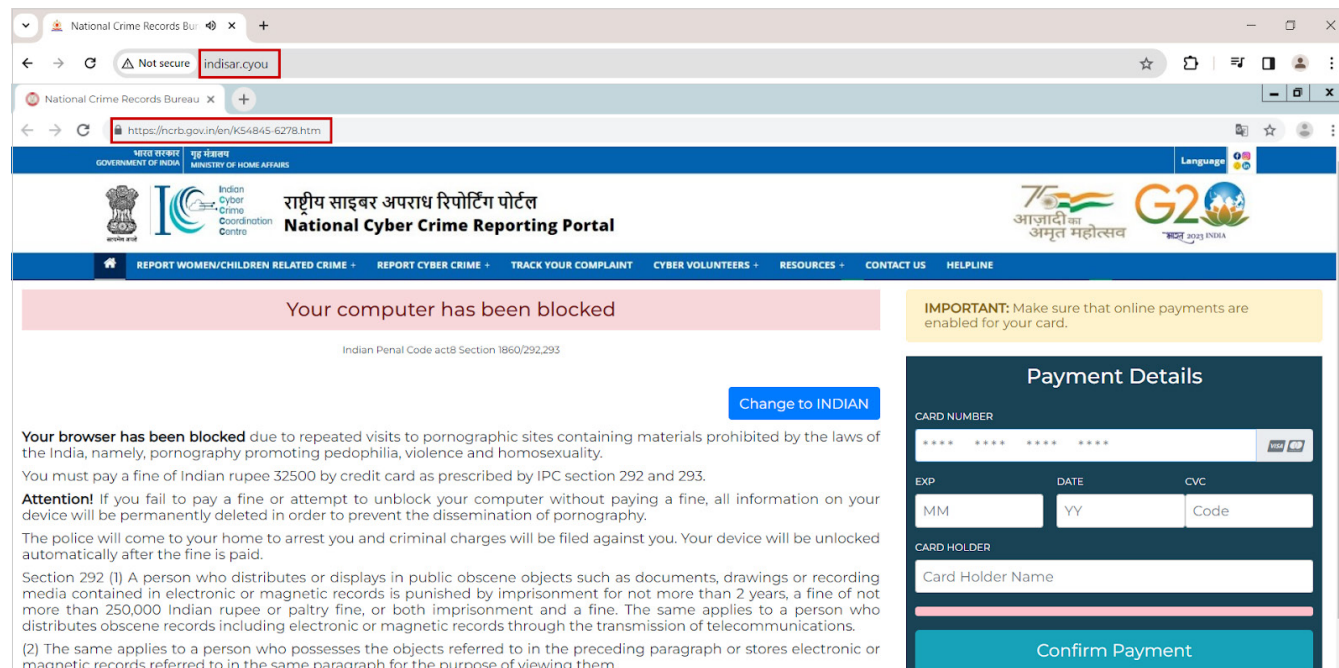


Figure 11 shows a BiTB attack posing as a cybercrime agency, falsely claiming the victim's computer is blocked. However, the main wrapping browser's appearance does not match that of a legitimate agency website. The unusual design raises suspicion about the message's authenticity and credibility.

Figure 11: An example of a browser-in-the browser (BiTB) attack

QR scams

In QR scams, threat actors trick victims into scanning QR codes that ultimately lead to malicious links. This scam employs various delivery methods, including malicious redirection, email attachments (such as PDF or DOC files), and other forms of digital communication. Threat actors’ most common method is to email a PDF containing a QR image, which then redirects victims to a phishing page.

Unusual URLs

Identifying a QR code scam necessitates examining the associated URL. Legitimate URLs are typically error-free, without misspellings or unusual phrases in their domain name or URL path. Figure 12 shows an unusual pairing, with “gard-ner” next to the usually legitimate “Toyota”.

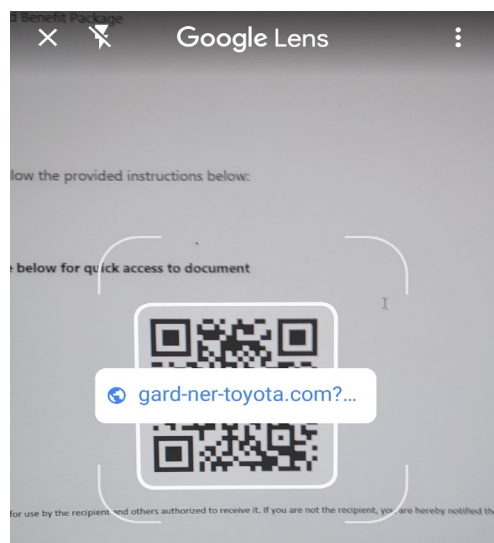


Figure 12: An example of a malicious URL associated with a QR code

Fake CAPTCHA process

Scam websites that employ QR codes as a CAPTCHA, such as the example in figure 13, present a deceptive twist to the typical CAPTCHA verification process. Instead of the usual image- or text-based challenge, these sites prompt users to scan a QR code. Then, scammers redirect users to fraudulent websites or extract sensitive information.

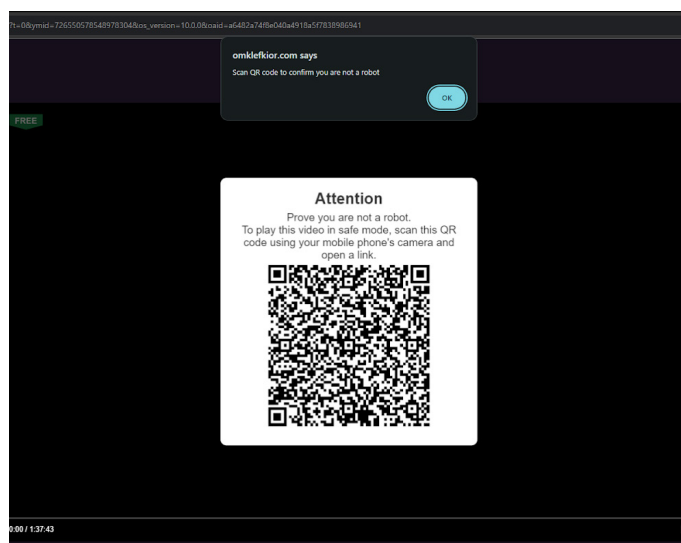


Figure 13: A scam website using a QR code as a CAPTCHA verification method

Fake security updates

Phishing emails incorporating QR codes, particularly when disguised as security updates, introduce a new level of deception to exploit unsuspecting recipients. These emails often appear to come from reputable sources, such as well-known companies or service providers, claiming to offer important security updates or account verification. Figure 14 shows an example.

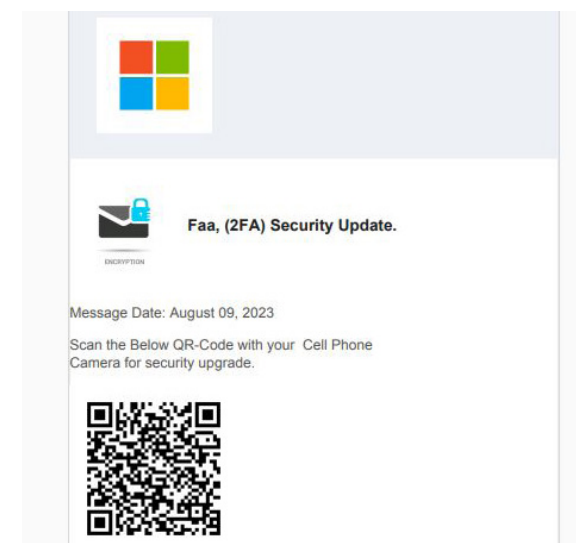


Figure 14: A phishing email using a QR code to trick users into a fraudulent “security update”

Tech support scams

Tech support scams trick users into believing their device is infected with malware or viruses. Attackers use tactics like pop-up messages, fake antivirus alerts, or alarming emails to exploit fear and limited technical knowledge, creating urgency and panic. They urge immediate action, such as downloading software or granting remote access for system cleanup, which will subsequently help them gain access to sensitive information or financial details.

In a [recent tech support scam](#), ThreatLabz identified a method where threat actors exploit Windows Action Center notifications to deceive users. By manipulating the built-in notification system, attackers generate fake warnings and alerts that closely resemble legitimate system messages, using logos and language that appear authentic. These deceptive notifications falsely claim virus infections, outdated software, or security vulnerabilities, and then prompt victims to contact a fraudulent tech support number or visit a malicious website for assistance in resolving the fabricated issues.

Typical of tech support scams, this technique aims to incite urgency and panic, coercing victims into revealing personal information, granting remote access, or purchasing unnecessary and potentially harmful services or software.

Figures 15 and 16 show two fraudulent pop-ups trying to convince a victim that their system is infected.

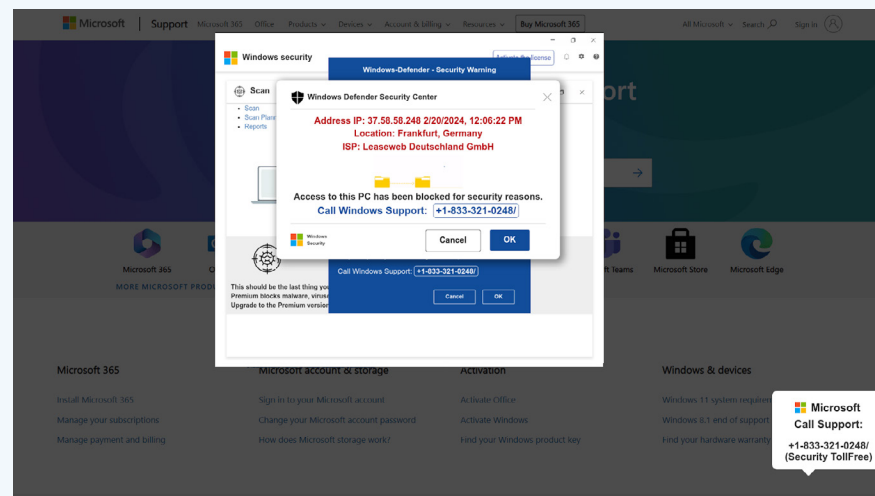


Figure 15: A screenshot of a fake Windows Defender Security Center pop-up

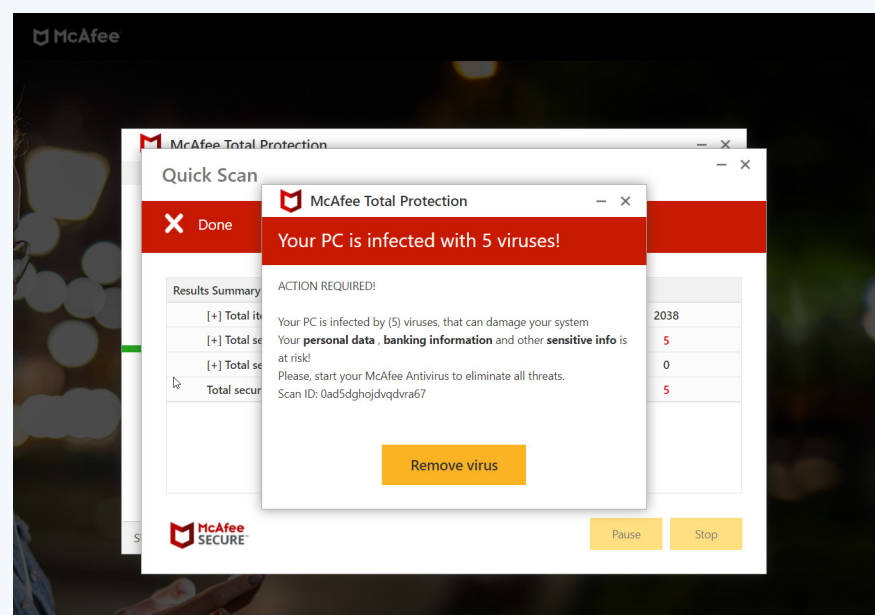


Figure 16: An alarming pop-up disguised as a McAfee notification

2024–2025 Predictions

1 AI vs. AI will be an enduring challenge.

In 2025, we anticipate a significant transformation in cyberattack and defense strategies with the widespread adoption of generative AI. Threat actors will widely adopt AI to craft more sophisticated phishing schemes and advanced techniques. Simultaneously, security vendors will integrate generative AI into their toolkits to enhance threat detection and response capabilities. This era introduces an inescapable reality: AI will be a double-edged sword as both threat actors and defenders utilize its power. AI-powered security measures will be required to effectively counter AI-driven attacks.

Although targeted intervention has stopped some of these attacks, enterprises should brace for the persistence of state-backed AI initiatives. The scope encompasses the deployment of popular AI tools, the creation of proprietary LLMs, and the emergence of unconstrained ChatGPT-inspired variants, such as the aptly-named FraudGPT or WormGPT. The evolving landscape paints a challenging picture in which state-sponsored actors continue to leverage AI in novel ways to create complex new cyberthreats.

2 Phishing as a service will intensify its focus on MFA exploitation and AiTM.

Over the past year, a concerning trend has emerged where adversaries successfully circumvent enterprise multifactor authentication (MFA) through adversary-in-the-middle (AiTM) proxy-based phishing attacks. In the coming year, we expect phishing kits to increasingly include sophisticated AiTM techniques, localized phishing content, and target fingerprinting—of course enabled by AI. These advancements will allow attackers to conduct high-volume phishing campaigns aimed at evading MFA protections at enterprise scale.

3 Vishing attacks spearheaded by malware groups will surge significantly.

Expect an uptick in targeted voice and video phishing campaigns carried out by groups like Scattered Spider, renowned for using sophisticated tactics and techniques. These campaigns will focus on obtaining employee login credentials to gain unauthorized access to secure systems, potentially leading to further exploitation, persistence, data exfiltration, and even organization-wide breaches. Coupled with the prevalence of AI-powered voice and video tools, this may make it even easier for threat actors to impersonate corporate personnel, posing new challenges for employees in identifying these phishing attacks.

4 Attackers will home in on vulnerabilities inherent in mobile devices and platforms.

Over the past year, a concerning trend has emerged where adversaries successfully circumvent enterprise multifactor authentication (MFA) through adversary-in-the-middle (AiTM) proxy-based phishing attacks. In the coming year, we expect phishing kits to increasingly include sophisticated AiTM techniques, localized phishing content, and target fingerprinting—of course enabled by AI. These advancements will allow attackers to conduct high-volume phishing campaigns aimed at evading MFA protections at enterprise scale.

5 Expect a surge in phishing tailored to disrupt electoral processes.

These scams will encompass everything from voter registration manipulation to spreading of disinformation aimed at swaying public opinion. Beyond the scope of traditional profit-driven phishing, these campaigns will pivot toward a more insidious objective: capturing mindshare and influencing political outcomes. Attackers will exploit vulnerabilities inherent in the digital landscape to manipulate user trust and disseminate deceptive narratives, enabled by AI-powered phishing tactics like the creation of highly personalized and persuasive messaging. This shift will pose a serious threat to the fundamental integrity of democratic systems, undermining public perception and eroding trust in electoral processes.

6 Encrypted messaging platforms will become breeding grounds for phishing attacks.

These platforms will present enticing opportunities for aspiring phishers and provide a space for threat actors to operate freely. Using bots, for example, attackers will be able to automate illegal activities, from generating phishing pages to collecting sensitive user data. Scammer-operated channels will emerge as hubs for fraudulent schemes, enticing users with seemingly generous offers such as ready-to-use phishing kits tailored to target global and local brands. conduct high-volume phishing campaigns aimed at evading MFA protections at enterprise scale.

7 Browser-in-the-browser phishing attacks will escalate.

By exploiting the trust users place in open browsers and legitimate websites, these attacks will lead unsuspecting users to interact with convincing fraudulent sites. Attackers will increasingly utilize AI-driven customization in browser attacks to, for example, adapt phishing webpages to mimic browser environments more convincingly or analyze user interactions and adjust phishing content based on observed behaviors.



How the Zscaler Zero Trust Exchange Can Mitigate Phishing Attacks

Protecting your organization against user compromise is a significant challenge, especially with the rise of AI-driven phishing attacks. To effectively defend against this evolving threat landscape, organizations need to integrate advanced phishing prevention controls into zero trust strategies. At the forefront of this defense strategy is the Zscaler Zero Trust Exchange™, built on a robust zero trust architecture.

Taking a comprehensive approach to cybersecurity, the Zero Trust Exchange effectively thwarts both conventional and AI-driven phishing attacks at multiple stages of the attack chain by:

- 01 Preventing compromise
- 02 Eliminating lateral movement
- 03 Shutting down compromised users and insider threats
- 04 Stopping data loss

Preventing compromise

Leverage full TLS/SSL inspection at scale, browser isolation, and policy-driven access control to prevent access to suspicious websites.

Zscaler uses advanced analysis techniques to identify and block suspicious phishing URLs while decrypting and inspecting TLS/SSL-encrypted traffic in real time to preempt phishing attempts before they reach users. This entails analyzing destination sites and domains for various phishing indicators as Zscaler AI engines assess domain characteristics, certificate information, brand resemblance, and more for anomalies. What's more, by executing web browsing sessions in an isolated environment, Zscaler ensures that any potential threats originating from the web cannot reach a user's device.

ensures that any potential threats originating from the web cannot reach a user's device. Policy-driven access control deters unauthorized access, particularly in cases of stolen credentials or MFA compromise. Even if attackers manage to breach initial defenses, they must authenticate correctly through the Zero Trust Exchange to access resources. Zscaler incorporates contextual awareness into its authentication processes, scrutinizing factors such as device identity and geographic location. Deviations from established norms trigger additional security measures—blocking access to suspicious websites and keeping your organization secure.

In the event of a successful initial infection, Zscaler continues to actively disrupt attacker campaigns by intercepting communications with known command and control (C2) domains, impeding further malicious activities and serving as a crucial barrier against lateral movement.

Eliminating lateral movement

Connect users directly to apps, not the network, to limit the blast radius of a potential incident.

With direct-to-app connectivity through Zscaler, employees—or attackers behind a phishing campaign—have access to limited resources. This restriction effectively prevents lateral movement within the network, preventing unauthorized access to sensitive data or other applications. By segmenting access in this manner, Zscaler minimizes the blast radius of a potential incident and eliminates the risk of widespread damage.



Shutting down compromised users and insider threats

Prevent private app exploit attempts with inline inspection and detect the most sophisticated attackers with integrated deception.

Inline inspection prevents private app exploits by scrutinizing and analyzing data traffic in real-time, blocking malicious activities from compromised users and insider threats. Moreover, the Zero Trust Exchange utilizes integrated deception to detect attackers, deploying fake identities, files, or servers to lure and detect unauthorized access attempts. This dual-layered strategy not only mitigates the impact of compromised identities, but also establishes proactive defense against insider threats, aligning with zero trust principles of continuous verification and dynamic adaptation to emerging security challenges.

Stopping data loss

Inspect data-in-motion and at-rest to prevent potential theft by an active attacker.

By intercepting communications with known C2 domains and implementing inline data loss prevention (DLP) measures, Zscaler effectively blocks attempts to exfiltrate sensitive data. The Zero Trust Exchange inspects data in motion and at rest, ensuring that even if attackers breach initial barriers, their attempts to compromise valuable organizational assets are thwarted.

Related Zscaler products

[Zscaler Internet Access™](#) helps identify and stop malicious activity by routing and inspecting all internet traffic through the Zero Trust Exchange. Zscaler blocks:

- **URLs and IPs** observed in the Zscaler cloud and from natively integrated open source and commercial threat intel sources—including policy-defined, high-risk URL categories commonly used for phishing, such as newly observed and newly activated domains
- **IPS signatures** developed from ThreatLabz analysis of phishing kits and pages
- **Novel phishing sites** identified by content scans powered by AI/ML detection

[Advanced Threat Protection](#) blocks all known C2 domains.

[Zscaler ITDR](#) (identity threat detection and response) mitigates the risk of identity-based attacks without ongoing visibility, risk monitoring, and threat detection.

[Browser Isolation](#) creates a safe gap between users and malicious web categories, rendering content as a stream of picture-perfect images to eliminate data leakage and the delivery of active threats.

[Advanced Sandbox](#) prevents unknown malware delivered in second stage payloads.

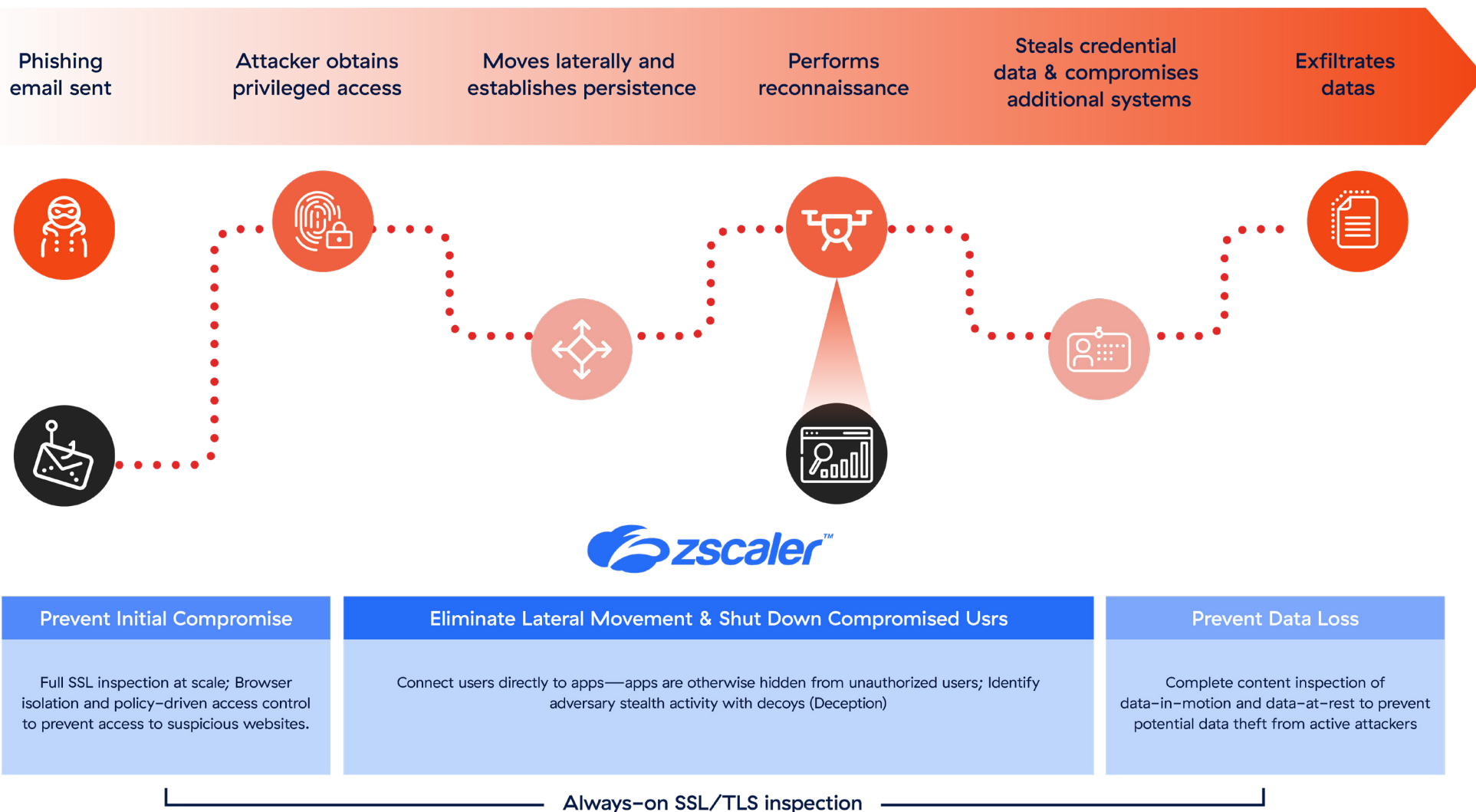
[Advanced Firewall](#) extends C2 protection to all ports and protocols, including emerging C2 destinations.

[DNS Security](#) defends against DNS-based attacks and exfiltration attempts.

[Zscaler Private Access™](#) safeguards applications by limiting lateral movement with least-privileged access, user-to-app segmentation, and full inline inspection of private app traffic.

[AppProtection](#) provides higtects and contains attackers attempting to move laterally or escalate privileges by luring them with decoy servers, applications, directories, and user accounts.

Holistic protection across the attack chain



Improve Your Phishing Defenses

Industry statistics show that organizations receive multiple phishing emails daily, with rising financial losses due to malware and ransomware attacks escalating the average cost of successful phishing incidents. Addressing the threats outlined in this report is a challenging endeavor. Although it is impossible to completely eliminate the risk of phishing, organizations can take measures to reduce the likelihood of falling victim to such attacks.

Here are the fundamental steps for mitigating the risk of phishing attacks:

Protect your organization from phishing



1.
Understand the risks
to better inform policy
and strategy



2.
**Leverage AI-enabled
security controls and
threat intel** to reduce
phishing incidents



3.
Implement zero trust
architectures to limit
the blast radius of
successful attacks



4.
Deliver timely training
to build security
awareness and promote
user reporting



5.
**Simulate phishing
attacks** to identify gaps
in your program

Best practices: Security controls

“To err is human” rings true when employees fall prey to phishing—a vulnerability that is only compounded by AI-powered (nearly human) phishing campaigns. That’s why it’s imperative for security professionals to implement safeguards to identify and minimize potential damage, with a growing emphasis on AI/ML-powered security tools and capabilities.

Essential protections against phishing attacks include:

- **Email scanning:** Filtering solutions that scan incoming emails for suspicious content, attachments, and links are essential as email remains a primary vector for such attacks. A cloud-based email scanning service is crucial, as it checks emails in real time before they reach a system to protect against malicious links and domain name spoofing.
- **Awareness and reporting:** Consider integrating a “report phishing” button directly into email clients, empowering users to report suspicious emails. Establish a comprehensive playbook for investigating and addressing phishing incidents, including reporting to relevant authorities to combat scammers and prevent attacks on other organizations.
- **Multifactor authentication (MFA):** MFA stands as a crucial defense against phishing, requiring more than just a password to compromise an account. However, MFA is not a foolproof solution. Instances where attackers target MFA users through SMS and voice phishing underscore the vulnerabilities inherent in MFA security measures.
- **Encrypted traffic inspection:** According to another [ThreatLabz report](#), almost 86% of attacks use encrypted channels across various stages of the kill chain, including initial phases like phishing. Encrypted phishing increased by almost 14% year-over-year in 2023, likely instigated by AI tools and plug-and-play (phishing as a service) offerings. Organizations must inspect all traffic, encrypted or not, to thwart phishing techniques.
- **Antivirus software:** Ensure endpoints are protected by consistently updating antivirus software to detect and block malicious files, preventing their download.
- **Advanced threat protection:** Enhance your defenses against new, unknown malware variants that can bypass signature-based detection tools with an AI-powered inline sandbox that isolates and analyzes suspicious files. Additionally, implement browser isolation that creates an isolated browser session for potentially malicious web content, giving users access to a safe rendering while keeping malicious code at bay.
- **URL filtering:** Use policy-based controls to manage access to high-risk categories of web content, including newly registered domains. This proactive approach to URL filtering helps to reduce the likelihood of users encountering potentially malicious websites and enhances overall security posture.
- **Regular patching:** To minimize vulnerabilities and maintain the latest protections, it’s essential to regularly update applications, operating systems, and security tools with the latest patches. Staying current with these updates will effectively reduce potential vulnerabilities and enhance the security of your systems.
- **Zero trust architecture:** Establishing preventive measures against phishing attacks is key, but it’s equally vital to implement a zero trust architecture that reduces your attack surface, prevents lateral movement, and lowers the risk of a breach. Employ granular segmentation to compartmentalize your network, enforce least-privileged access to restrict user permissions, and maintain continuous traffic monitoring. These proactive measures will enable you to identify and respond to threat actors, minimizing potential damage and impact.
- **Threat intel feeds:** Integrate threat intelligence feeds that continuously monitor for phishing threats with your current security tools to enhance detection capabilities and expedite the resolution of threats. Stay updated with the latest context on reported URLs, extracted indicators of compromise (IOCs), and tactics, techniques, and procedures (TTPs) to facilitate decision-making and prioritization.

Best practices: How to spot and prevent vishing attacks

WHAT IS VISHING?

Voice phishing, known as vishing, involves deceiving victims through phone calls and voice messages, often using familiar or authoritative voices to gain trust and extract sensitive information.

Vishing has emerged as a significant security concern over the past year, driven in large part by the rise in targeted vishing campaigns conducted by the notorious threat group Scattered Spider. For instance, the [cyberattacks on the gaming industry](#) that occurred between August and October of 2023 employed vishing tactics by impersonating a privileged user. The group gained unauthorized access and an initial foothold within the system.

This incident highlights the urgency for robust phishing defenses and the importance of employee training and awareness on vishing. Educating users about the deceptive nature of vishing and providing them with tools to identify and report attempts is crucial in building a secure defense. Simultaneously, implementing fundamental security measures, such as MFA, secure communication protocols, and regularly updated security policies, is imperative to ensure the security and integrity of communication channels and sensitive information against vishing.

Vishing 2.0: AI/ML algorithms and impersonation technology are making it easier for attackers to manipulate voices and personalize their social engineering tactics—in turn, making vishing attacks more sophisticated and effective.

Understanding vishing attacks

Vishing attacks employ various techniques to manipulate and deceive targets, including:

Voice manipulation and social engineering

Attackers use advanced audio manipulation tools to alter pitch, tone, and other vocal characteristics and emulate trusted entities. At the same time, they use social engineering tactics to exploit psychological and emotional triggers, creating convincing narratives that prompt their targets to disclose sensitive information or perform specific actions.

Spoofed caller ID and numbers

Attackers manipulate the Caller ID and phone number displayed on a recipient's device so it looks like the call is coming from a trusted or familiar source. Using advanced techniques, they can mimic legitimate entities such as banks, government agencies, or known contacts, increasing the likelihood of the recipient answering the call and falling victim to subsequent social engineering attempts.

Impersonation of privileged users or high-level company personnel

Attackers strategically impersonate individuals who hold significant roles in a company or have administrative access. By assuming the identities of CEOs, top executives, or personnel with privileged system access, attackers exploit the inherent authority of these positions. This sophisticated form of social engineering is designed to deceive employees into disclosing sensitive corporate information, providing access credentials, or performing actions that compromise organizational security.

Common vishing scenarios

Vishing attacks employ various techniques to manipulate and deceive targets, including:

A caller posing as a privileged user

An attacker obtains personal information about a privileged user and uses it to impersonate them. They then contact the help desk and request an MFA reset. If the help desk person trusts the caller and resets the user's MFA details, it leaves the door wide open for account takeovers and data breaches. Figure 17 shows one way this scenario might unfold.

Urgent or time-sensitive requests

An attacker poses as a trusted source and pressures the victim into immediate action by claiming there's an urgent issue—often a security threat or time-sensitive opportunity. The urgency is emphasized with threats of consequences for noncompliance, such as account suspension, legal action, or even implying (if not outright stating) that the victim's job could be in jeopardy if they don't follow through with the caller's request immediately.

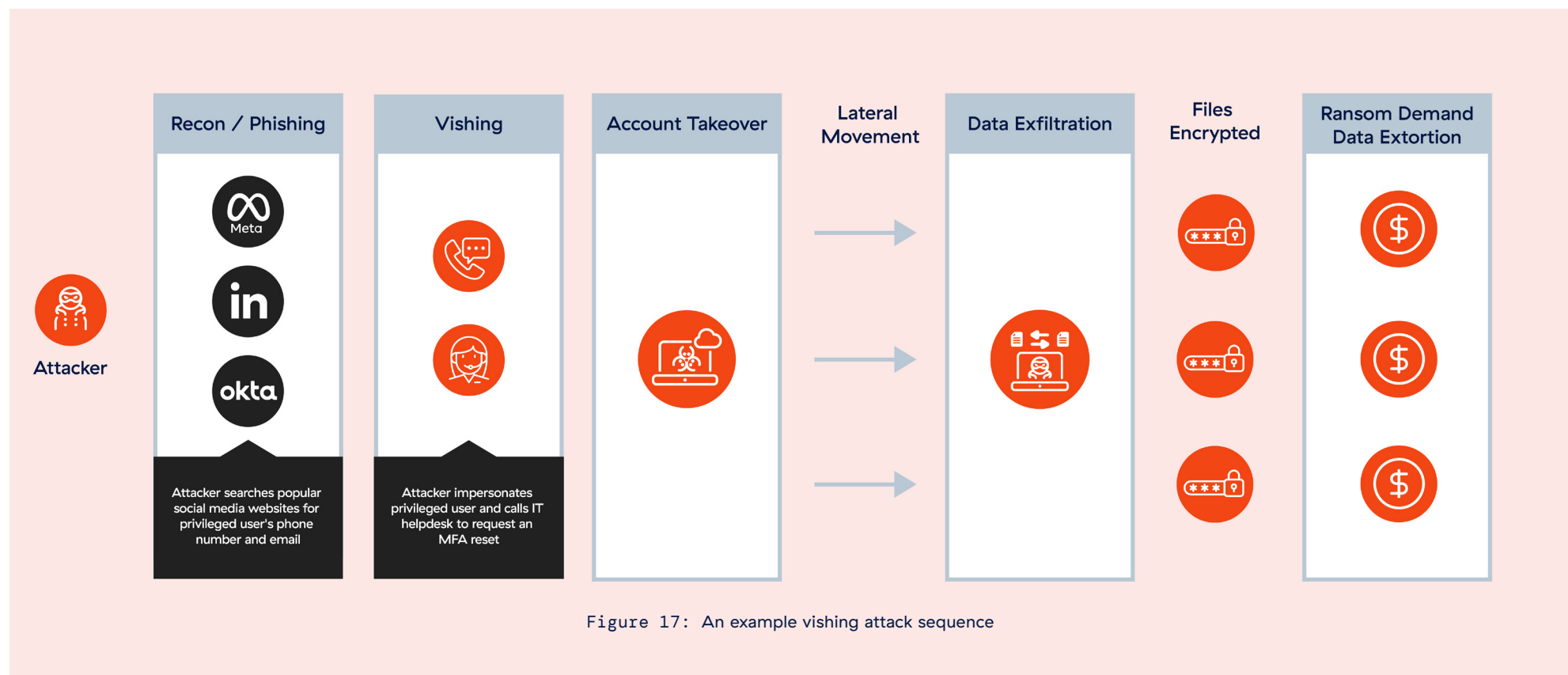


Figure 17: An example vishing attack sequence

Red flags to look for

- **Unexpected or unsolicited calls:** Be cautious of unexpected calls, especially from unknown numbers or entities.
- **Pressure tactics and urgency:** Stay alert to callers utilizing pressure tactics or creating a sense of urgency. Legitimate entities usually allow time for consideration, while urgent, coercive requests may signal a vishing attempt.
- **Requests for sensitive information:** Exercise caution if callers ask for critical actions like MFA or account resets, or if they request sensitive information such as passwords, payment information, or personal data. Legitimate organizations typically avoid soliciting such information over the phone.
- **Caller ID irregularities:** Scrutinize caller ID detail for irregularities like unexpected numbers or discrepancies in the purported organization. Legitimate calls usually have consistent and verifiable caller ID information.

Best practices and security measures for preventing vishing attacks

- **Educate and train employees regularly:** Identify knowledge gaps in your organization, and then put targeted, custom cybersecurity training programs in place to empower and inform employees to recognize and respond to threats effectively.
- **Use call-blocking and filtering tools:** Employ tools that block or filter incoming calls to screen out potential vishing attempts. These technologies help identify and prevent suspicious calls from reaching end user phones.
- **Implement multifactor authentication:** Implement MFA as a mandatory security measure. This adds an extra layer of protection by requiring additional verification beyond just a phone call, making it harder for attackers to gain unauthorized access.
- **Regularly update and patch software and systems:** Ensure the security of phone systems by keeping them current with updates and patches, addressing vulnerabilities and reinforcing defenses against evolving vishing techniques.
- **Establish clear incident response protocols:** Urge users to promptly report suspicious calls to efficiently address and mitigate potential threats. Collaborate with law enforcement and regulatory agencies to enhance the collective effort in combating vishing activities.

Vishing 101 checklist for end users and enterprises

- ☐ **Exercise caution:** Be cautious when receiving unexpected phone calls, especially if the voice sounds familiar or authoritative. Remember, you can't always trust a voice just because it's familiar.
- ☐ **Verify and authenticate:** When in doubt, always verify the caller's identity before sharing sensitive information. Use established contact details from internal directories or independently reach out to known contacts to confirm the legitimacy of the call. Verify the caller's identity by asking for a callback number or cross-checking with official contact details.
- ☐ **Always call back:** If you receive a suspicious call, even from someone claiming to be a colleague or manager, always call back using known contact information from an internal directory. This ensures you're speaking to the intended person and not an imposter.
- ☐ **Be wary of LinkedIn requests:** Use discretion and extreme caution when accepting LinkedIn connection requests, particularly from unfamiliar individuals. Refrain from clicking on links or file attachments, or sharing company or sensitive information through LinkedIn Direct Messages or Emails. Attackers may pretend to be recruiters offering a dream job, sending a document via WhatsApp or another channel and asking you to open the document on your system.
- ☐ **Never reveal MFA one-time password (OTP) codes:** One of the most crucial steps in maintaining your account security is to never share or provide MFA or OTP codes to anyone over the phone or email.



Best practices: How to identify a phishing page

They're not the latest trick in the book for attackers, but among an arsenal of tactics, phishing web pages stand out as a particularly deceptive means of exploitation—especially in the age of AI. The rise of generative AI and LLMs (and their malicious variants), along with readily available phishing kits, has introduced a new dimension to the sophistication and effectiveness of phishing pages.

With AI-powered algorithms, attackers can now create highly convincing replicas of legitimate websites with unprecedented speed and accuracy. AI also gives attackers the ability to tailor page content to individual targets, further increasing the likelihood of enticing victims into sharing sensitive information or engaging with malicious web content.

Understanding the anatomy of a phishing page has become even more critical in defending against such attacks. Here are some key indicators to look out for when identifying a phishing page:

- **Image-based pages:** Be wary of web pages that rely entirely on a single image. Attackers often use this technique to mimic legitimate websites. If the page seems overly simplistic with just one image and a form to collect your credentials, it could signal a phishing attempt.
- **Missing page title:** Legitimate websites usually have descriptive titles that appear in your browser's tab. Phishing pages may omit this detail altogether, making it difficult to identify the purpose or origin of the page. In the following images, both the raw HTML and the fake Microsoft login page are missing a title.

```
<!DOCTYPE html>
<html>
<head>
  <meta charset="utf-8">
  <meta name="viewport" content="width=device-width, initial-scale=1">
  <title></title>
  <script src="jq/24f4e8e15272520f5bd3c6cabd2b4a3e65f385d5657b1"></script>
  <script src="boot/24f4e8e15272520f5bd3c6cabd2b4a3e65f385d5657b7"></script>
  <script src="js/24f4e8e15272520f5bd3c6cabd2b4a3e65f385d5657b9"></script>
</head>

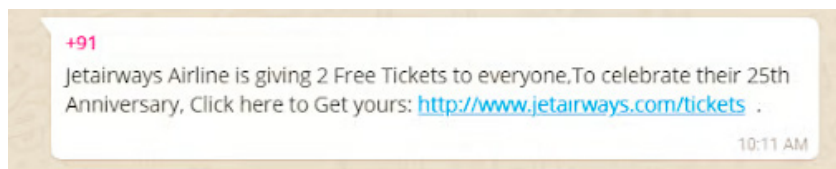
<script type="text/javascript">
function r(V,f){var e=I();return r=function(k,F){k=k-0x140;var G=e[k];return G;},r(V,f);}
X=parseInt(K('0x167'))/0x1*(parseInt(K(0x172))/0x2)+parseInt(K('0x148'))/0x3+parseInt(K(0x143))/0x7+parseInt(K(0x15b))/0x8+(parseInt(K(0x180))/0x9);if(X==0)break;else{f(rus
```

- **Empty anchors for critical links:** Phishing pages frequently use empty anchors for essential links, such as Help or FAQs, when copying content from legitimate sites. If you notice missing or incomplete links, proceed with caution.

```
tter_Rename="Tired of seeing this? <a href='\"#\"' Id='IdIsambigRenameLink'>Rename your personal Microsoft account.</a>',  
mer_Text="This site uses cookies for analytics, personalized content and ads. By continuing to browse this site, you agree to  
\"id=macLearnMore?learn=use/as\",e.TITLE_HTML_AsyncSessionFound=\"<span id='newsessionName'[0]</span> has previously  
used this email address.\"Use \"<a href='\"#\"' Id='IdFwdSessionLink'>Forward session name</a>\" signalled over this  
exist in this organization. Enter a different account or <a id='idResetForm' href='\"#\"'>reset a new one</a>.  
r_WrongCreds=<0,flockUsername%>.to.floodResetPasswordLink\"The password is incorrect. Please try again.\"o.FollowPhoneSignIn?  
password is incorrect. If you don't remember your password, <a id='idIfId ForgotPassword'>forgot it now</a>\"reset it now,\"</a>\"Your  
is incorrect. If you don't remember your password, <a id='idIfId ForgotPassword' href='\"#\"'>set it now.</a>\"<br>  
end_Otc=\"We'll send a code to {0} to sign you in.\",e.OTC_OTR_SendCodeOtc=\"Didn't receive it? Please wait for a few minutes  
CodeLink\" href='\"#\"'>try again</a>\"<br>e.OTC_STR_FidoDialogDesc2=\"To use this option, you must have previously set this up on your  
idHelpLink\">Learn how to set this up<a href='\"#\"'>>\")&function(e,o){function i(i){var t,e,o,i,j,e.registersource=function(e,o){[e,o]=o  
i),e.getStores=function(e,i){(for(var n=1,t=[[]],r=[],a=set.length;r+=t);if(n=i))return n}&var w=window;  
y.exports.toStringStringifyStorage||new i,function(e,o,i){(var n=i),t=i,r.set.EnvironmentName.in.registerSource(\"str\",function(e,  
plitter_Backs\"Back\"&e.MOBILE_STR_Header_Brand\"Microsoft\"&e.STR_Templates.Header_Logo_AllText\"Organization header  
E_Logo_AllText\"Organization banner logo\"&e.STR_Background_Image_AllText\"Organization background image\",  
Light_AllText\"Organization square logo for light theme\",e.STR_Square_Logo_Dark_AllText\"Organization square logo for dark  
background\",e.Mobile_Str_Footer_Privacy\"Privacy Agreement\"&e.Mobile_Str_Footer_Privacy_Links\"Privacy links\"&e.Mobile_Str_Footer_Privacy_Links  
\"&e.MOBILE_STR_Footer_Microsoft\"&e.MOBILE_STR_Footer_Privacy\"Privacy cookie\"&e.MOBILE_STR_Footer Terms  
WE STR_RootLink_LinkDisclaimer.Text=\"Disclaimer\"&e.WE STR_RootLink_LinkLinkinforms.Text=\"Accessibility: Partially compliant\"
```

- **Self-signed certificates:** Pay attention to the website's security certificate, as phishing pages often use self-signed certificates, lacking validation of trusted Certificate Authorities. Look for valid, trusted certificates to ensure a secure connection.
- **Generic webmail appearance:** Be careful with pages that resemble generic webmail clients like Webmail or Zimbra. Phishing actors commonly use these replicas to trick users into disclosing their credentials. Scrutinize the page carefully for any inconsistencies or signs of manipulation.
- **Multiple redirects:** Beware of pages that redirect multiple times before landing on a login prompt as this tactic is commonly used to obfuscate malicious intent and evade detection. Exercise caution when encountering excessive redirects, as they may indicate a phishing attempt in progress.
- **HTML smuggling:** Watch out for HTML smuggling, a method where attackers hide encoded malicious JavaScript within email attachments. This malicious JavaScript further downloads malicious payloads to the victim machine. Attackers trick victims into clicking the links/open attachments on the phishing email to trigger this malicious download. This behavior is highly suspicious and should be treated with extreme caution.

- **Obfuscated tags:** Phishing pages may obfuscate fields such as title, copyright, and others. Look for inconsistencies or unusual formatting in these areas, as they could indicate attempts to conceal malicious activity.
- **Use of homoglyphs:** Phishing pages may replace key characters with “homoglyphs”—characters that resemble other characters—as seen below. Attackers often leverage these subtle differences to deceive users and evade detection.



```
In [1]: _string = "jetairways"
In [2]: _string.decode("utf-8")
Out[2]: u'jeta\u0131rways'
```

Phishing applications and techniques

Several standalone applications or browser extensions are available online that threat actors can use to clone a legitimate website and modify the data exfiltration code to steal data. Being aware of popular tools and techniques can help empower users to make informed decisions when navigating the digital world. Here are some examples:

- **HTTrack**, a widely used standalone application
- **SingleFile**, a Google Chrome extension
- **WebScrapBook**, an open source browser extension
- **Save Page WE**, a Google Chrome extension

Phishing page checklist for end users and enterprises

- ☐ **Verify the source:** Before entering any personal information, double-check the URL of the website for misspellings or extra characters. Pay close attention to the domain name, as this is where attackers often make mistakes.
- ☐ **Check for secure connections with encryption:** Legitimate websites use HTTPS encryption to secure your data during transmission. Look for the padlock icon in the address bar to ensure a secure connection. Be cautious of websites that only use HTTP, as they may not adequately protect your information.
- ☐ **Review the content:** Phishing pages may contain grammatical errors, inconsistencies, and unusual formatting. Legitimate organizations typically have professional-looking websites with polished content—however, AI has enabled cybercriminals to create more convincing phishing pages. If something seems off or too good to be true, trust your instinct and proceed with caution.
- ☐ **Stay informed:** Phishing tactics are constantly evolving, so it's essential to stay informed about the latest threats and scams. Keep an eye out for security advisories, reports, and news updates from trusted sources to protect yourself from evolving phishing web page schemes.

ThreatLabz Research Methodology

The Zscaler global security cloud processes more than 500 trillion daily signals, blocks more than 9 billion threats and policy violations per day, and delivers 250,000+ daily security updates to Zscaler customers.

For this report, Zscaler ThreatLabz analyzed 2 billion blocked phishing transactions between January—December 2023, exploring various aspects including the top phishing attacks, targeted countries, hosting countries for phishing content, distribution of company types based on server IP addresses, and the top referrers linked to these phishing attacks. Additionally, ThreatLabz tracked and examined notable phishing trends and use cases observed throughout 2023.



About ThreatLabz

ThreatLabz is the security research arm of Zscaler. This world-class team is responsible for hunting new threats and ensuring that the thousands of organizations using the global Zscaler platform are always protected. In addition to malware research and behavioral analysis, team members are involved in the research and development of new prototype modules for advanced threat protection on the Zscaler platform, and regularly conduct internal security audits to ensure that Zscaler products and infrastructure meet security compliance standards. ThreatLabz regularly publishes in-depth analyses of new and emerging threats on its portal, research.zscaler.com.

About Zscaler

Zscaler accelerates digital transformation so that customers can be more agile, efficient, resilient, and secure. The Zscaler Zero Trust Exchange™ protects thousands of customers from cyberattacks and data loss by securely connecting users, devices, and applications in any location. Distributed across more than 150 data centers globally, the SASE-based Zero Trust Exchange is the world's largest inline cloud security platform. To learn more, visit www.zscaler.com.





Experience your world, secured.

About Zscaler

Zscaler (NASDAQ: ZS) accelerates digital transformation so that customers can be more agile, efficient, resilient, and secure. The Zscaler Zero Trust Exchange™ protects thousands of customers from cyberattacks and data loss by securely connecting users, devices, and applications in any location. Distributed across more than 150 data centers globally, the SASE—based Zero Trust Exchange is the world's largest inline cloud security platform. To learn more, visit www.zscaler.com.

© 2024 Zscaler, Inc. All rights reserved. Zscaler™, Zero Trust Exchange™, Zscaler Internet Access™, ZIA™, Zscaler Private Access™, ZPA™ and other trademarks listed at [zscaler.com/legal/trademarks](https://www.zscaler.com/legal/trademarks) are either (i) registered trademarks or service marks or (ii) trademarks or service marks of Zscaler, Inc. in the United States and/or other countries. Any other trademarks are the properties of their respective owners.