



Cloudflare Security Brief

Threats against people, apps, and infrastructure



Table of contents

[Table of contents](#) >

3	Executive summary	10	API exposure
4	Key findings	11	LLM risks
5	DDoS attacks	12	Zero Trust adoption
8	Vulnerability exploitation	13	Recommendations
9	Phishing threats	21	Appendix

Executive summary

< [Table of contents](#) >

From Q2 2023 – Q1 2024, Cloudflare observed threat activity moving bigger and faster. This included record-breaking DDoS attacks, vulnerability weaponization just 22 minutes after a proof-of-concept was published, and hundreds of millions of phishing attacks.

Threat actors weren't the only ones to evolve; organizations changed too. Increased investments in cloud, applications, and, especially, artificial intelligence (AI) fundamentally changed how organizations need to think about defining and defending their attack surface. Maintaining security in turbulent times will require both knowledge and action.

The Cloudflare Security Brief, created by observations, experience, and data from Cloudflare's global cloud network, provides security leaders with insights into threats including DDoS attacks, vulnerability exploitation, and phishing. It also spotlights emerging risks that security leaders need to prepare for: APIs and AI. Lastly, this brief examines how organizations are implementing Zero Trust to keep their people, applications, and infrastructure safe.

Security leaders can use the insights and recommendations provided to prioritize their efforts on effective controls that ensure resiliency in the coming years.



Methodology

Operating one of the largest global cloud networks across 310 cities, Cloudflare protects approximately 20% of the web. This unique vantage point across the Internet provides extensive visibility into threat activity, enabling Cloudflare to stop an average of 182 billion cyber threats each day.

The findings in this brief are primarily based on aggregated traffic patterns observed across Cloudflare's global network between April 1, 2023 and March 31, 2024.

Key findings

[Table of contents](#)



DDoS attacks grew in size and complexity

In 2023 Cloudflare detected multiple record-breaking DDoS campaigns, with hyper-volumetric attacks peaking at 201 million requests per second.



Organizations must prepare for rapid weaponization

Cloudflare observed attempted exploitation of a new vulnerability just 22 minutes after a proof-of-concept was published.¹



Phishing attacks show no signs of slowing down

After nearly 30 years of attacks, phishing remains a top threat. 48% of phishing attacks try to get users to click on a deceptive link.²



Organizations have larger API attack surfaces than they think

APIs are a prime target for attackers targeting data. Cloudflare found 33% more API endpoints through machine learning than what customers self-reported, on median.³



Generative AI investments create new attack targets

Custom large language models (LLMs) are a veritable goldmine for attackers. Organizations that invest in LLMs must understand the three types of LLMs and the risks they create.



VPN replacement is frequent in Zero Trust adoption

75% of cyber security & IT professionals implementing Zero Trust report that they have replaced or plan to move away from VPNs for all employees.⁴

DDoS grew in size and complexity

[Distributed denial of service](#) (DDoS) attacks are evolving. No longer a low-level annoyance, Cloudflare observes DDoS used as a tool for external threat actors to interfere with government functions and business continuity. For example, in November 2023, OpenAI’s ChatGPT faced outages due to DDoS attacks, impacting millions of users.⁵

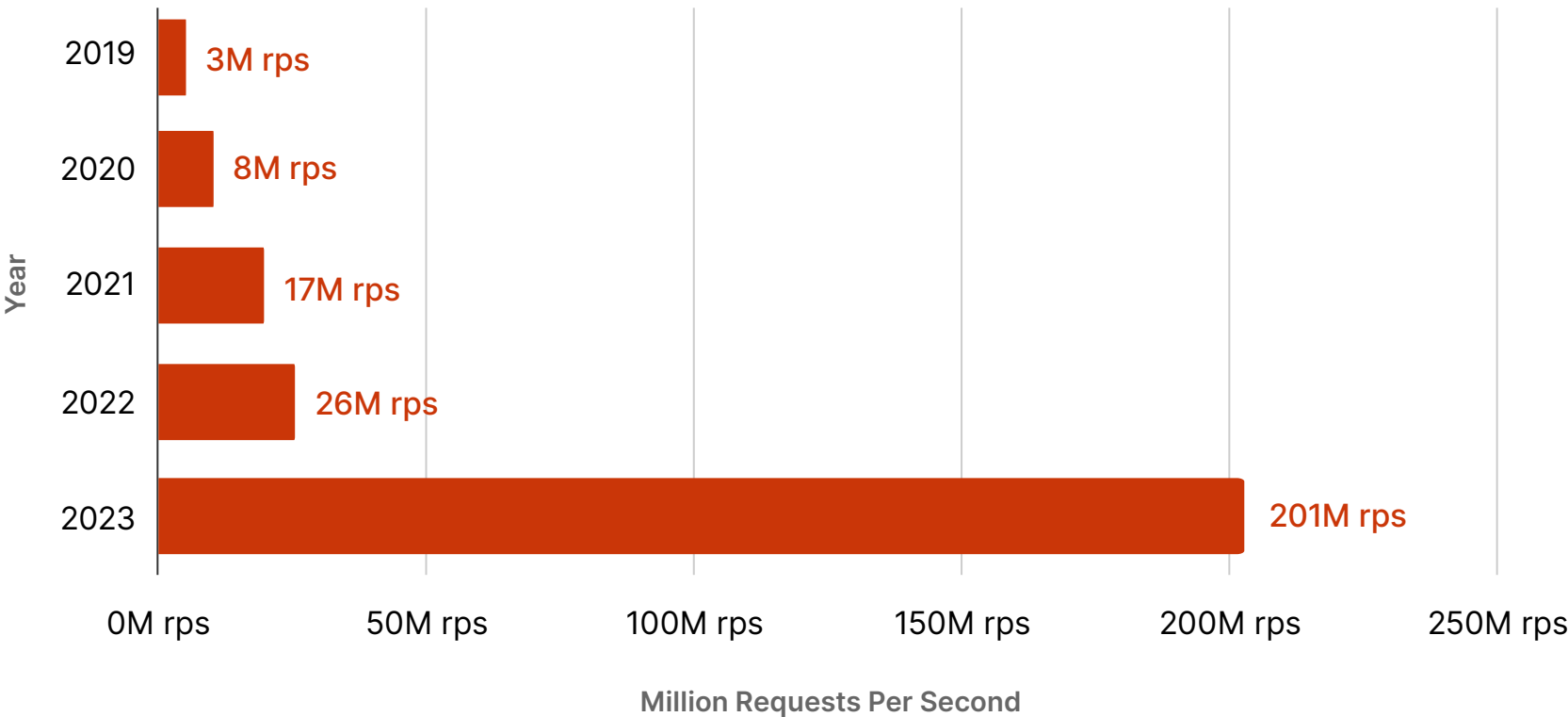
DDoS remains a top attack method for threat actors, and in 2023 those attacks reached new heights. In Q3 2023, Cloudflare [discovered](#) a persistent and deliberately engineered campaign of thousands of hyper-volumetric DDoS attacks. Our systems detected and mitigated the largest attack we’ve ever seen — 201 million requests per second (rps) — which was almost eight times larger than our previous record in 2022 of 26 million rps.

The attacks exploited a vulnerability in HTTP/2 protocol called Rapid Reset which leverages HTTP/2’s stream cancellation feature by sending a request and immediately canceling it over and over. It is likely that the threat actors deliberately tested the exploit against Cloudflare’s network because of its size – very few organizations can absorb DDoS attacks at this scale.

One crucial thing to note is that this technique dispels the myth that a large number of bots are necessary to participate in an attack. It used a modestly-sized botnet, consisting of roughly 20,000 hosts. Cloudflare regularly detects botnets that are orders of magnitude larger than this.

HTTP/2 vulnerabilities continue to be discovered, such as the April 2024 discovery of CONTINUATION Flood, which has the potential to crash a web server with a single TCP connection. Cloudflare expects that threat actors will continue to exploit vulnerabilities in Internet protocols to scale DDoS attacks and wreak havoc.

Largest HTTP DDoS attacks



Most attacked industries

Segmenting application-layer DDoS attacks by industry, we see the five most targeted industries were gaming/gambling, IT/Internet, cryptocurrency, computer software, and marketing/advertising. Gaming/gambling frequently [tops our list](#) for most attacked industry due people seeking to gain a competitive advantage in gaming.

However, DDoS attacks are widespread – these are not the only industries that should prioritize mitigations. Normalizing DDoS attacks by an industry’s total volume of traffic allows us to identify which industries have outsized targeting relative to their smaller size. Biotechnology, transportation, cryptocurrency, events service, and chemicals are the most attacked industries relative to their total traffic volume.

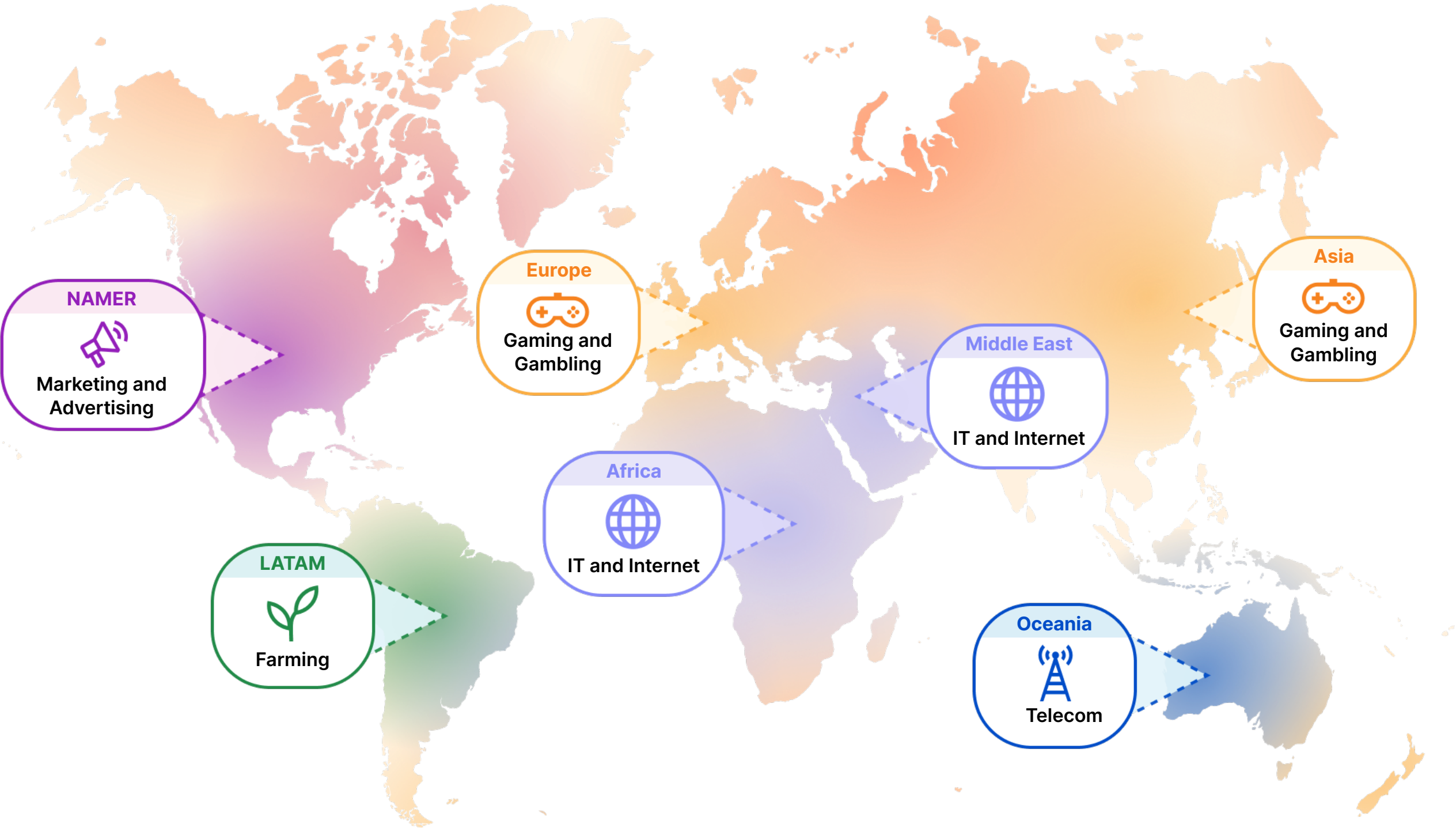
Application-layer DDoS attacks by top industries
Divided by volume of worldwide DDoS traffic

1	Gaming and Gambling	6	Telecommunications
2	IT and Internet	7	Retail
3	Cryptocurrency	8	Adult entertainment
4	Computer software	9	Banking, Financial services, and Insurance
5	Marketing and Advertising		
		10	Manufacturing

Application-layer DDoS attacks by top industries
Divided by traffic of each industry

1	Biotechnology	6	Accounting
2	Transportation	7	Wholesale
3	Cryptocurrency	8	Farming
4	Events services	9	Gaming and Gambling
5	Chemicals	10	Environmental Services

Most attacked industries by region



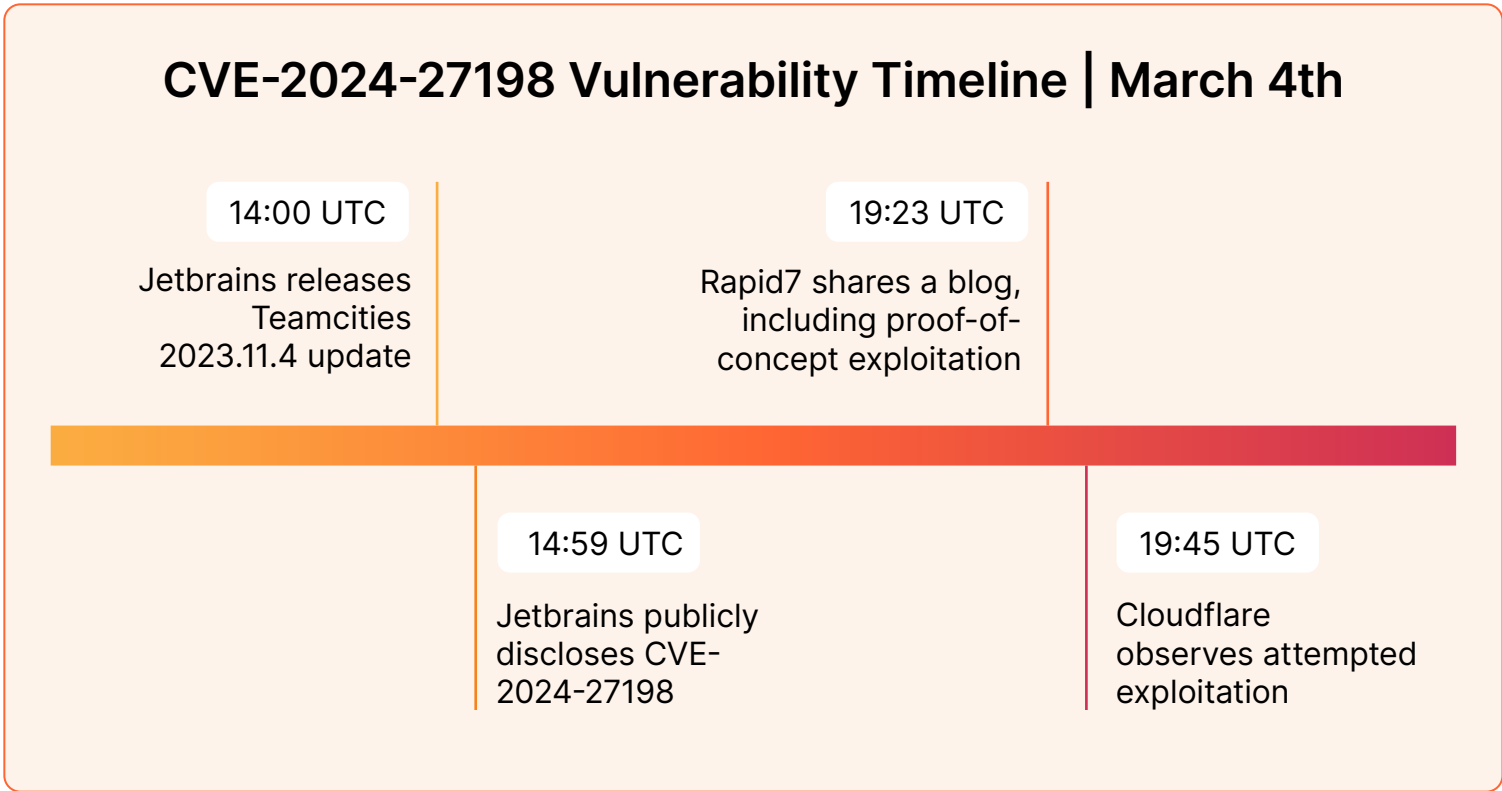
Threat actors exploit vulnerabilities faster than organizations can patch them

Thanks to advancements in cloud and AI, organizations build code faster than ever. Unfortunately, Cloudflare also sees threat actors weaponizing vulnerabilities at breakneck speed.

Take [CVE-2024-27198](#), a vulnerability that allows for a complete compromise of a vulnerable TeamCity server, including unauthenticated remote code execution. **Cloudflare observed attempted exploitation just 22 minutes after proof-of-concept code was published.**

Meanwhile, vulnerabilities are discovered at an overwhelming pace. There were more than 5000 critical vulnerabilities disclosed in 2023⁶, yet the mean time to remediate a critical severity web application vulnerability is 35 days.⁷ Organizations can't keep up pace with patching. Combine that with the fact that there were nearly 100 [zero-day](#) vulnerabilities in 2023, up 50% from 2022,⁸ and it's no surprise that vulnerability management is such a struggle for many organizations.

What does this add up to? An overwhelming number of emergency response moments for organizations when new zero-days and critical vulnerabilities are disclosed.



Stopping zero-days before day zero

It's not all bad news. Advances in machine learning are providing critical capabilities to zero-day prevention. Using machine learning, it is possible to stop previously unknown zero-day attacks from the very moment they're first attempted, as Cloudflare displayed when blocking the [Ivanti Connect Secure VPN vulnerability](#).

Phishing attacks show no signs of slowing down

Phishing is an attempt to get a user to take an action such as downloading malware, clicking links to harvest credentials, or transferring money. Phishing has plagued organizations for decades and shows no signs of slowing down. In fact, 9 out of 10 successful cyber attacks start with phishing.

Over a 12-month period, Cloudflare processed more than 15 billion emails, providing visibility into popular phishing tactics and trends.

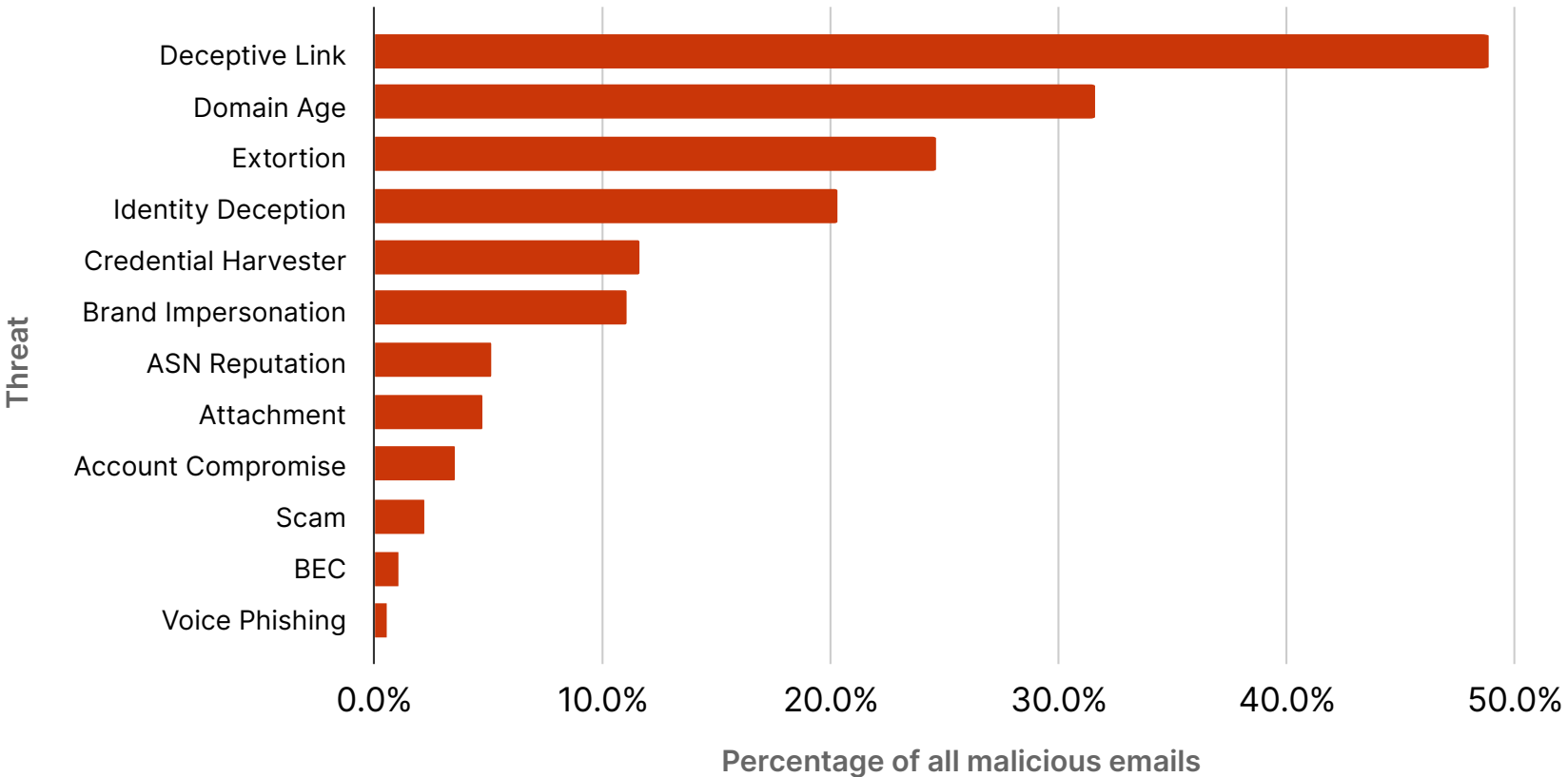
Deceptive links were the #1 phishing tactic, included in 48% of all malicious emails. Links remain popular because of the various ways in which links can be masked to retain authenticity, including link shorteners for deferred attacks or [QR codes](#) that pivot users to a less secure mobile experience.

Business email compromise (BEC) represents a much smaller percentage of phishing volume at 1%, but its financial impact cannot be ignored. In 2023, the Federal Bureau of Investigations (FBI) reported BEC losses of more than \$50 billion in a 1 year period.



9 out of 10
successful cyber attacks
start with phishing

Top threat categories in malicious emails



Detailed descriptions of the above-noted categories can be found in the Appendix.

AI-driven phishing attacks

Recently, AI has been much hyped as an offensive boon to increase personalization and speed in phishing emails. While Cloudflare’s email security analysts observed an increase in LLM-written phishing emails, AI’s impact on successful phishing attempts against enterprises is relatively minimal.

Advancements in machine learning and artificial intelligence help defenders too. Employing a range of analytic techniques across technical email structure, text sentiment, and historic communication patterns can [reliably stop LLM-generated phishing attacks](#).

The growing risks and rewards of APIs

APIs help organizations integrate and operate their environments to fuel competitive advantages — with greater business intelligence, swifter cloud deployments, and more. And they're heavily used: today, APIs outpace other Internet traffic, **comprising more than half (58%) of the dynamic Internet traffic** processed by Cloudflare last year.⁹

However, as the API economy grows, so do the problems of loss of control and complexity with API development, management, and security.

APIs are increasingly complex to manage and protect against abuse. Unprotected APIs can lead to data exposure, unpatched vulnerabilities, data compliance violations, lateral movement, and other threats.

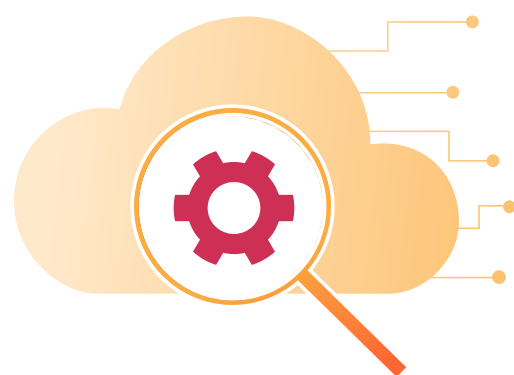
APIs are also a cornerstone of AI implementations, as the primary mechanism for interacting with generative AI models. Protecting AI models will require a strong handle on the existence, permissions, and usage of APIs across an organization.

Unfortunately, many organizations lack accurate API inventories. Cloudflare found 33% more API endpoints through machine learning-based discovery, compared to what organizations self-reported.

APIs that have not been managed or secured by the organization using it — also known as 'Shadow' APIs, are often introduced by developers or individual users to run specific business functions.

While they are not inherently malicious, shadow APIs are unprotected attack surfaces that introduce new risks.

Organizations cannot properly defend what they cannot see. And those that implement API security without an accurate picture of their API landscape can also unintentionally block legitimate traffic



Cloudflare found
33% more API endpoints
than what organizations self-reported



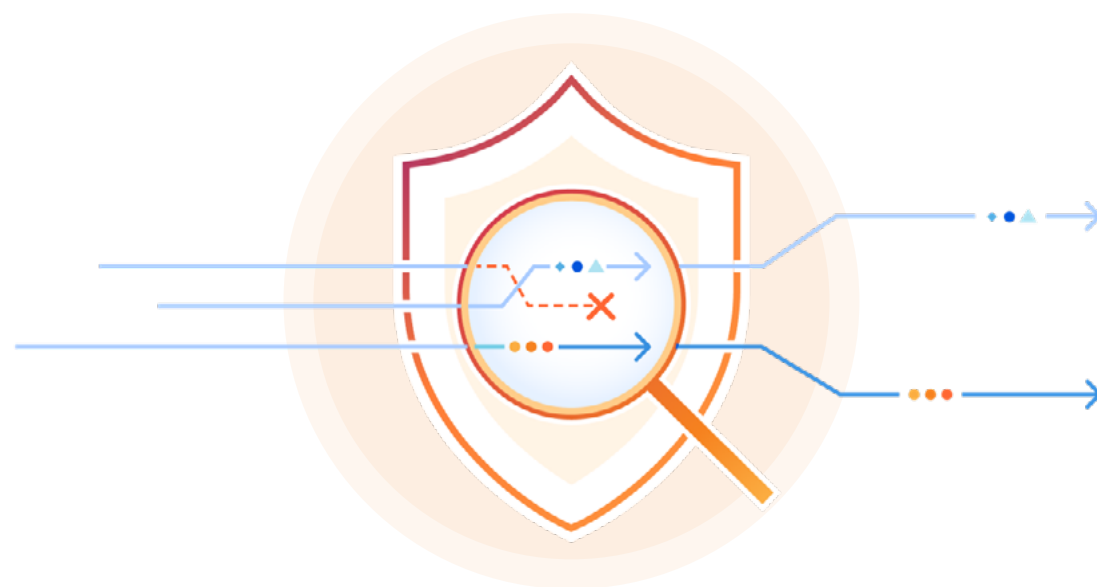
The security implications of LLM adoption

Organizations are investing heavily in [generative AI](#), with high expectations that it will drive revenue and increase efficiency. IDC forecasts spending on GenAI solutions will reach \$143 billion in 2027.¹⁰

But investments in [LLMs](#) are creating highly attractive targets for threat actors. Internal LLMs are likely to have wide-ranging access to sensitive information and intellectual property, and, since the models run on high-powered machines, financially motivated actors may target them for their computing power.

This isn't hypothetical – we've already seen attacks on LLMs. Oligo discovered active exploitation of a widely used open-source AI framework that granted threat actors access to production AI workloads (meaning they could steal or even tamper with models and data sets).¹¹ The actors stole production database credentials, data, passwords, and cloud access privileges, and even installed crypto mining malware.

The level of risk exposure AI creates for an organization will vary depending on how it is used, but now is the time for every security organization to be assessing LLM risk.



Understanding AI Risk The 3 types of LLMs



Internal LLMs

Custom models trained on internal data to assist employees and boost productivity.

Example: an AI co-pilot trained on sales data and customer interactions used to generate tailored proposals.

Key risk: access to sensitive data and intellectual property



Product LLMs

Part of a product or service offered to customers.

Example: a customer support chatbot built to interact with company resources

Key risk: reputational risk



Public LLMs

LLM accessed outside the boundaries of a corporation, often for free.

Examples: GPT from OpenAI or Claude from Anthropic.

Key risk: sensitive data leakage

Increased network attacks demonstrate need for Zero Trust

In 2023 there was a significant [uptick](#) in attacks and zero-days against network and security products such as VPNs, firewalls, and load balancers. The trend [continues](#) in 2024, with critical zero-days for Ivanti's Connect Secure VPN as well as Palo Alto Networks' firewall-based GlobalProtect VPN product.

Supply chain attacks against enterprise public-facing infrastructure are moving from likely to inevitable. A [Zero Trust](#) approach to security is increasingly essential to reduce the exploitable attack surface and add defense-in-depth.

ESG asked 200 cyber security and IT professionals currently adopting Zero Trust about their adoption path.⁴ The two highest ranked use cases for initial implementation were enforcing Zero Trust application access (ZTAA) policies for SaaS applications and deploying Zero Trust Network Access ([ZTNA](#)) for private applications.

Enterprise resource planning (ERP) and communication and collaboration tools were the two highest ranked applications that organizations have secured or intend to secure in the initial deployment of their Zero Trust journey.

Moving from VPNs to ZTNA

With VPN in particular being all-too-frequently compromised, organizations are increasingly evaluating alternate remote access options. 98% said that remote access solutions that directly connect users to applications (rather than the broader network) were important. **And 75% of cyber security & IT professionals currently using ZTNA report that they have replaced or plan to move away from VPNs for all employees.**

One barrier to adopting ZTNA tools is installation of an endpoint agent. Installation can be time-consuming and costly, hindering an organization's ability to quickly and successfully deploy ZTNA.

85% of cyber security & IT professionals agreed that agentless ZTNA tools simplified the deployment process, reducing the administrative burden and potential points of failure associated with agent-based solutions.



75% of cyber security & IT professionals

have replaced or plan to move away from VPNs for all employees

Recommendations



1

Implement modern DDoS mitigation best practices

Best practice	Action
Deploy threat intelligence and in-line, automated DDoS mitigation solutions	Manual scrubbing centers do not scale with modern, high-volume attacks. Use multiple detection techniques to optimize security posture: 1. Dynamic stateless fingerprinting 2. Machine learning-based classification 3. Anomalous traffic detection 4. Traffic profiling and stateful mitigation 5. Threat intelligence on current DDoS activity and trends
Build a disaster recovery scenario for a continuous, long-lasting DDoS attack	Start by identifying critical infrastructure vulnerable to DDoS attacks and the impact of their downtime. Be sure to include commonly forgotten, but still vital aspects of network stacks, such as DNS servers and VPN endpoints.
Update network, DNS and application infrastructure to be more resilient for your traffic profile	Ensure your DDoS mitigation capacity is large enough to handle twice the largest attacks on record and twice the max rates of your legitimate traffic. Ensure your security vendor can mitigate the latest network and application layer protocol vulnerabilities. Offload DNS traffic to compliant and secured cloud platforms with traffic routed through edge networks closest to the user.
Improve network and application performance to avoid bottlenecks	Leverage a digital waiting room to ensure real users and visitors are gracefully informed of the waiting period without overwhelming application servers. Optimize caching, manage loads better with a content delivery network (CDN) and cloud based loading balancing solutions.
Use a positive security model: Ensure that desired traffic gets in reliably	Keep business critical protocols, IPs, ASNs, ports and user-agents open to clean traffic. Use schema validation and an API gateway for API traffic.

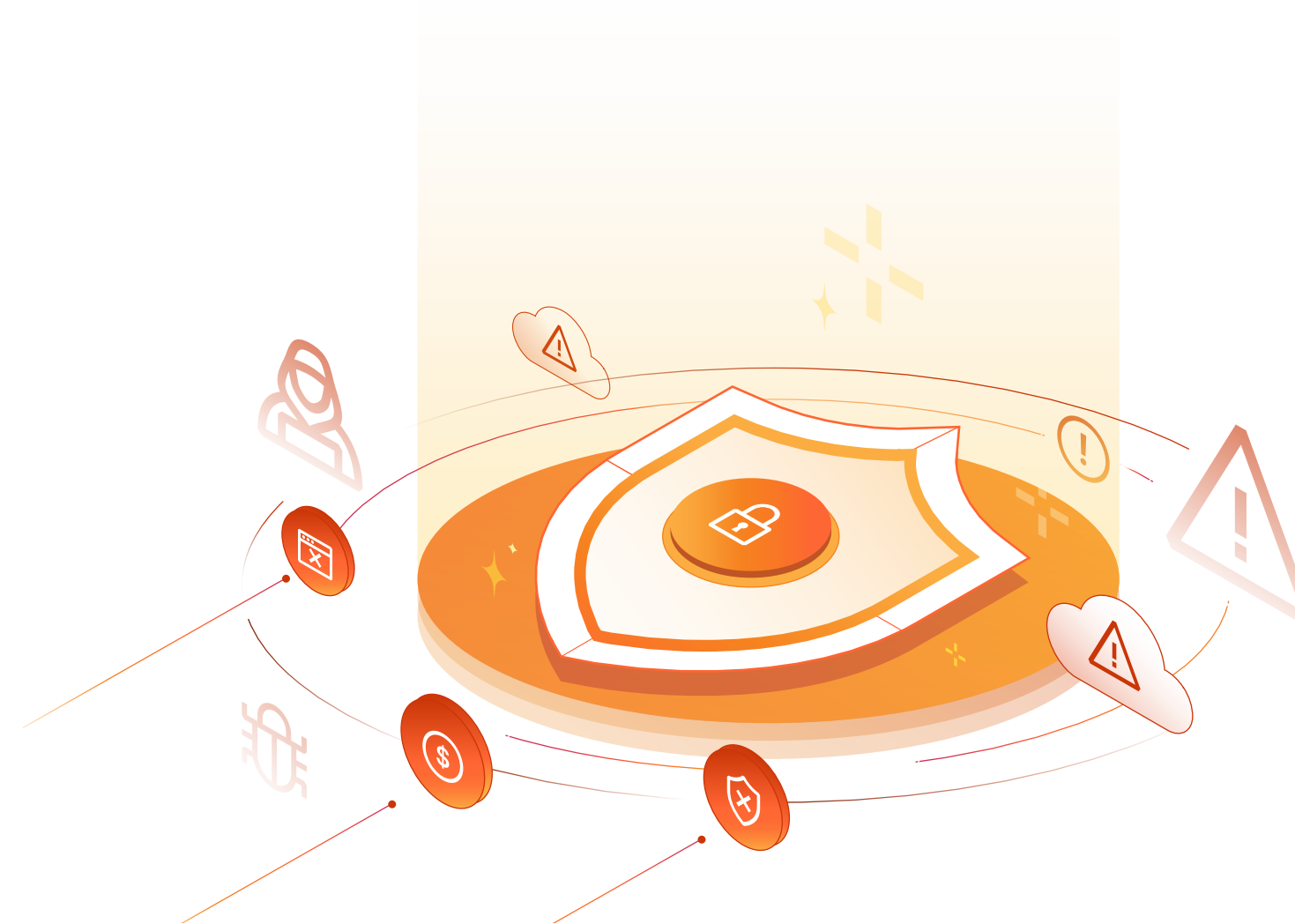
2 Adopt a multi-pronged approach to vulnerability response

Patching is often the best method to mitigate vulnerability risk, but it can't be the only method. Organizations must develop strategies to address both the window of exposure before a patch becomes available, as well as the time it takes to fully vet and deploy a patch into production systems.

The best place to start is attack surface reduction. Ensuring assets are properly protected, with network segmentation using Zero Trust principles, can dramatically reduce what threat actors can discover and compromise.

A web application firewall (WAF) that stops threats based on up-to-date threat intelligence is essential for protecting Internet-facing assets that can't live behind the network firewall. And in the case of major vulnerabilities with no patch or patching resources, creating (or using vendor-provided) rules to stop targeted exploitation is an effective mitigation until a patch can be applied.

Finally, given limited resources, and thousands of vulnerabilities disclosed each month, organizations must prioritize patching based on active exploitation. There are over 20,000 vulnerabilities disclosed each year, but most are never exploited. Focusing on known exploitation is the best way to quickly reduce the risk of compromise. A good start point is CISA's [Known Exploited Vulnerabilities Catalog](#).

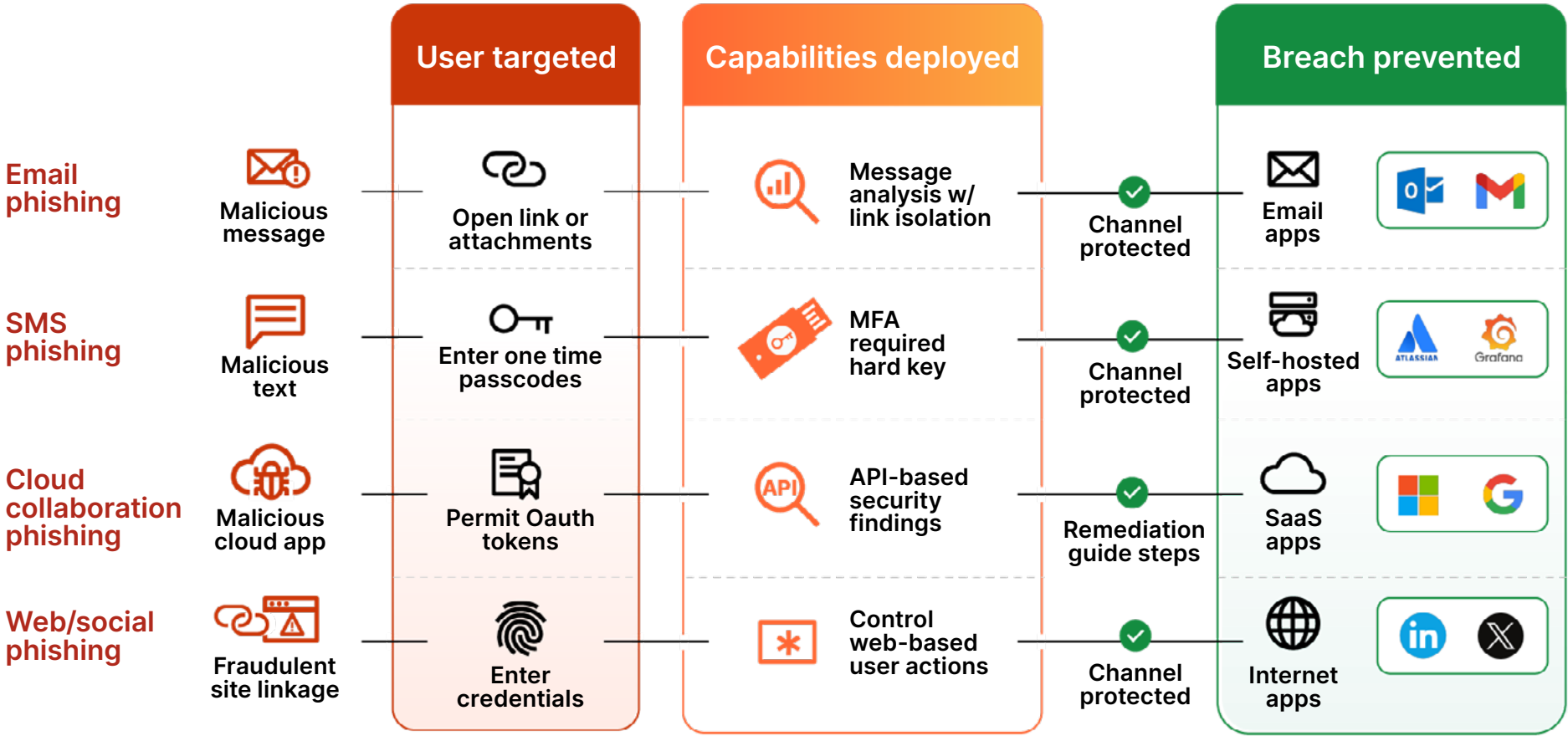


3

Assess and apply phishing controls that span all exposed user channels

Organizations cannot not rely solely upon security functions implemented within a cloud security provider. As the recent [investigation](#) by the DHS discovered, nation states are actively looking to disrupt security operations at major cloud providers such as Microsoft. Instead, organizations should identify a complimentary phishing solution to their email provider that protects against attacks beyond the inbox and implement a set of controls that are capable of addressing every point of exposure for users.

1. Start by assessing your current level of phishing risk within email to identify gaps in your existing defenses.
2. Implement email security capabilities that employ AI/ML-driven message analysis while providing pre- and post-delivery protection for continuous coverage. These capabilities should be paired with a low-touch approach for handling malicious links, even those hidden in QR codes or activated post-delivery.
3. Extend your phishing protection and provide another layer of defense by deploying ZTNA-enabled hard keys to prevent unauthorized access in the event that credentials are compromised.
4. Layer on capabilities that can help identify unauthorized applications that can act as a front door for attackers.
5. Complete your phishing solution by adopting technology that can isolate and control user actions on web-based apps to prevent malware and credential theft.



4 Measure and improve API maturity level over time

The most comprehensive approach for protecting APIs is to implement a holistic web application and API protection (WAAP) platform. However, an organization that is just beginning to acknowledge their API exposure may not find this feasible overnight. Progress needs to start somewhere. Cloudflare recommends implementing API protection in three phases.



Phase ① API visibility

Companies must first track and formally manage all their API endpoints, including any shadow APIs. However, when they do find APIs, it is difficult to accurately build a unique schema for each of potentially hundreds of API endpoints. With an API visibility service, organizations can both automatically discover API endpoints and identify who owns that API and how that API should be used.

Phase ② General web attack protection

Web applications and APIs often work together (for example, an ecommerce website using an API to process payments). However, the global nature of the Internet exposes websites and other applications to attacks from many locations, at various levels of scale and complexity. The following are 'table stakes' services to directly protect web applications and the APIs behind them from DoS and DDoS attacks, credential stuffing, zero-day vulnerabilities, and other threat types:

- DDoS mitigation services sit between a server and the public Internet to prevent surges of malicious traffic from overwhelming the server
- A Web Application Firewall (WAF) filters out traffic known (or suspected) to be taking advantage of web application vulnerabilities
- Encryption certification management helps manage key elements of the SSL/TLS encryption process
- Rate limiting protects endpoints from DoS attacks, brute-force login attempts, and other API traffic surges — without penalizing legitimate users.

Phase ③ API-specific attack protection

Tools like WAFs and DDoS are critical for web security and the (human) app user's experience, but these services were designed to protect applications — not APIs specifically. As an organization exposes more services via APIs, they should augment web app security with specific API security measures.

Advanced API security, using unsupervised machine learning, is capable of developing separate baselines for each API, and predicting the intent of API requests (whether legitimate or malicious) as they are made.

5

Get involved in LLM projects early

It is critical that the security team understand the various risks associated with LLM usage and development, and then work to be actively involved in any LLM deployments.

An internal LLM breach could be disastrous for an organization, and unmitigated access to public LLMs opens potential for sensitive data leakage, which is why organizations like Samsung [have banned](#) employee access after they discovered sensitive code had been leaked.

A good place to start is the [OWASP Top 10 for LLMs](#). Organizations should determine which vulnerabilities contribute the most risk to their organization's planned AI projects and then begin to prioritize protection measures. Protections can range in type and scope from data loss protection to implementing a [dedicated firewall for AI](#).

Overreliance is a key vulnerability to watch for, especially as it relates to developers using LLMs. AI-generated code still needs to be validated by humans first, and should be run through the same quality and security testing processes as human-generated code.

Top 10 vulnerabilities for Large Language Models

Model denial of Service Excessive resource-heavy requests leads to service degradation and costs	Prompt Injection Manipulation of model through crafty inputs to influence decision	Sensitive Information Disclosure Sensitive data being exfiltrated from the model	Supply Chain Vulnerabilities Vulnerable component embedded in the model	Insecure Plugin Design Plugins can be insecure inputs and insufficient access control
Insecure Output Handling Output accepted without validation. XSS, CSRF, SSRF	Model Theft Exfiltration of proprietary LLM model	Excess Agency Models can perform actions due to excessive permissions	Training Data Poisoning LLM training data is tampered bias and vulnerabilities are introduced	Overreliance Excessive trust on the output of LLM leading to misinformation

Define a roadmap to adopt Zero Trust

Time and time again have proven that perimeter-based security models do not work for the modern threat landscape. Increasingly, organizations are turning towards Zero Trust security best practices. The premise of never trust, always verify, is simple, and the benefits are clear — improved security outcomes, reduced breach costs, higher operational efficiency, and improved user experience — but successfully deploying Zero Trust is much easier said than done.

Typically, Zero Trust adoption is broken out into several phases. Implementations consistently start with a smaller use case or targeted set of users, prove out the value, and then expand from there. The plan will invariably evolve along the implementation path, but having a plan is essential to success.

There is no perfect answer for how to choose a starting point, but frequently-cited internal decision factors include:



Speed of implementation — contractors, in particular, have limited remote access needs that can often be fulfilled without installing end-user software, which can simplify a project rollout and provide a “quick win” for strengthening security



Flexibility and openness to change — for example, the security team might be the best first customer if a pilot project can be tightly scoped and done in tandem with existing infrastructure



Users/roles/apps that are at greater risk for attacks — developers who have access to valuable intellectual property, security/risk professionals, executives may be prime targets, or sensitive internal apps housing customer or financial data



Employee experience feedback — consider end-user complaints to determine which internal workflows could benefit the most from efforts to improve business productivity



Existing contract timing/logistics — upcoming contract renewals for current point solutions could steer your focus toward a relevant use case to address, and help create goal timelines for legacy solution augmentation or replacement

Unify security everywhere with Cloudflare

Cloudflare offers composable, scalable Everywhere Security to help reduce complexity and accelerate business innovation. Our cloud platform unifies many security capabilities and harnesses real-time threat intelligence to enforce low-touch, high-efficacy protections across users, applications, and corporate networks.

 [Explore how to enforce security without compromising innovation](#)

 [Contact us today for a consultation](#)



Phishing definitions

Account compromise — When an attacker takes control of a user’s email account. This is also referred to as Email Account Compromise (EAC), which is a close relative of Business Email Compromise (BEC). Attackers use a wide array of techniques such as dictionary brute forcing, credential harvesting attacks, and credential theft. The essential details are that a user’s email account credentials become compromised through malicious actions. Subsequently, the attacker uses that account to send malicious content to new targets.

ASN reputation — The overall score assigned to an Autonomous System Number (ASN) based on behavior. For example, ASNs from which high volumes of spam or malicious emails originate, will tend to have poorer reputations and thus lower scores. ASNs with low reputation scores are often used in attacks.

Attachment — Any file attached to an email that, when opened or executed in the context of an attack, includes a call-to-action (e.g., lures target to click a link) or performs a series of actions set by an attacker. If the intended victim opens an attachment or clicks a malicious attachment link, they may ultimately install a piece of malware that could lead to ransomware or follow-on operations through backdoors and RATs.

Brand impersonation — A form of identity deception where an attacker sends a phishing message that impersonates a recognizable company or brand. Brand impersonation is conducted using a wide range of techniques. A common one is display name spoofing, where the sender display name in the visible email headers includes a legitimate brand. In addition, attackers might use domain impersonation. In this case, the attacker registers a domain that looks similar to the impersonated brand’s domain, and uses it to send phishing messages.

Attackers often use various forms of obfuscation, such as homograph spoofing, in brand impersonation attacks. They might also register the exact same domain name as that used by the

impersonated brand but with a different top level domain (TLD). These techniques can be leveraged throughout all sections of an email, including the sender display name, email address (including the sender domain name), subject line, body content (HTML and plaintext), hypertext for links, and hyperlinks themselves (i.e., the actual URLs).

Business email compromise (BEC) — An increasingly common, effective, and costly targeted email attack designed to trick recipients into transferring funds, typically through forged invoices, to scammer accounts. BEC falls into various categories based on its sophistication, ranging from using a spoofed email to compromising a vendor in a supply chain attack.

Credential harvesters — Sites set up by an attacker to deceive users into providing their login credentials. This particular attack presents the user with a page that imitates an email or other account login page. Unwitting users may enter their credentials, ultimately providing attackers with access to their accounts. Because people often reuse passwords for multiple accounts, a member of your organization providing credentials to a harvester may give an attacker access to many accounts.

Deceptive link — When clicked, a deceptive link will open the user’s default web browser and render the data referenced in the link, or open an application directly (e.g. a PDF). Since the display text for a link (i.e., hypertext) in HTML can be arbitrarily set, attackers can make a URL appear as if it links to a benign site when, in fact, it is actually malicious. Malicious links can lead to arbitrary code execution or Remote Code Execution (RCE), credential harvesting, click fraud, unwanted installs, and other compromises.

Domain age (related to domain reputation) — The overall score assigned to a domain. For example, domains that send out a large number of new emails immediately after domain registration will tend to have a poorer reputation, and thus a lower score. Whereas older, known domains tend to have a positive reputation, and thus a

higher score. Domains with low reputation scores are often used in attacks.

Extortion — This tactic is commonly used to force a person or organization to perform a set of actions they would not otherwise normally perform. This is typically done under duress; for example, asking the intended victim to pay a ransom during a DDoS attack. The level of extortion can lead to a wide range of compromise depending on the attacker’s intentions and resources. Identity deception — This occurs when an attacker or someone with malicious intent sends an email claiming to be someone else. The mechanisms and tactics of this vary widely. Some tactics include registering domains that look similar (aka domain impersonation), are spoofed, or utilize display name tricks to appear to be sourced from a trusted domain. Other variations include sending email using domain fronting and high-reputation web services platforms.

Scam — A broad category of phishing fraud. The foundation is to entice a victim to provide money under a promise of a product, service, good, or even significant sum of money in return. The common theme is the transfer of money in a method that is atypical for the sender. Changes in common payment practices or sudden demands to pay sums via wire transfer can also be indicators.

Voice phishing — Also called “vishing,” this usually refers to the practice of leaving fake voice messages in hopes that victims will call back to provide personal information (such as) bank and credit card details), which will be used in other attacks. In our email security detections, we have observed attackers combining email and voice vectors by sending emails with attachments of a voicemail recording, media file or a link to a file. We have also observed attackers sending emails that had no malicious payloads, just a phone number

Endnotes

1. JetBrains disclosed CVE-2024-27198 on March 4th, 2024 at 14:59. Rapid7 published a [proof-of-concept analysis](#) of CVE-2024-2178 several hours later at 19:23 UTC. At 19:45 UTC, Cloudflare observed attempted exploitation of the vulnerability.
2. Based on a sample of threat indicators (“categories”) detected by the Cloudflare email security service between April 1, 2023 - March 31, 2024,. These indicators lead to email dispositions of malicious, BEC, spoof, or spam. Individual messages may contain multiple threat categories such as “Identity Deception”, “Brand Impersonation”, “Link”, and others that are described in the appendix.
3. For REST API endpoints, Cloudflare’s API Discovery found on median 33% more endpoints through machine learning than we discovered via customer-provided session identifiers across all customers’ domains/zones, per account, over the time period April 1, 2023 - March 31, 2024
4. Source: Enterprise Strategy Group, a division of TechTarget, Inc. Research Survey, Cloudflare Zero Trust for the Workforce Survey, May 2024.
5. Source: [Open AI Status Page](#)
6. Source: [CVE Details](#)
7. Source: Edgescan, [2024 Vulnerability Statistics Report](#)
8. Source: Google, [We’re All in this Together A Year in Review of Zero-Days Exploited In-the-Wild in 2023](#)
9. Between April 1, 2023 - March 31, 2024, API traffic with successful responses (200 status code) represented a median 58% of Cloudflare’s dynamic HTTP traffic. Dynamic content is content that changes based on factors specific to the user, such as time of visit, location, and device.
10. Source: [IDC Forecasts Spending on GenAI Solutions Will Reach \\$143 Billion in 2027 with a Five-Year Compound Annual Growth Rate of 73.3%](#)
11. Source: Oligo, [ShellTorch: Multiple Critical Vulnerabilities in PyTorch TorchServe Threatens Countless AI Users](#)



[Table of contents](#)

© 2024 Cloudflare Inc. All rights reserved.
The Cloudflare logo is a trademark of Cloudflare. All other
company and product names may be trademarks of the
respective companies with which they are associated.

Call: 1 888 99 FLARE
Email: enterprise@cloudflare.com
Visit: cloudflare.com

REV: BDES-5586.2024MAY01