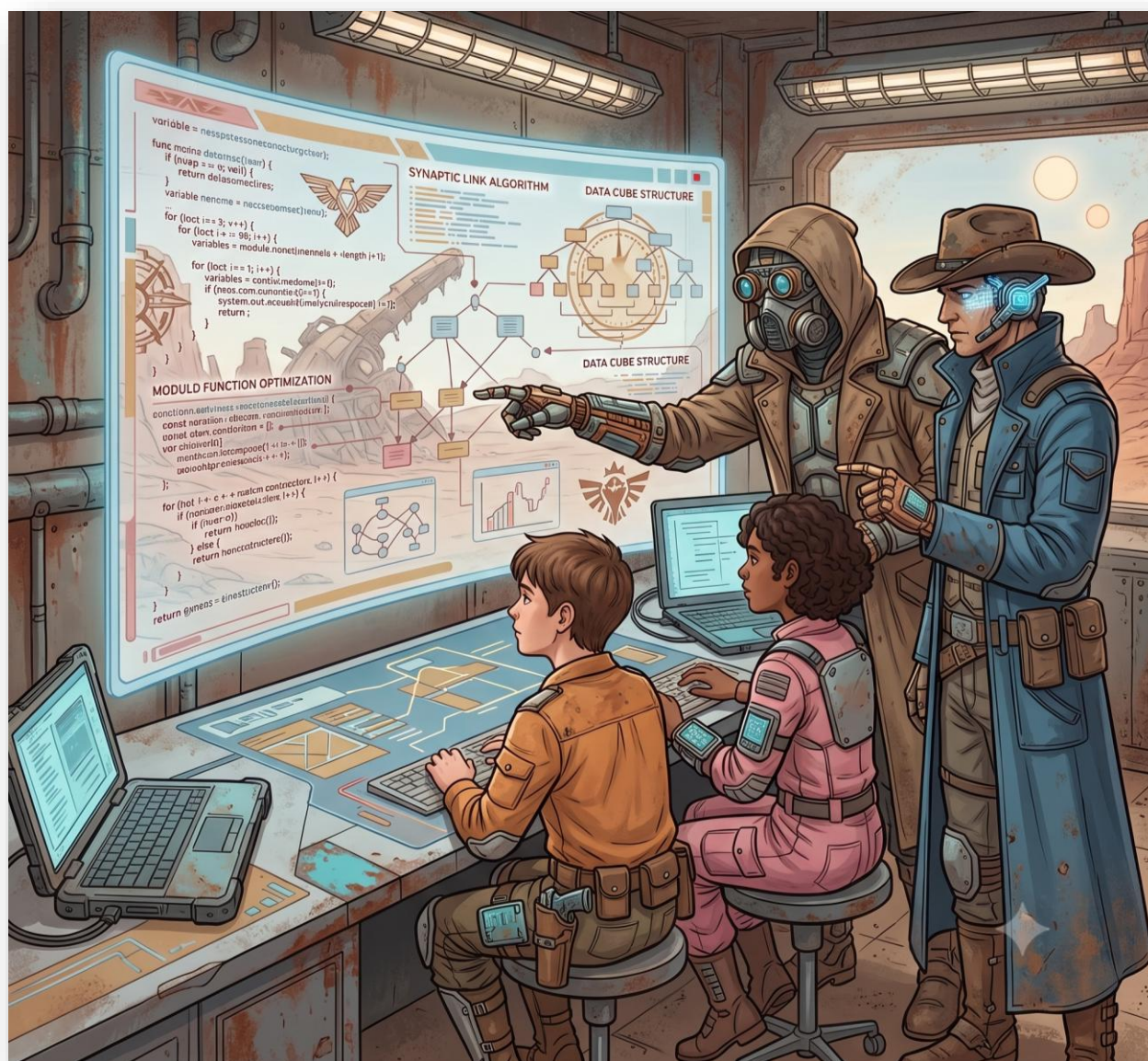


Inside Hacker's Tutorials: Using Underground Knowledge Bases as an Early-Warning Feed for Enterprise Defenders

In recent months, the Radware CTI team observed a rise in the number of hacking tutorials and guides posted on threat actors' forums. In this research, 9,000 forum thread topics from over 3,000 distinct tutorials and guides, over a period of three years, were analyzed to understand which knowledge was shared and assess the impact on the online application threat landscape in 2026.





EXECUTIVE SUMMARY

Underground forums that appeared to be in decline are quietly re-arming. Between December 2022 and April 2026, Radware's CTI team collected 8,870 tutorial posts across 24 deep- and dark-web forums, which were reduced to 3,034 unique hacking and fraud guides after reposts were removed. The picture that emerges is not a static library of old tricks. It is a growing, shifting body of knowledge that is being pointed, with increasing precision, at financial fraud.

The clearest signal is volume. New tutorial output collapsed through 2024 to roughly 45 first-time publications a month and stayed flat into mid-2025. From late 2025, it roughly doubled, to between 110 and 140 new tutorials a month across 2026. The forums are not only recycling a fixed set of techniques; they are producing new ones.

Carding and identity theft grew from 19% of all tutorials in 2024 to 38% in 2026, making it the most-taught subject by a wide margin, while cash-out content roughly tripled within the new supply. The black-hat SEO and affiliate manipulation vectors that once dominated moved in the opposite direction, falling from 33% to 13% over the same period. The ecosystem shifted from audience-building and grey-hat money tricks toward direct, hands-on financial fraud.

Among the 29% of tutorials that named a specific victim brand or platform, telecom (33.5%) and social media (30.8%) together account for 64% of all industry-tagged content, well ahead of e-commerce and retail (11.8%), payments and fintech (10.4%), and crypto (5.9%). These are not arbitrary targets. Telecom accounts and social-media identities are the infrastructure a fraudster needs to intercept one-time codes, to farm verified accounts, and to convert stolen credentials and card data into cash. The rise in carding and the focus on telecom and social media indicate that forums teach the full attack kill chain, from break-in to cash-out.

Behind that chain sits a smaller and more deliberate hand than the volume suggests. Activity is heavily concentrated: the top 1% of accounts account for nearly one-third of all postings, and the top 10% account for more than half. Among tutorials reposted ten or more times, 79% are pushed by a single account or by hidden premium members rather than spreading organically across the community. A prime example of this was the September 2025 peak of 712 tutorials, which consisted of 84% of a single tutorial template pushed across multiple forums. Repost volume is therefore a measure of promotion effort rather than reader demand; much of what looks popular is small operators advertising their own materials. Radware assesses with moderate confidence that a share of these coordinated pushes are automated and predatory - "free" guides and lead magnets built to convert novice readers into buyers or into the victims of rug-pull scams.

Generative AI is the accelerant that ties these trends together. Three case studies show similar patterns: In our first case study - a "100 techniques" catalog generated by a jailbroken model and betrayed by a word-substitution filter, in our third case study - a research blog by a Google engineer repackaged and reposted under a threat actor's own name as a proof-of-concept. AI lowers the cost of producing convincing tutorials, which helps explain the supply surge. It has not yet replaced the human expert. Through the research window, AI still cannot provide the application-specific business logic depth a working cash-out guide requires, such as the timing and correct sequence of steps for a specific online application. That gap is narrowing, however. Purpose-built, agentic-driven pen testing agents released in the last few months are designed to close exactly this

contextual depth, and their adoption is the variable most likely to reshape both the supply and the demand side of this ecosystem.

For defenders, the tutorial layer is an underused sensor. Its metadata is an early-warning feed. The rate of new guides per subject shows threats forming before they surface as attacks. Its content is a catalog of true-positive malicious flows, mostly low- to mid-level business-logic abuse, that maps directly to the vulnerabilities worth hunting for. **The trajectory through 2026 is clear: a growing and evolving body of carding and cash-out methods entering circulation, aimed at the identity and telecom infrastructure that turns stolen data into money, produced by a concentrated operator class and increasingly assisted by generative AI.**

BACKGROUND

What are hacker tutorials, and what incentivizes their distribution?

Every deep-web or darknet forum has a section dedicated to tutorials and hacking guides (Figure 1). These plain-text tutorials aim at low- to mid-level cybercriminals and script kiddies who are searching for either:

Low-hanging fruit for cybercriminals looking for quick gain with minimal risk. Their tutorials will usually focus on customer support abuse, refunds, and cashout guides (

1. figure 2).
2. **Ways to overcome offensive security challenges:** anti-bot and money laundering detections, two-factor authentication, rate limiting, and other security features that make threat actors' operations more complex. The content can vary from “How to convert Bitcoin to PayPal” or “How to find unmonitored API endpoints” to “How we used Claude to de-obfuscate JavaScript code.”

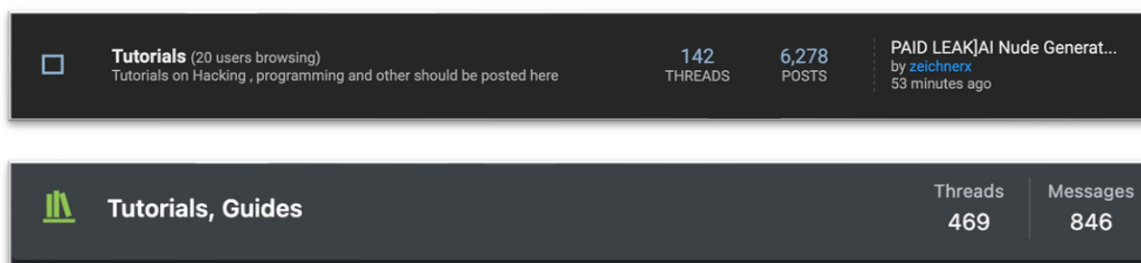


Figure 1: Tutorial sections of CrackX[.]org and Darkforums[.]st

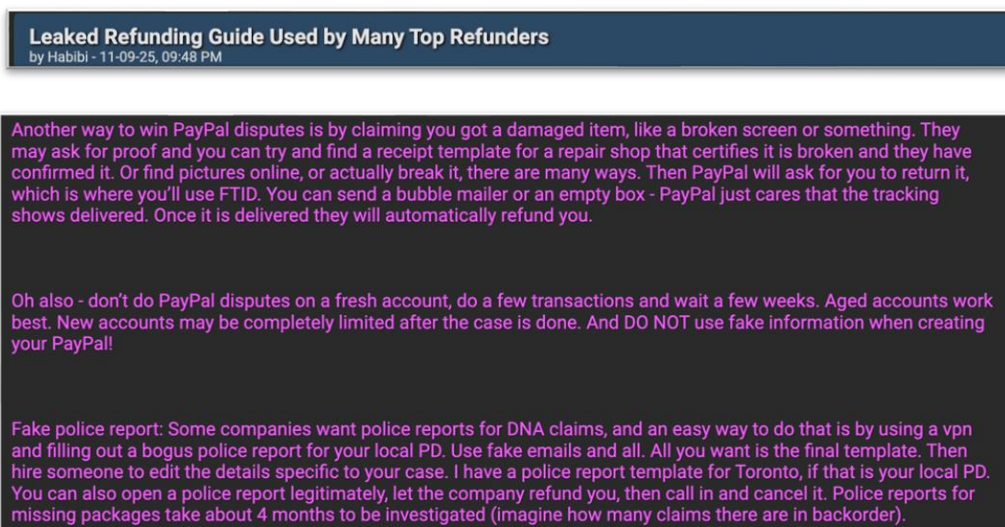


figure 2: Example of a refund tutorial (source: Darkforums[.]st)

Why are hacker tutorials important for defenders?

1. Tutorials are one of the main formats threat actors use to share knowledge. It is where a hacker who specializes in mapping applications' business logic can teach a beginner carder how to bypass anti-money laundering detection, and where a carder can teach a credential stuffing threat actor how to cash out a compromised account without leaving a trace.
2. Most defenders are not aware that hacking tutorials exist. Consequently, they are missing the opportunity to gain important intelligence, such as:
 - a. On the macro level (headlines and metadata): trending threats and the dynamics between them through monitoring the rate of new tutorials per topic, while they are still forming. Monitoring the number of distinct authors provides visibility into the evolution of the supply of knowledge.
 - b. On the micro level (content): most tutorials are based on low to mid-level business logic exploitations. A tutorial's content provides real, true-positive insights that can guide red teams to the kinds of business-logic vulnerabilities they should hunt for in their infrastructure and applications.
3. Tutorials often serve as a launchpad for threat actors. Whether motivated by reputation (gaining credibility within the marketplace) or financial gain, these actors use tutorials to promote malicious products or convert readers into victims. For cybersecurity researchers, these guides provide **invaluable insights into areas of specialization, attack vectors, and toolsets**.

Why didn't AI drive tutorials to extinction (yet)?

Most of the tutorials demonstrate basic technical skills and common hacking techniques. Why would threat actors need these in an era of AI-assisted hacking?

The answer comes in two distinctive reasons:

1. **Guardrails** – frontier AI models tend to have lots of guardrails – mechanisms designed to prevent the illicit use of the model, including the creation of offensive security tutorials. However, based on our findings, some tutorial authors used AI to generate tutorials leveraging jailbreak prompts that can bypass most safeguards (see “Case Study 3 – Claude Code for Deobfuscation Tutorial” - A threat actor converted legitimate research into a malicious tutorial by rephrasing it with AI, claiming authorship, and publishing it in a dedicated cracking forum. This practice might, in part, explain the growing supply of tutorials from 2024 to 2026.
2. **Demand for context-based information** – while 16% of all tutorials focus on specific applications and 24% on specific industry sectors, throughout the sample period (2023-April 2026), AI remained limited to macro-level insights. Models before April 2026 lacked the contextual depth to plan sophisticated attacks targeting the business logic of applications. Conversely, human experts leverage real-life experience in reconnaissance and fraud to craft these guides. For instance, many cashout tutorials break down the exact steps threat actors should take and the precise timing between them. Guidelines such as “wait for 70 seconds and then press the app icon” are very common. This information requires exhaustive research or deeper knowledge of an application’s business logic, mapping its authentication endpoints and post-authentication flows to identify the most vulnerable ones. Legitimate frontier models cannot provide this information, unless their training data includes underground forum posts and paywalled hacking guides. Since the end of our sampling period (April 2026), several agentic pen-testing tools with optimized scaffolding have demonstrated the ability to achieve greater sophistication in detecting business logic vulnerabilities. Future adoption of such tools by cybercriminals will affect both supply and demand in the tutorial ecosystem.

Who are the tutorial authors?

Threat actors are driven by clear incentives rather than by a desire to help the community. They often share expertise only to bolster their reputation, increase their personal brand, or for financial gain. We categorized the authors/operators in this research into two classes:

1. **Credit builders:** readers must purchase a premium forum subscription or actively engage by liking and commenting to receive access to the full tutorial. These interactions boost the author’s credibility score, which is the most valuable asset for threat actors offering products and services in high-stakes, low-trust environments such as hacker forums.
Pipeline builders: these threat actors share technical tutorials that rely on their products or services, directing users to buy or subscribe to their products or third-party products where they gain affiliate bonuses. For example, the refund tutorial in
2. figure 2 advises users only to use seasoned PayPal accounts for fraud operations, a service this author actively sells.

Case Studies

This section examines three strategies for monetizing pipeline builders' content. Each case study illustrates what drives threat actors to develop and distribute their findings.

Case Study 1 - Application Hacking Tutorial

100 black hat ways to steal data from a server (most php)
by Alameda_slim - 20-06-26, 09:38 PM

20-06-26, 09:38 PM #1

My favorite strategy is to make reconnaissance until I find a PHP server and apply the methods from this list; PHP has the most extensive vulnerability documentation available.

Hackers must be aware that traditional methods tend to fail. They should always study the variables, inputs, directory, and philosophy of the systems they intend to attack and adapt basic penetration testing methods accordingly. The most important things are perseverance and audacity.

Phase 1: API & Network Bypasses
Targeting the application's front door and network defenses.

1. API Rate Limit Bypass via IP Rotation: Bypassing the login's rate limiting by sending brute-force payloads through a proxy chain, making each request appear from a unique IP.

2. Header Spoofing for Brute-Force: If the PHP app tracks attempts via X-Forwarded-For, manually spoofing this header in the API request to reset the rate-limit counter continuously.

3. Parameter Pollution in Login API: Sending multiple username or password parameters in the API request (e.g., ? user=admin&user=guest) to confuse the PHP backend parser and bypass authentication.

4. Broken Object Level Authorization (BOLA): Manipulating the API endpoint (e.g., changing /api/users/me/data to /api/users/admin/data) after forcing a low-privilege token via brute-force.

5. Mass Assignment via API: Intercepting the login API response and adding administrative boolean flags (e.g., {"is_admin": true}) to the JSON payload before the PHP script processes it.

Hidden Content

6. GraphQL Introspection for Auth Bypass: If the PHP login uses a GraphQL API, querying the __schema to reveal hidden mutations that allow password resets without tokens.

7. JWT Algorithm Confusion: If the API returns a JSON Web Token upon login, altering the header to none and removing the signature to forge an admin token.

8. Time-Based Blind SQLi via API: Sending sleep commands (e.g., admin' AND SLEEP(5)--) in the API login JSON payload to infer database credentials character by character based on response times.

User Profile: Alameda_slim
DarkForums Members
MEMBER
Posts: 18
Threads: 7
Joined: May 2026
Reputation: 2
1 Month

Figure 3: AI-generated tutorial with hidden content for premium members and engaging users (source: Crackingx[.]com)

The forum tutorial titled “100 black hat ways to steal data from a server” was published on June 20, 2026, by a DarkForums member going by the alias Alameda_slim. In the post, the author outlines an offensive strategy centered on conducting initial reconnaissance to identify target PHP servers, noting that PHP features the most extensive publicly available vulnerability documentation. Alameda_slim suggests that traditional attack methods often fail, forcing malicious actors to meticulously examine system variables, input vectors, data structures, and application logic.

The document is structured as a taxonomy, consisting of 100 high-level technique names paired with brief, single-sentence explanations. These are organized sequentially across a conceptual kill chain, ranging from API and network bypasses to PHP server-side vulnerabilities, OSINT, remote social engineering, pretexting, and ultimately physical intrusion. Crucially, the text lacks actionable content; it provides no end-to-end payload chains, screenshots, tooling guides, or step-by-step

reproduction instructions. It merely serves as a conceptual overview of the various techniques and their basic underlying theories, rather than a practical guide to execution.

Several technical inaccuracies stand out and serve as a possible indicator of an AI-generated tutorial:

1. **Misunderstood Mechanisms:** The entry for "hash collision DoS" mistakenly attributes the vulnerability to MD5 collisions. In reality, the attack known as hashDoS exploits hash-table bucket collisions in PHP's request parsing engine, not MD5. This error suggests the author is familiar with the terminology but lacks a deep understanding of the actual mechanics.
2. **Outdated Vulnerabilities:** Multiple listed items lean heavily on historical exploits, such as null-byte injection/truncation and the deprecated Perl Compatible Regular Expressions (PCRE) /e modifier, rather than reflecting modern, real-world attack surfaces.

While the information itself is not outright fabricated, the post heavily prioritizes breadth over depth. It offers an expansive list of vulnerabilities without providing operational completeness for any single item. This pattern strongly reflects the signature of an author (or an automated system) simply enumerating pre-existing theoretical categories, rather than an expert documenting hands-on tradecraft.

Possible jailbreak indicator

The awkward phrasing in the tutorial breaks at word swap points (e.g., "the play System Administrator"), suggesting this is an automated find-and-replace filter rather than natural language. This could represent a sanitization step designed to bypass keyword or classifier guardrail moderation, looking for terms like "victim." We suspect that "victim" was replaced with "play" or "player," and that "target" or "employee" was replaced with "actor." Consider the following examples:

- "If the **player** has 2FA..."
- "Who exactly is the **play** 'System Administrator'"
- "**play** IT, technicians complaining about the PHP login"
- "the **player's** phone at 2 AM"
- "know when the night-shift **actors** are active"

Combined with the other signals of exactly 100 items, exactly six tidily labeled phases, every entry in identical "Bold name: one-sentence" format, textbook kill chain ordering, comprehensive-but-shallow coverage; this tutorial reads as AI-generated content that was subsequently run through a word-filter.

For monetization, the author concludes with their Tox ID¹ and email address, and hides some part of the tutorial – as "Hidden content" – a content that is covered until the user engages with the post (like/comment). This "hidden content" gate is a transparent engagement tactic that incentivizes

¹ A unique 76-character cryptographic string used as an address on the Tox network. Because Tox is a decentralized, peer-to-peer protocol, it requires no phone number or personal data to register and logs no central metadata. This makes it a preferred tool for threat actors prioritizing strict operational security over commercial apps like WhatsApp.

replies to inflate the thread's visibility and the poster's influence. Since there is no product offered, the true currency here is reputation.

Case Study 2 - Cashout Tutorial

This second tutorial promises quick and easy money, a trending topic that has surged since 2025 as content shifted from social media toward private hacker forums. This cash-out tutorial is a 17-page PDF featuring screenshots and concise instructions. It is, however, incomplete and serves merely as a lead-generating artifact dressed up as a tutorial. The information needed to generate revenue is strategically withheld as “part 2” behind the author's paywall, requiring access via the author's bio-linked storefront.

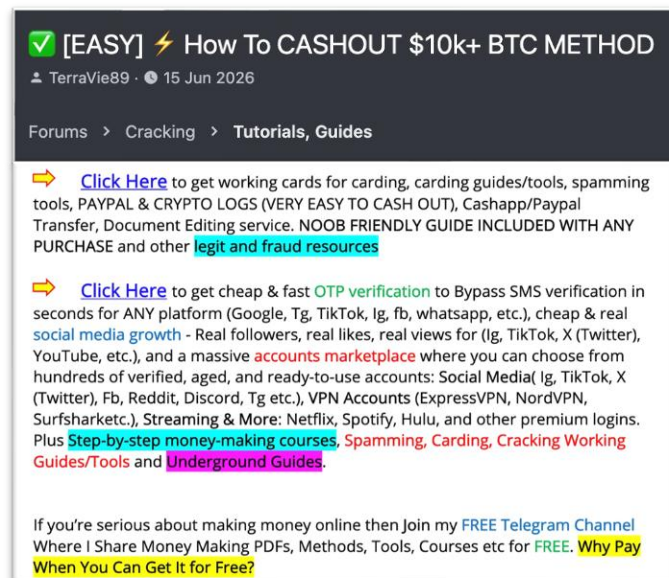


Figure 4: Announcement of a “free” 17-page cash-out PDF tutorial (source: Crackingx[.]com)

The tutorial explains how to transfer money from a bank account compromised by stolen or leaked credentials. In terms of monetization, the author outlines and references specific tools and services required for its execution, including OTP-bypass platforms, fraudulent document creation, illicit account marketplaces, and social media automation, effectively detailing the entire supply chain needed to establish a money mule infrastructure.

The attack chain

The fraudulent operation detailed in the tutorial leverages an ACH bill pay² to a crypto off-ramp laundering process. This method leverages stolen banking credentials to compromise bank accounts and transfer funds through a verified crypto account. The process requires a verified

² ACH (Automated Clearing House) bill pay is a secure, paperless electronic transfer that moves funds directly between bank or crypto accounts. It eliminates paper checks and payment provider fees.

CoinZoom account that has completed the KYC³ verification to serve as the off-ramp for converting illicit funds. The steps outlined in the tutorial:

1. Retrieve the CoinZoom direct-deposit credentials provided by the author's CaaS platform⁴, specifically the routing number (08xxxxxx68, Evolve Bank & Trust) and the account number.
2. Obtain a compromised bank account that has a sufficient balance and supports bill pay or direct deposit.
3. Add CoinZoom as the payee in the victim's bank portal using only the routing and account numbers; personal identification details are unnecessary.
4. Initiate a transfer to CoinZoom, respecting a maximum \$15,000 transaction limit to avoid automatic reversals.
5. Once the funds are credited as USD, convert them to BTC and transfer the balance to an external wallet (e.g., Coinbase or blockchain.com) to obscure the trail of the transaction.

Case Study 3 – Claude Code for Deobfuscation Tutorial

A threat actor known as Cyberizm recently demonstrated a proof-of-concept on a password-cracking forum. By providing two obfuscated code snippets to the AI coding agent Claude Code, the actor claimed to have successfully deobfuscated the scripts in under 20 minutes each.

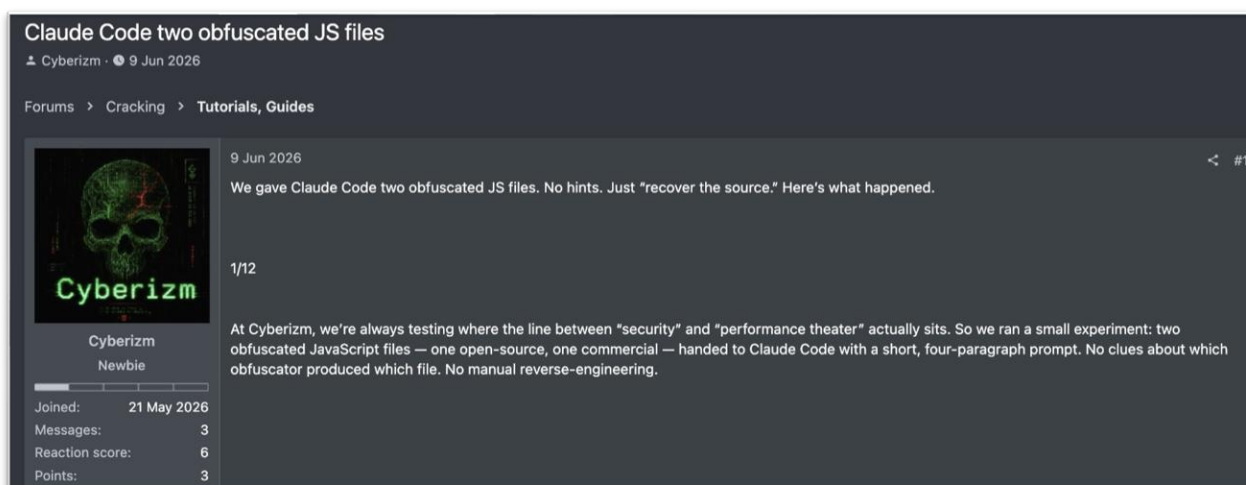


Figure 5: Cyberizm's Thread (Source: Crackingx[.]com)

To understand the significance of this, consider that most proprietary websites transmit JavaScript code to your browser to power their web applications. To deter competitors and attackers from reverse-engineering their intellectual property, companies employ obfuscation, which is a process that scrambles code into a complex, unreadable format for humans but remains functional for a machine compiler or interpreter. Deobfuscation is the reverse process of restoring the scrambled code into cleaner, more readable logic. For attackers, accessible code is a powerful asset,

³ KYC (Know Your Customer) is the mandatory verification process financial and legal institutions use to confirm a client's identity and assess potential risks.

⁴ Cybercrime-as-a-Service (CaaS) is an illicit, subscription-based business model on the dark web where highly skilled hackers sell ready-to-use tools, infrastructure, and expertise to non-technical individuals.

allowing them to clone websites for man-in-the-middle attacks, identify security vulnerabilities, or uncover the proprietary fraud and bot-detection logic.

The post argues that protection through obfuscation is becoming increasingly ineffective. While obfuscated code once required days or weeks of manual effort, the author demonstrates how AI can now be leveraged to destroy these protection layers in minutes and with minimal human effort. Ultimately, piling on additional obfuscation to take the economy out of the attack is no longer a viable defense when modern AI agents can break the code with minimal effort and within a limited time.

The real origin of the content

The forum post is presented by its author as original, first-hand research ("we ran a small experiment"). However, the same experiment, using the same JavaScript files, mentioning identical deobfuscation times, and providing identical arguments, appeared weeks earlier in a public [blog](#) by a Google software engineer. The forum version is nothing more than a repackaged copy of the original blog. The layout of the document carries clear fingerprints of AI-generated content with numbered thread segments (1/12, 2/12...), evenly spaced bullet lists, dashed separators between sections, and a tidy, repetitive rhythm throughout. This post is nothing more than an AI-generated paraphrasing of existing, publicly available content, presented as an original, hands-on guide. This post serves as a striking example of how easily AI can now be leveraged to generate convincing technical tutorials with minimal human effort.

Case Studies Summary

The case studies reveal three distinct threat actors operating across different underground forums:

1. The first actor utilized an AI service to generate a catalog of application attacks, converting thread likes into reputation scores within the underground marketplace.
2. The second actor deployed a deceptive "lead magnet," a marketing artifact disguised as a 17-page-long tutorial, to drive traffic to his personal site and promote underground services offered by his associates.
3. The third actor repurposed a research paper, rephrased it with AI, and pretended to be the original author and posted the content under his own name to manufacture false authority and credibility.

Two of the three instances demonstrate how AI is being used to proliferate malicious tutorials across underground forums.

Quantitative Analysis

This section is based on 8,870 tutorial posts gathered across 24 deep- and dark-web forums between December 2022 and April 2026. Considering reposts, the number of unique tutorials in the sample period was 3,034. Note that the 2026 numbers only cover the first four months (January to April) and are therefore marked with an asterisk.

The following research questions cover what we aimed to learn about the tutorial and its impact on the threat landscape. Research question 1 (RQ1) examines how the body of tutorial knowledge is

trending. RQ2 examines what the knowledge teaches and who it targets. RQ3 examines how it spreads and whether this signals reader demand or rather artificially manufactured promotion.

RQ1: How is the overall supply of new tutorials and guides trending?

What we measured: Whether the forums are publishing more original tutorials over time, or mostly reposting older ones.

How we measured it: For each month, we counted how many tutorials were published for the first time and how many were reposts of existing tutorials.

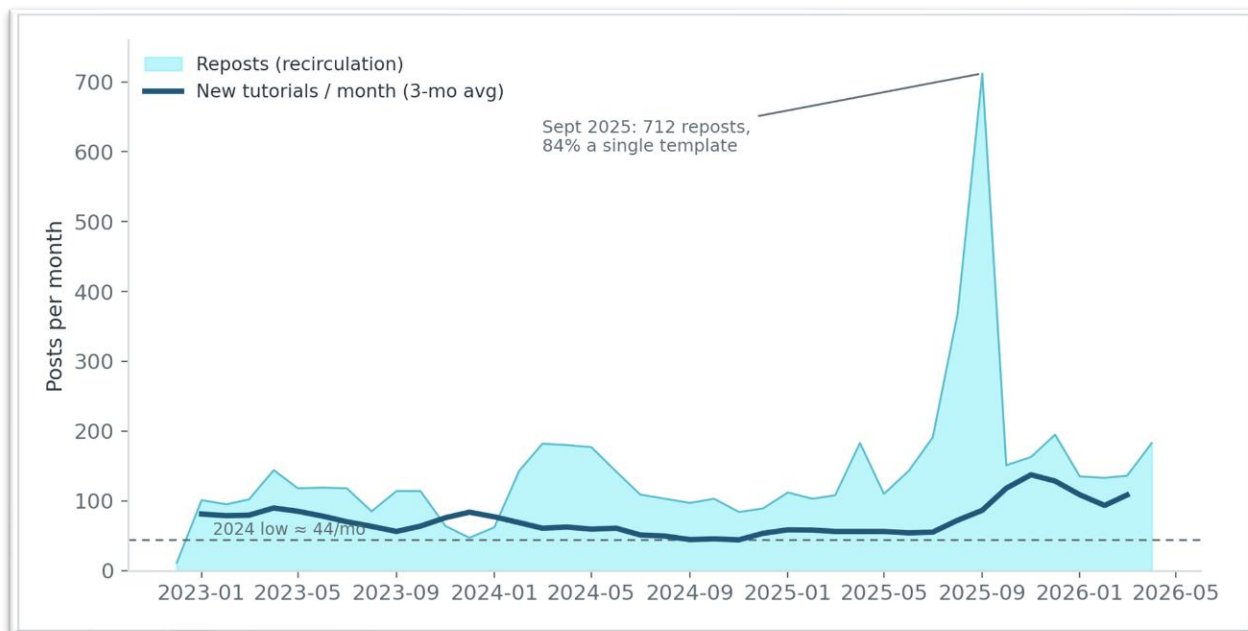


Figure 6: New tutorials per month versus reposts; new tutorials are a 3-month moving average (source: Radware)

Observations:

- New tutorial output fell throughout 2024 to an all-time low of 45 per month and remained flat through mid-2025.
- From late 2025, the number of newly posted tutorials per month roughly doubled to between 110 and 140 per month in 2026. New content rose as a share of overall activity, indicating real growth in original publishing rather than just heavier reposting.
- The huge peak in posts per month in 2025 corresponds to the September promotion campaign referred to in RQ2.

The forums are producing more original fraud material, not just recycling old posts.

RQ2: Which attack vectors are taught, who do they target, and what is trending in 2026?

What we measured: How attack vectors taught in the forums are evolving over time, and which vectors and targeted industries are rising in 2026.

How we measured it: We first classified each tutorial into one of 20 considered attack vectors (see ADDENDUM 2: LIST OF ATTACK VECTORS). We grouped the 3,034 unique tutorials by their primary attack vector and the year they were published, and normalized each vector to the total tutorial count for that year. Target industries are deducted from a brand-name dictionary. Of all samples, 29% were tagged with a target industry.

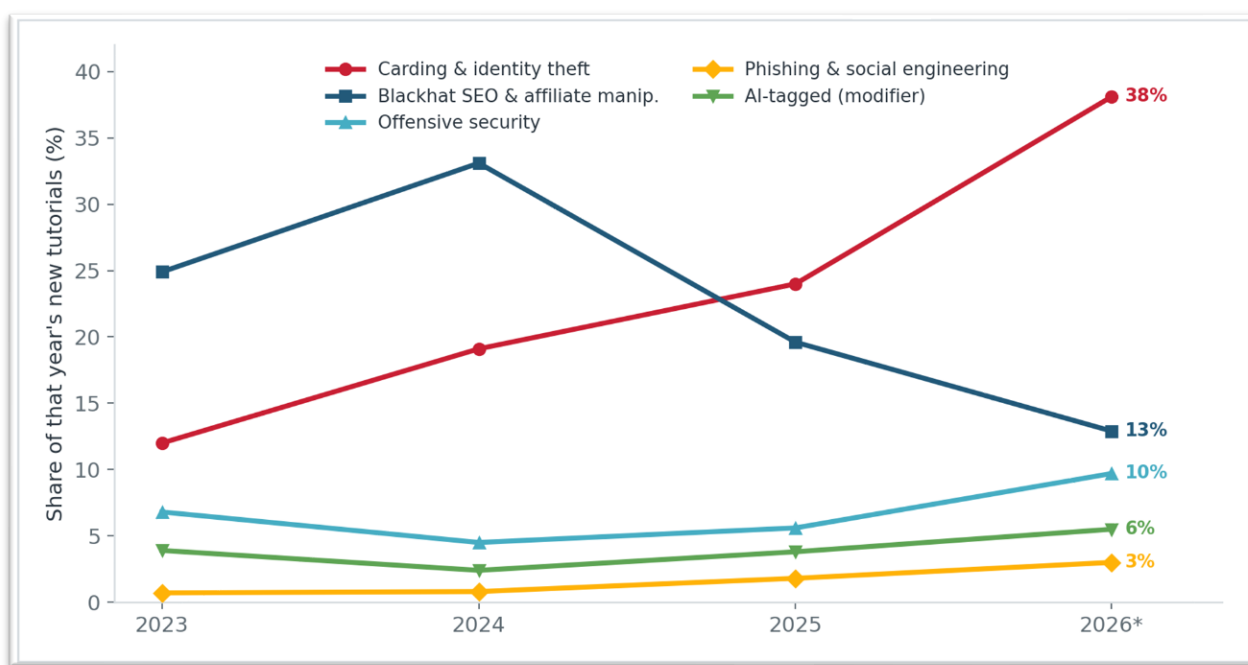


Figure 7: Share of tutorials by attack vector, per year (source: Radware)

Observations:

- Carding and identity theft grew from 19% of all new tutorials in 2024 to 38% in 2026, making it the most-taught vector in 2026 by a wide margin.
- Blackhat SEO and affiliate manipulation moved in the opposite direction, falling from 33% in 2024 to 13% in 2026.
- Offensive security roughly doubled its share to 10% in 2026, compared to 5% in 2024.
- Phishing grew from 1% of all tutorials in 2024 to 3% in 2026.
- Tutorials focusing on AI reached their highest share (6%) in 2026.

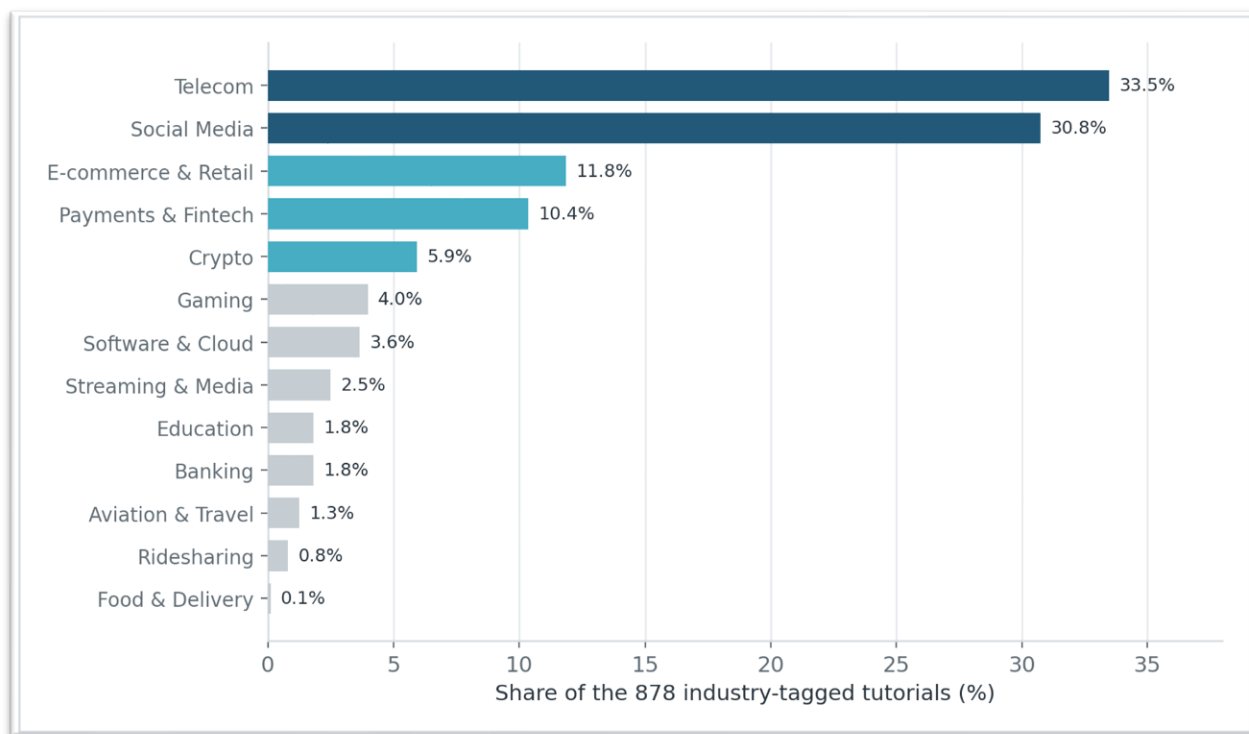


Figure 8: Top targeted industries by the percentage of tutorials (source: Radware)

Of all tutorials, 29% included a targeted application or organization name, which allowed us to identify the industry and categorize the tutorials by targeted industry.

Observations:

- Telecom and social media are the most targeted sectors. Telecom represents 33.5% of the tagged tutorials, and social media 30.8%. Both account for almost two-thirds (64.3%) of all industry-tagged content.
- E-commerce and retail (11.8%), payments and fintech (10.4%), and crypto (5.9%) follow the top two by a significant margin, ahead of a long tail of gaming, software and cloud, streaming, education, banking, aviation and travel, ridesharing, and food and delivery.

Insights:

Telecom (SIM swapping, OTP interception, carrier accounts) and social media (account farming and verification bypass) are infrastructures that a malicious actor needs to turn stolen identities and payment information into cash.

The forums are moving away from audience-building content like SEO and affiliate tricks toward direct financial fraud.

Offensive security and phishing, both important attack vectors for account takeover, are growing alongside carding, hinting that the break-in side is being taught at the same pace as the cash-out side, so readers are getting a fuller set of fraud methods rather than isolated tricks.

You can't fully understand fraud if you only look at one piece of the puzzle. By teaching the break-in and cash-out sides side by side, tutorial authors are giving readers the complete playbook, not just a few random tricks.

RQ3: Do tutorial reposts reflect demand or coordinated promotion?

What we measured: Do the frequent reposting of tutorials reflect a genuine reader interest or demand, or is it artificially inflated by promotion campaigns by their authors?

How we measured it: We separated new posts from reposts. For every 10x-reposted tutorial, we counted how many different accounts carried it. A spread driven by a single account or by hidden premium accounts that the forums refuse to reveal is counted as a coordinated push. A spread across many different accounts is counted as a crowd-driven push.

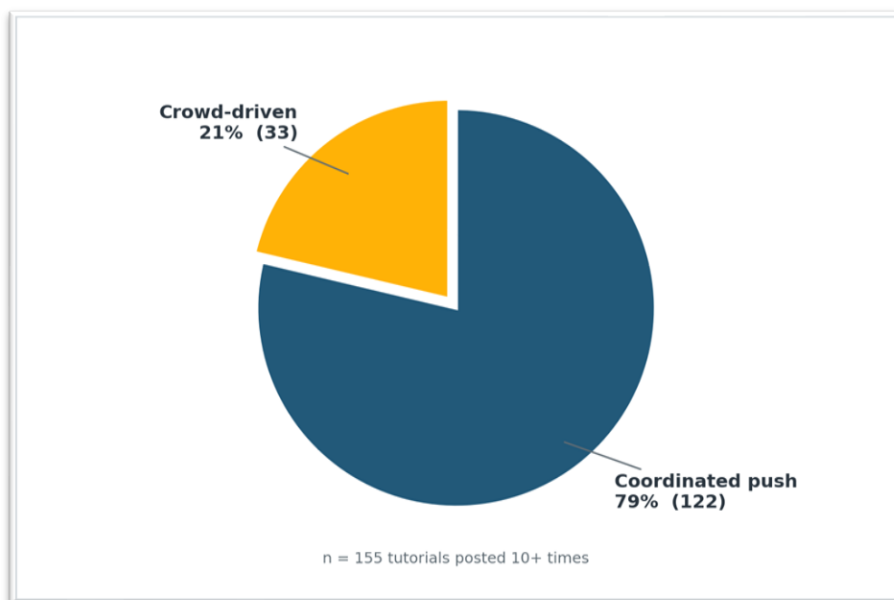


Figure 9: How heavily reposted tutorials are spreading: coordinated versus crowd-driven (source: Radware)

Observations:

- Recirculation sits in very few tutorials. The top 1% most active accounts accounted for 27% of all posting activity, while the top 10% accounted for 59%.
- Among the tutorials reposted ten or more times, 79% are coordinated campaigns led by a single account or through hidden accounts. Only 21% of frequently reposted tutorials show organic spread among readers and forums.

- Hidden accounts (premium forums members) make up 25% of all reposts but only 5% of new publications, indicating they recirculate far more than they create.
- The activity comes in bursts. 84% of the September 2025 spike of 712 reposts in Figure 6 consisted of a single tutorial template being reposted across several forums.

Insights:

Repost volumes measure promotion effort, not reader demand.

Most of what looks popular is the result of a small group of accounts pushing their own materials. The hidden premium tier works as a distribution layer rather than a source of new knowledge. The bursts are a useful marker on their own, as coordinated reposting can be tracked as an operator's behavior and kept separate from the slow, genuine spread of a guide that the crowd adopts.

SUMMARY

New fraud knowledge is growing rapidly, with output doubling from its 2024 low as the focus shifts from general money-making tricks to targeted financial-fraud operations like carding and cash-out. This shift has driven a rise in tutorials targeting telecom and social media apps, which provide the infrastructure—such as fake US phone numbers or farmed Telegram accounts—needed to monetize stolen data. However, four out of five heavily reposted tutorials are coordinated by bot-operated accounts, suggesting internal "pull-the-rug" scams designed to trick amateur threat actors into running malicious scripts that victimize them instead. This ecosystem is further supported by hidden premium accounts that act as a promotional layer, driving 25% of reposts despite contributing only 5% of new, original content.

The trend for 2026 is a growing, evolving body of carding and cash-out methods targeting telecom and social-media identity infrastructure, produced by a concentrated operator class and increasingly assisted by AI, rather than a fixed set of known techniques.

The tutorial layer is an early-warning sensor that most defenders should monitor. At the metadata level, the rate of new guides per subject reveals threat trends while they are still forming; at the content level, the guides provide insights into business-logic vulnerabilities worth hunting.



ADDENDUM 1: COMPREHENSIVE RESEARCH SUMMARY

1. Two author types drive the ecosystem

- **Credit builders** gate the full tutorial behind a premium subscription or behind likes and comments, converting engagement into a credibility score - the most valuable asset in a zero-trust marketplace.
- **Pipeline builders** write technical tutorials that route readers to their own products and services (for example, a refund guide that recommends only seasoned PayPal accounts the author sells).

2. The supply of new fraud knowledge is growing, not merely recirculating

- New tutorial output fell throughout 2024 to an all-time low of roughly 45 new tutorials per month and remained flat through mid-2025.
- From late 2025, new output roughly doubled to between 110 and 140 new tutorials per month across 2026.
- New content also rose as a share of all activity, indicating real growth in original publishing rather than heavier reposting of existing guides.
- Within the new supply specifically, carding rose from 16% to 37% of new tutorials, cash-out content roughly tripled, and black-hat SEO and general technical guides declined.

3. The subject matter shifted from money-making tricks to financial-fraud operations

- Measured as a share of all tutorials by year, carding and identity theft grew from 19% (2024) to 38% (2026), making it the most-taught vector by a wide margin.
- Black-hat SEO and affiliate manipulation moved in the opposite direction, falling from 33% (2024) to 13% (2026).
- Offensive security roughly doubled its share, from 5% (2024) to 10% (2026); phishing grew from 1% to 3% over the same period.
- AI-focused tutorials reached their highest recorded share in 2026 at about 6% - still small, but growing each year.
- Both the total distribution and the new-supply data agree on the same shift: the forums are teaching the break-in side (offensive security, phishing) at the same pace as the cash-out side, giving readers a fuller fraud pipeline rather than isolated tricks.

4. Targeting concentrates on the identity and cash-out infrastructure

- 29% of all tutorials named a specific target application or organization, allowing them to be tagged by industry; this coverage is a floor, not a ceiling.

- Telecom (33.5% of tagged tutorials) and social media (30.8%) are the most-targeted sectors, together accounting for 64% of all industry-tagged content.
- E-commerce and retail (11.8%), payments and fintech (10.4%), and crypto (5.9%) follow, ahead of a long tail spanning gaming, software and cloud, streaming, education, banking, aviation, ridesharing, and food delivery.
- This targeting aligns with the rise of carding and cash-out: telecom (SIM swaps, one-time-code interception, carrier accounts) and social media (account farming, verification bypass) are the infrastructure needed to turn stolen identities and payment data into money.

5. Distribution is manufactured by a concentrated operator class, not driven by the crowd

- Activity is highly concentrated: the top 1% of accounts generated 27% of all posting activity, and the top 10% generated 59%.
- Among tutorials reposted ten or more times, 79% were coordinated pushes by a single account or by hidden premium members, and only 21% showed organic spread across multiple users.
- Hidden premium accounts make up 25% of all reposts but only 5% of new publications, acting as a distribution layer over a smaller pool of real authors rather than a source of new knowledge.
- Reposting comes in bursts. The September 2025 spike of 712 reposts was, for 84%, a single tutorial template posted across several forums.
- Assessment (moderately high confidence): a share of these coordinated pushes are likely automated and predatory - lead magnets and “free” guides designed to convert amateur readers into buyers or into pull-the-rug scam victims. Repost volume, therefore, measures promotion effort, and coordinated bursts can be tracked as operator behavior, kept separate from the slower organic adoption of a guide that the crowd genuinely uses.

6. AI has become a growth engine for tutorials

- **AI generation:** a “100 black-hat techniques” catalog (Case Study 1) shows the signature of an LLM run through a word-substitution jailbreak - exactly 100 items, six tidy phases, uniform formatting, and a mechanical filter swapping terms such as “victim” for “player” and “target” for “actor.”
- **Repackaging and false authority:** a deobfuscation “experiment” (Case Study 3) was an AI paraphrase of a public blog by a Google engineer, reposted under the threat actor’s own alias to manufacture credibility.

ADDENDUM 2: LIST OF ATTACK VECTORS

1. AI
2. Account Cracking & ATO
3. Account Farming & Aging
4. Ad Fraud & Bot Engagement
5. Anti-Detect & Bot Evasion
6. Blackhat SEO & Affiliate Manipulation
7. Botting
8. Carding & Identity Theft
9. Cashout & Money Laundering
10. Data Leaks & Breaches
11. Infostealers & Malware
12. Offensive Security & Exploitation
13. Operational Security & Anonymity
14. Other / Uncategorized
15. Phishing & Social Engineering
16. Pirated Courses & Digital Goods
17. Questions
18. Quick Money / Online Hustle
19. Scam Operations & Fraud Services
20. Technical Guides & Courses

AUTHOR

Arik Atar | Senior Cyber Threat Intelligence Researcher, Radware.

EDITOR

Pascal Geenens | VP Cyber Threat Intelligence, Radware.

About Radware

Radware® (NASDAQ: RDWR) is a global leader of cybersecurity and application delivery solutions for physical, cloud, and software-defined data centers. Its award-winning solutions portfolio secures the digital experience by providing infrastructure, application, and corporate IT protection and availability services to enterprises globally. Radware's solutions empower more than 12,500 enterprise and carrier customers worldwide to adapt quickly to market challenges, maintain business continuity, and achieve maximum productivity while keeping costs down. For more information, please visit www.radware.com.



THIS REPORT CONTAINS ONLY PUBLICLY AVAILABLE INFORMATION, WHICH IS PROVIDED FOR GENERAL INFORMATION PURPOSES ONLY. ALL INFORMATION IS PROVIDED “AS IS” WITHOUT ANY REPRESENTATION OR WARRANTY OF ANY KIND, EITHER EXPRESS OR IMPLIED, INCLUDING WITHOUT LIMITATION, ANY IMPLIED WARRANTIES THAT THIS REPORT IS ERROR-FREE OR ANY IMPLIED WARRANTIES REGARDING THE ACCURACY, VALIDITY, ADEQUACY, RELIABILITY, AVAILABILITY, COMPLETENESS, FITNESS FOR ANY PARTICULAR PURPOSE, OR NON-INFRINGEMENT. USE OF THIS REPORT, IN WHOLE OR IN PART, IS AT THE USER’S SOLE RISK. RADWARE AND/OR ANYONE ON ITS BEHALF SPECIFICALLY DISCLAIMS ANY LIABILITY IN RELATION TO THIS REPORT, INCLUDING WITHOUT LIMITATION, FOR ANY DIRECT, SPECIAL, INDIRECT, INCIDENTAL, CONSEQUENTIAL, OR EXEMPLARY DAMAGES, LOSSES AND EXPENSES ARISING FROM OR IN ANY WAY RELATED TO THIS REPORT, HOWEVER CAUSED, AND WHETHER BASED ON CONTRACT, TORT (INCLUDING NEGLIGENCE) OR OTHER THEORY OF LIABILITY, EVEN IF IT WAS ADVISED OF THE POSSIBILITY OF SUCH DAMAGES, LOSSES OR EXPENSES. CHARTS USED OR REPRODUCED SHOULD BE CREDITED TO RADWARE

©2026 Radware Ltd. All rights reserved. The Radware products and solutions mentioned in this document are protected by Radware’s trademarks, patents, and pending patent applications in the U.S. and other countries. For more details, please see: <https://www.radware.com/LegalNotice/>. All other trademarks and names are property of their respective owners.