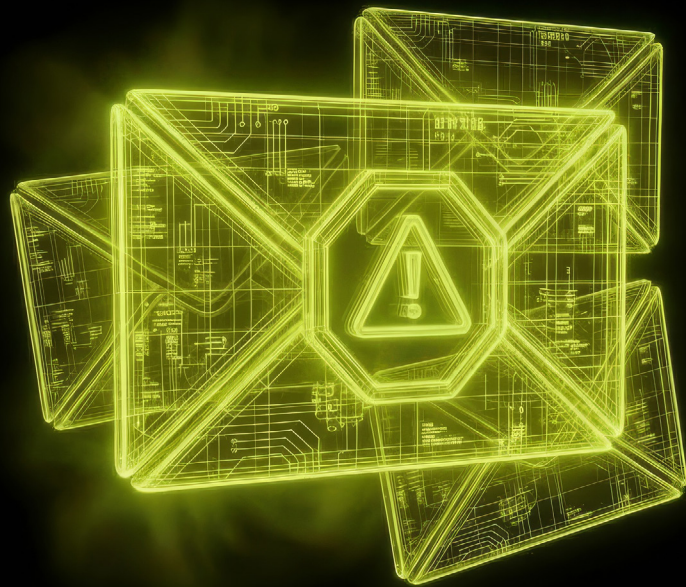


Abnormal

THREAT REPORT

2026 Threat Outlook:

5 Email Attacks You Need to Know



Executive Summary

In every organization, from sole proprietors to global conglomerates, email serves as the primary communication channel for business operations—and the most reliable entry point for threat actors.

What makes email particularly effective for attackers isn't just its ubiquity, but its direct access to the human element. Employees must constantly engage with the inbox, creating countless opportunities to exploit psychology, familiarity, and routine behavior through seemingly benign interactions. This dynamic rewards cybercriminals who can convincingly operate within legitimate correspondence rather than outside it.

The five attacks highlighted in this report reflect a continued shift toward high-effort, identity-centric threats that abuse everyday workflows rather than technical weaknesses. Multi-stage QR code phishing and thread-spoofed vendor impersonation demonstrate how threat actors increasingly rely on layered pretexts, fabricated context, and realistic artifacts to condition targets and bypass detection. OAuth consent phishing and lateral phishing illustrate the growing value of persistence and post-compromise access, enabling adversaries to maintain a foothold within internal environments while evading signature-based controls. Finally, AI-generated payroll fraud underscores how generative AI is accelerating reconnaissance, personalization, and execution across business email compromise attacks, producing high-impact outcomes with minimal technical indicators.

The potency of these threats lies in their ability to seamlessly blend into normal business activity. By impersonating known contacts, leveraging compromised accounts, and weaponizing trusted platforms, attackers generate messages that are structurally and contextually consistent with ordinary enterprise communication, eroding the reliability of traditional indicators of compromise.

The result is a threat landscape in which conventional security solutions, such as legacy secure email gateways, struggle to distinguish malicious intent from authorized activity. Built to identify known indicators and overtly suspicious signals, these tools are ill-suited to stop attacks designed to appear routine, credible, and contextually appropriate. Closing this gap requires AI-native defenses that understand identity, behavior, and context across interactions and automatically detect anomalies, rather than relying on predefined rules or static signals.



Table of Contents

5 Email Attacks Set to Increase Enterprise Exposure in 2026

04

▪ 01 Multi-Stage QR Code Phishing

05

▪ 02 Thread-Spoofed Vendor Impersonation

08

▪ 03 OAuth Consent Phishing

11

▪ 04 Lateral Phishing

14

▪ 05 AI-Generated Payroll Fraud

17

Predictions for 2026 and Beyond

19

Defending Against New and Emerging Threats

20

About Abnormal

21



5 Email Attacks Set to Increase Enterprise Exposure in 2026



▶ EMAIL IS THE CORNERSTONE OF BUSINESS COMMUNICATION, UNIVERSALLY ADOPTED ACROSS INDUSTRIES AND LOCATIONS.

That widespread reliance has long made it an attractive attack vector, providing threat actors with a consistent environment in which they can test, refine, and repeat tactics.

Modern attackers continually probe for ways to evade organizational security measures and manipulate targets via the inbox. Rather than relying on a single technique, they adapt their approach, experimenting with new pretexts, formats, and delivery mechanisms until they achieve the desired outcome.

The following are real-world examples of threats Abnormal AI customers received in 2025. They illustrate how attackers are actively applying these methods today and highlight the strategies organizations must be prepared to detect and defend against in 2026.



Multi-Stage QR Code Phishing

During the first part of the decade, malicious QR codes seemed to be everywhere, and threat actors went all in on QR code phishing attacks. Initially, these attacks relied on simple pretexts and involved minimal steps. Typically, threat actors sent a single email containing a malicious QR code that, when scanned, led the target directly to a fake login page—often an impersonated Microsoft or Google portal—with a prompt to enter their credentials.

But over the past year or two, threat actors have begun utilizing malicious QR codes in considerably more complex ways. Multi-stage QR code phishing incorporates additional steps into the attack flow to better support the appearance of legitimacy and evade detection. Moreover, whereas traditional QR code phishing attacks primarily revolved around fraudulent multi-factor authentication expiration notices and shared document notifications (which at one point accounted for nearly 50% of all QR code attacks detected by Abnormal), multi-stage QR code phishing attacks leverage a variety of pretexts.

Thus, by the time targets reach the credential harvesting page, they've been conditioned through multiple interactions to trust the workflow. Often, the final page also includes sophisticated personalization, such as company branding or pre-populated email addresses. This layered approach not only minimizes the target's suspicions but also helps attackers bypass security controls that may flag direct redirects to known malicious domains.



This steady development in strategy from simple, single-scan lures to multi-stage workflows reflects a broader move toward deliberate, high-effort phishing campaigns designed to outmaneuver both user skepticism and traditional detection methods—an approach we expect to accelerate further in 2026.

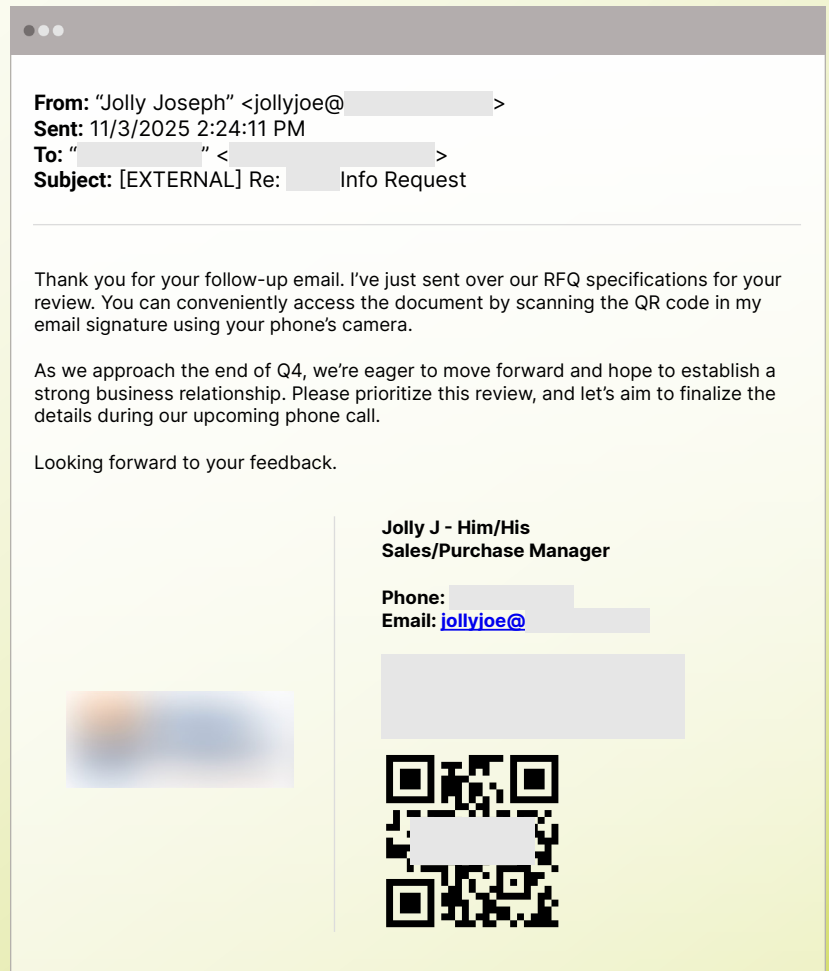
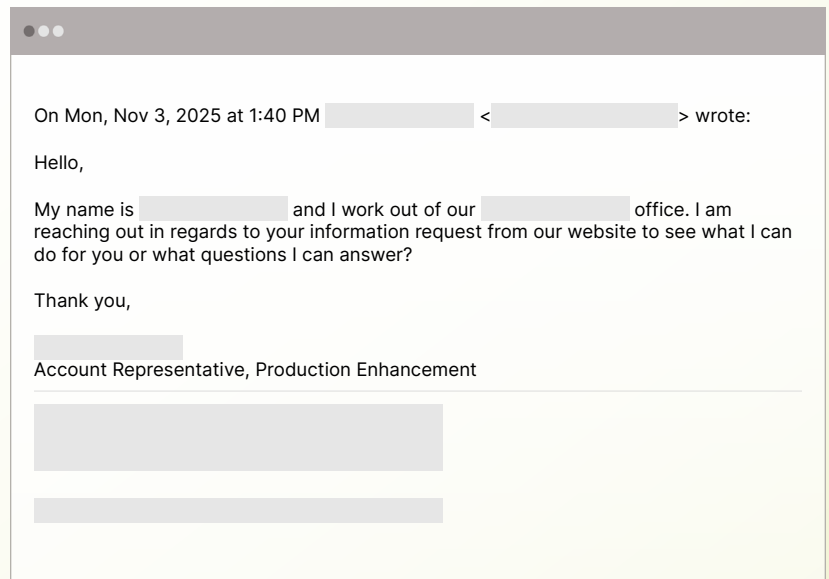


Real-World Example of Multi-Stage QR Code Phishing

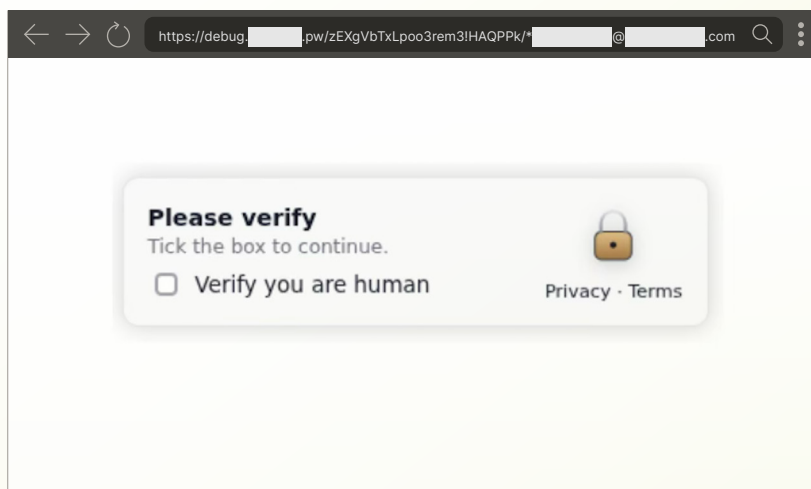
In this attack, the threat actor initiates contact by submitting an information request via a form on the targeted organization's website. That submission prompts a legitimate employee to reach out, creating a real thread the attacker can then exploit.

In their reply, the attacker assumes the identity of a vendor-side representative. The message presents a request for quote (RFQ)-related pretext and includes end-of-quarter urgency, encouraging the recipient to prioritize the request. To increase the appearance of legitimacy, the threat actor uses a recently registered lookalike domain and incorporates official-looking branding, a corporate address, and professional signature elements.

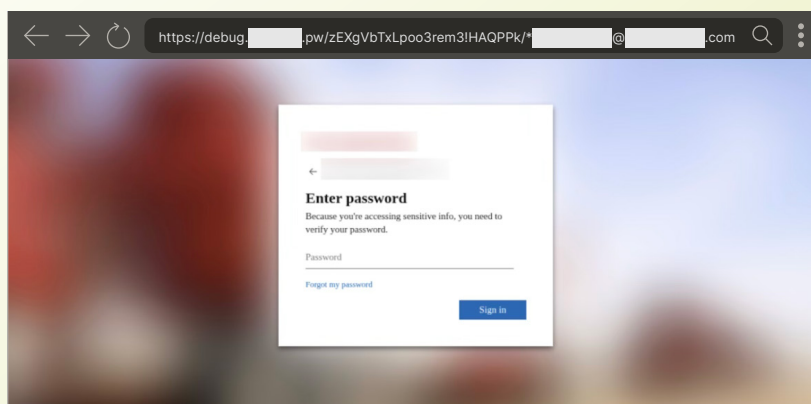
The attacker instructs the recipient to scan a QR code "using your phone's camera" to access the supposed RFQ specifications. This shifts the interaction to a mobile device outside corporate email protections and bypasses link-scanning controls.



When scanned, the QR code redirects the victim to a “Verify you are human” page. The final URL includes the target’s email address as a path component, indicating the flow is personalized and likely used to prefill downstream content.



After this verification step, the target is presented with a branded login page mimicking the company’s authentication portal. The page displays the correct email address already prefilled and requests a password, positioning the attacker to harvest credentials if the target submits their information.



Detecting Multi-Stage QR Code Phishing

Legacy secure email gateways struggle to detect multi-stage QR code phishing not simply because the email looks legitimate, but because the attack deliberately fragments visibility across environments. The email itself resembles a routine RFQ follow-up within an existing thread, complete with professional branding and a plausible vendor identity, offering little for traditional scanners to analyze. The only actionable element is a QR code embedded in the signature, leaving URL- and attachment-based scanners with little to nothing to inspect. Once scanned, the interaction immediately moves to a mobile device, where redirects and credential harvesting occur outside the corporate email perimeter, preventing legacy tools from observing the full attack sequence.

Abnormal detects this threat by treating the email as the first step in a broader, multi-stage attack rather than a self-contained event. Its behavioral AI understands normal vendor relationships and communication patterns, allowing it to identify anomalies, such as a first-time sender using a newly registered lookalike domain and prompting an unusual action. Abnormal also decodes QR codes, inspects downstream URLs, and correlates those findings with contextual signals. By identifying deviations across identity, behavior, and workflow, Abnormal can confidently classify the message as malicious and stop the attack before an employee engages.

Thread-Spoofed Vendor Impersonation

In thread-spoofed vendor impersonation attacks, threat actors use fabricated email threads that appear to be legitimate correspondence between a vendor and an internal employee to establish a believable pretext. The attacker inserts the fake message chain into their initial email to act as supporting evidence that validates the authenticity of their inquiry, which is usually related to a financial workflow—such as updating bank details or resolving an overdue invoice. In the thread, the impersonated employee grants permission to proceed with the impersonated vendor's request.

Typically, the employee the attacker is posing as is an authority figure, which creates pressure to comply, and the fictitious exchange is specific enough to feel relevant but also generic enough to avoid accidentally creating flags that indicate the request is fraudulent. To create a semblance of authenticity, attackers commonly use look-alike domains and, when possible, leverage compromised legitimate domains to bypass blocklists. They will also include fabricated documentation, such as doctored invoices and even altered tax forms, to deceive targets into completing the transaction.



Thread-spoofed vendor impersonation represents a notable evolution in financial fraud tactics, in which attackers invest more effort in constructing believable pretenses and prioritize realism over volume. In 2026, pretext-heavy attacks like this are expected to feature more prominently in vendor impersonation campaigns, reflecting a preference for manipulating trusted business processes over exploiting technical vulnerabilities.



Real-World Example of Thread-Spoofed Vendor Impersonation

At its core, thread-spoofed vendor impersonation is quintessential social engineering. It exploits human psychology, relies on the imitation of trusted parties, and leverages a convincing pretext to manipulate the target into making a harmful decision. The example below is an excellent demonstration of these elements in action.

Sent from a compromised legitimate address, the message purports to be from ZoomInfo, a well-known business intelligence platform, and uses a subject line that appears as a forward (Fw: Prospect Intelligence Platform - Discover Hidden Opportunities). The attacker informs the target that the CEO, "P.Z.," has asked them to send an outstanding invoice for payment and invites the recipient to "Please see the conversation below for more details."

From: [redacted] <no-reply@govquest.com>
Sent: 10/21/2025 6:33:42 PM
Reply-to: [redacted] <[redacted]@zoom-info-us.com>
To: [redacted]
Subject: Fw: Prospect Intelligence Platform - Discover Hidden Opportunities

Hello,

Attached is the outstanding invoice #9F3ERQKJ. [redacted] asked me to send it to you. Please see the conversation below for details.

Could you kindly let us know when we can expect payment for this invoice? We aim to resolve this matter promptly to ensure the payment is processed as soon as possible.

We appreciate your prompt attention to this to avoid any service interruptions.

Let me know if you need more info.

Best regards,

[redacted]
Accounting

Zoominfo Technologies LLC
805 Broadway St, Vancouver, WA 98660, United States

Begin forwarded message

From: [redacted]
Sent: Thursday, October 16, 2025 10:34 AM
To: [redacted]
Subject: Re: Prospect Intelligence Platform - Discover Hidden Opportunities

Hi [redacted],

Thank you for the bill, everything looks good.

Could you please forward a copy of your bill to [redacted] for swift payment? We'll ensure the payment is made promptly to avoid service interruptions.

I'm looking forward to use your database once everything is set up.

Thank you,
[redacted]

Begin forwarded message

From: [redacted]
Sent: Monday, October 13, 2025 09:12 AM
To: [redacted]
Subject: Re: Prospect Intelligence Platform - Discover Hidden Opportunities

Hi [redacted],

Attached is the final bill for invoice (#9F3ERQKJ) covering six user licenses (one complimentary), Access to all companies and contacts in the US and Canada, and 10,000,000 exports credit.

The total amount due is noted on the bill with payment due upon receipt.

Timely payment will help avoid any service interruptions.

Let me know if you have any questions.

Best regards,

[redacted]
Accounting

Zoominfo Technologies LLC
805 Broadway St, Vancouver, WA 98660, United States



As is typical of thread-spoofed vendor impersonation attacks, the phrasing is vague and oriented toward immediate payment. The body emphasizes urgency with phrases like “We aim to resolve this matter promptly to ensure the payment is processed as soon as possible” and “We appreciate your prompt attention to this to avoid any service interruptions.” The attack included two attachments: an invoice requesting nearly \$50,000 via ACH payment and a W-9 form presented as vendor verification.

Should the target proceed with processing the invoice, they will unknowingly transfer tens of thousands of dollars directly to the threat actor.

Detecting Thread-Spoofed Vendor Impersonation

In financial workflows, familiarity and apparent approval are often treated as signals of legitimacy—an assumption thread-spoofed vendor impersonation exploits directly. Many traditional security tools are unlikely to flag this as an attack because the message closely resembles a routine invoice follow-up from a reputable provider, with a believable subject line, professional signature, and attached invoice and W-9 that contain no malware or obviously malicious links. Because the email originates from a legitimate domain (govquest.com) and includes what appears to be a real internal approval thread from leadership, signature- and reputation-based filters see little risk. As a result, the fraudulent payment request can land directly in inboxes.

Detecting this attack requires validating business intent, not just message authenticity. Abnormal evaluates vendor communications in the context of established financial workflows, modeling known relationships, payment behaviors, and approval patterns. It recognizes that the sending domain has no prior relationship with the organization and that replies are funneled to a ZoomInfo lookalike domain (zoom-info-us.com), which conflicts with known vendor records. Natural language processing and header analysis surface urgency, invoice language, and unusual thread formatting, elevating the message as high-risk vendor fraud and blocking it before payment is initiated.



OAuth Consent Phishing

Unlike traditional phishing attacks that attempt to steal usernames and passwords, OAuth consent phishing manipulates users into authorizing malicious applications, granting attackers persistent access that bypasses multi-factor authentication.

OAuth (Open Authorization) is a legitimate authorization protocol that allows users to grant third-party applications access to protected resources via tokens rather than login credentials. Attackers exploit this mechanism by abusing users' familiarity with OAuth consent prompts—the "This app would like to..." permission screens that appear when connecting trusted applications. Because these prompts look routine and legitimate, users may approve them without recognizing the risk.

Once a user approves the request, the attacker can obtain token-based access to the account. Depending on the permissions granted, this may allow them to read emails, modify mailbox settings, or send messages as the user. OAuth tokens remain valid until explicitly revoked, which means attackers can maintain persistent access to the account—enabling reconnaissance, data exfiltration, and the potential launch of additional attacks against other users within the organization, even if the user later changes their password.



OAuth consent phishing has emerged as a distinct alternative to traditional credential theft, particularly in environments where MFA has reduced the effectiveness of password-based attacks. As organizations deepen their reliance on cloud-native applications in 2026, this attack class is poised to expand, providing attackers with durable, low-friction access that legacy defenses were never designed to detect.

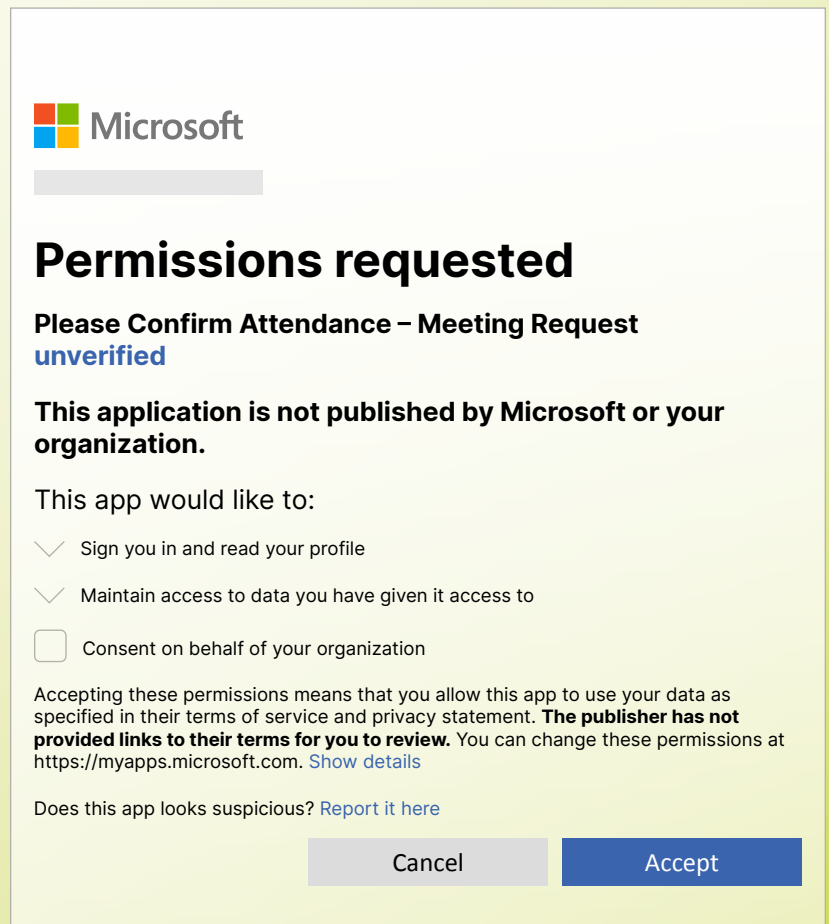
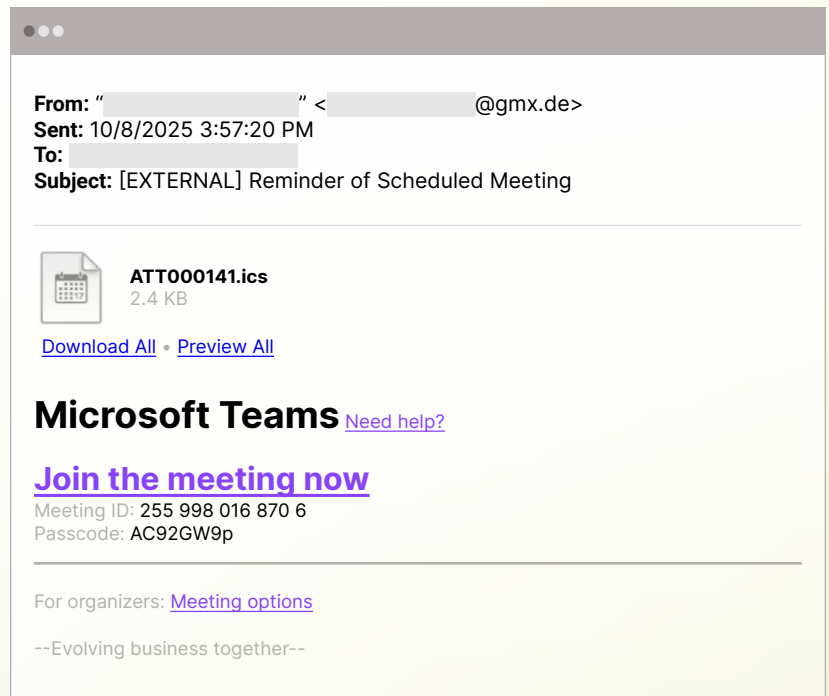


Real-World Example of OAuth Consent Phishing

Meeting invites are among the most trusted messages in a corporate inbox. They're part of the background noise of modern work—routine, ordinary, and rarely questioned. But that familiarity makes them an ideal delivery vehicle for OAuth phishing attacks, as in the example seen here.

The initial email purports to originate from "Teams HR Customer" but is actually sent from an external address hosted by GMX Mail, a free consumer email service based in Germany. It is crafted to mimic a Microsoft Teams meeting invite, with a prominent "Join the meeting now" call-to-action, Meeting ID, Passcode, and a "For organizers: Meeting options" link.

If the target clicks "Join the meeting", they are taken to a compromised Azure Web App that displays a Microsoft-branded OAuth consent page prompting them to authorize an unverified application titled "Please Confirm Attendance – Meeting Request." This fake application requests permissions to sign in and read the user's profile, maintain access to data it has been given access to, and consent on behalf of the target's organization.



Should the user consent, Azure AD issues an authorization code, then an access token and a refresh token, to the attacker-controlled app. With `offline_access`, the attacker can silently refresh tokens without user interaction, bypassing MFA after the initial grant.

What makes OAuth consent phishing so insidious is that malicious applications use the same legitimate authorization flows as trusted apps when requesting permissions. Each element of the OAuth protocol—including the consent prompt and token issuance—functions exactly as designed. The critical difference lies in how the granted permissions are abused after consent is given.

Detecting OAuth Consent Phishing

OAuth consent phishing exploits a blind spot created when legitimate security protocols are assumed to be safe by default. In this attack, the sender uses a legitimate consumer email domain that passes SPF, DKIM, and DMARC, and the email body closely mirrors a standard Microsoft Teams meeting invite with no attachments or obvious phishing language. The embedded link initially routes through a real Microsoft login endpoint before redirecting to a compromised Azure Web App. As a result, signature-based detection and sandboxing tools see only trusted cloud infrastructure, not known credential-harvesting or malware sites. Because the attack hinges on a user granting consent—rather than entering credentials—traditional secure email gateways often fail to recognize the risk unfolding after delivery.

Stopping this attack requires detecting when a legitimate authorization flow is being abused. Abnormal flags that a consumer email account with no prior relationship is impersonating a Microsoft Teams system sender, and that the content, timing, and call-to-action deviate from the recipient's normal communication patterns. It inspects the full URL chain to the compromised Azure Web App and suspicious OAuth consent request from an unverified application, treating these anomalies as high risk and automatically remediating the email.



Lateral Phishing

Lateral phishing attacks occur when threat actors gain access to a legitimate internal email account and use it to target other users within the same organization. After compromising the account—either through a successful phishing attempt or by abusing previously exposed credentials, such as those obtained from a data breach or reused across services—attackers send messages to other internal employees that closely mirror familiar communication styles. These emails are crafted to appear routine and trustworthy, manipulating recipients into sharing sensitive information or fulfilling fraudulent requests.

What makes lateral phishing particularly dangerous is its stealth and built-in credibility. Messages originate from a real internal account, allowing them to bypass many traditional security controls and lowering recipients' natural suspicion. This trust-based exploitation not only increases the success rate of individual phishing attempts but also enables attackers to move laterally through the network, escalate privileges, and expand their access. By operating from within the organization, lateral phishing increases dwell time, complicates detection and response, and heightens the risk of data theft, financial fraud, or downstream compromise.

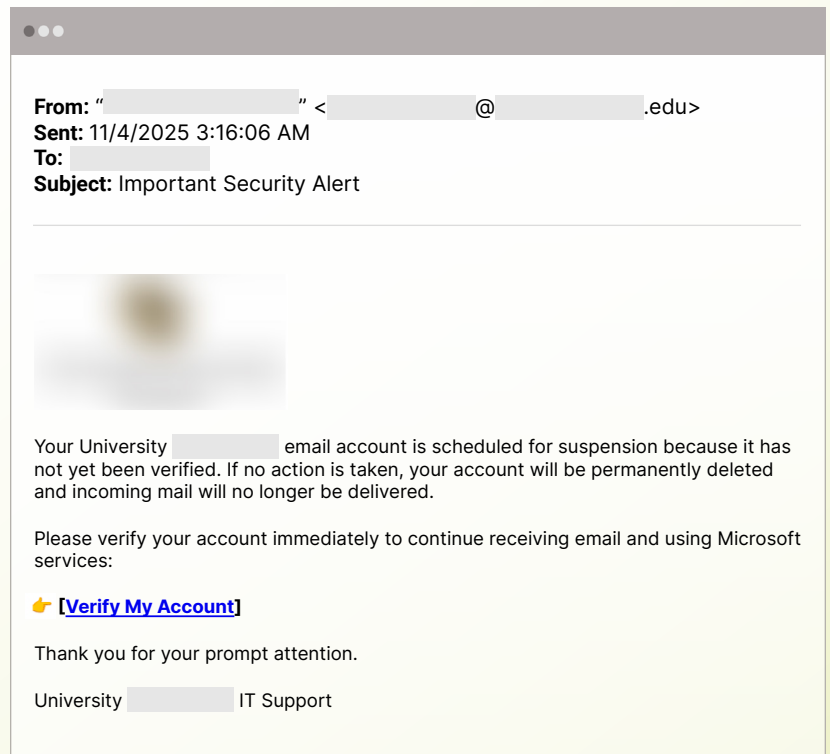


Lateral phishing illustrates a shift in threat actor focus, with attacks that prioritize persistence and post-compromise expansion over initial access alone. This trend is expected to intensify in 2026, as adversaries increasingly leverage trusted internal identities to prolong dwell time and amplify the downstream impact of a single compromised account.

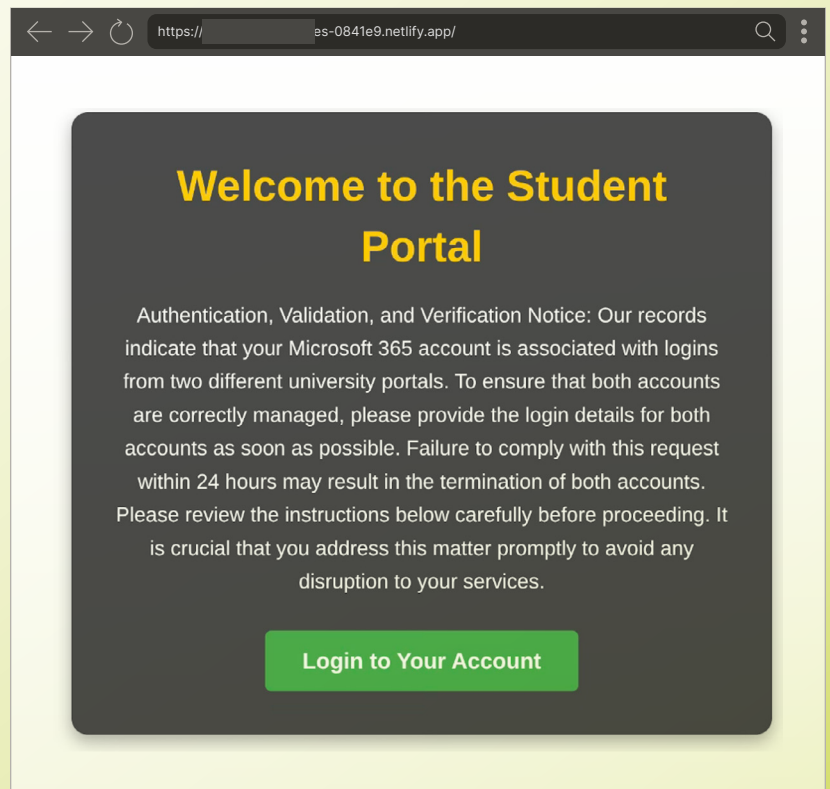


Real-World Example of Lateral Phishing

This lateral phishing attack leverages a compromised university email account to execute a convincing credential-harvesting campaign. Posing as “[Redacted] IT Support” and using the subject line “Important Security Alert,” the message warns the recipient that their email account is scheduled for suspension. The email references Microsoft services and includes institutional branding to enhance legitimacy, pressuring recipients to click a “Verify My Account” link to avoid permanent deletion and loss of incoming mail.



Upon clicking the link, users are directed to a spoofed student login portal splash page that mimics an official university authentication workflow. The notice claims that the user's Microsoft 365 account is associated with two different university portals and instructs them to provide the login credentials for both accounts. The user is warned that failure to do so may result in termination of account access.





Proceeding further leads the target to a Netlify-hosted form designed to harvest sensitive information, including personal email addresses, phone numbers, and—critically—two separate sets of Microsoft 365 credentials under the pretext of validating both “present” and “former” college email accounts.

The use of a legitimate .edu email address, familiar branding, and high-pressure consequences that discourage scrutiny make this lure especially effective. Additionally, by hosting the credential-capture page on Netlify rather than on platforms that prohibit password collection—such as Google Forms—the attacker maximizes the likelihood of data theft.

If successful, the stolen credentials enable account takeover, inbox rule manipulation, and further lateral phishing from newly compromised accounts, significantly increasing dwell time, blast radius, and the risk of broader data exposure or follow-on attacks across the organization.

Detecting Lateral Phishing

Lateral phishing exposes a fundamental weakness in email security solutions that treat internal senders as inherently trustworthy. Because these messages originate from legitimate, authenticated accounts, they pass SPF, DKIM, and DMARC checks and carry no external reputation risk. The content mirrors familiar IT notifications, with accurate branding and service references, leaving static rules and keyword-based detection largely ineffective against this kind of attack. Low sending volume and internal targeting further reduce suspicion, allowing legacy tools to overlook malicious intent simply because the sender is already trusted.

Detecting lateral phishing requires evaluating how an identity behaves—not just whether it is valid. Abnormal models the sender’s historical email behavior and the university’s normal IT communications. It recognizes this user does not typically send mass security notifications, and the “Important Security Alert” template has never been used by this account. Abnormal also inspects the email itself, identifying a never-before-seen external URL that imitates a Microsoft login page and language consistent with credential harvesting rather than prior IT notifications. These behavioral, content, and URL anomalies drive a high risk score, resulting in the email being quarantined.



AI-Generated Payroll Fraud

Payroll fraud is a form of business email compromise (BEC) in which threat actors impersonate employees to redirect paychecks to attacker-controlled bank accounts. These attacks exploit human trust and payroll processes rather than technical vulnerabilities, typically beginning with a message to HR or payroll teams, falsely claiming a need to update direct deposit information, often under the guise of urgency or confidentiality. Because the emails contain no links or attachments to analyze, have no obvious spelling or formatting errors, and closely resemble legitimate internal communications, they often evade both traditional security controls and employee suspicion.

Generative AI has further amplified the effectiveness of these threats by automating both reconnaissance and execution. Attackers can quickly identify the appropriate target, determine which employee to impersonate, and tailor messages to specific organizational roles using AI-driven research tools. They then use generative AI to craft highly personalized, grammatically flawless emails that mirror an individual's tone and communication patterns. The resulting impact is immediate and personal: diverted wages may go unnoticed until a missed paycheck, while organizations face reputational harm, erosion of internal trust, and potential regulatory consequences.

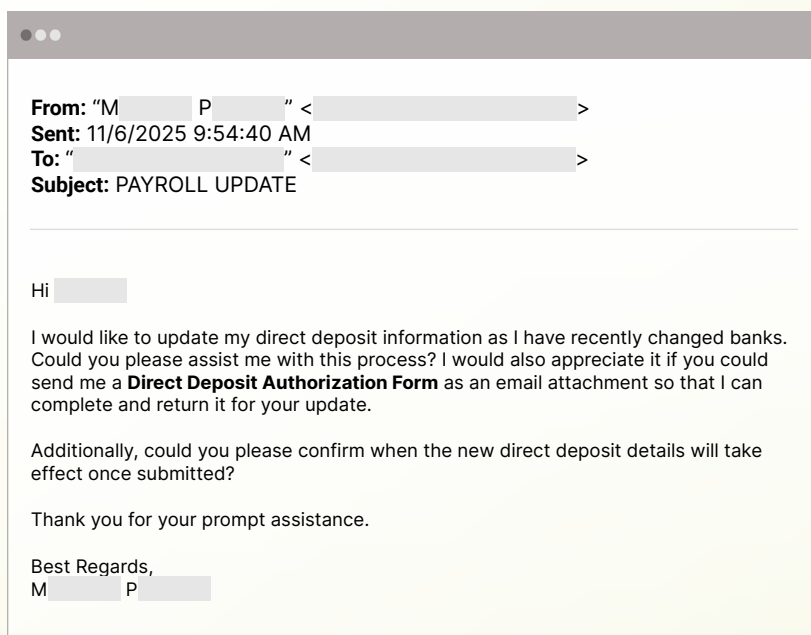


AI-generated payroll fraud reflects the emergence of a more targeted, research-driven form of business email compromise, in which attackers use automation to tailor impersonation at scale. Looking ahead to 2026, this attack type signals a wider trend toward AI-enabled social engineering that delivers high-impact outcomes while leaving little technical evidence for traditional tools to analyze.



Real-World Example of AI-Generated Payroll Fraud

The attack begins with a simple, text-only email sent from a compromised legitimate account. While the sender address itself is unrelated to the organization, the attacker obscures this mismatch by setting the display name to a Senior Vice President at the target company. This subtle manipulation creates the illusion of an internal executive reaching out directly to payroll, immediately establishing authority and familiarity. The subject line, "PAYROLL UPDATE," paired with a personalized greeting ("Hi Ron"), further signals relevance and increases the likelihood of engagement.



The pretext centers on a routine administrative task: the supposed executive claims to have recently changed banks and asks for help updating direct deposit information. This exploits a common, low-friction business process that payroll staff regularly handle, making the request appear normal and non-urgent on the surface.

To advance the fraud, the attacker asks the recipient to send a Direct Deposit Authorization Form as an attachment. This shifts the burden of initiating the process to the target and keeps the interaction within normal payroll procedures. It also gives the attacker a straightforward way to submit fraudulent banking details under the guise of standard documentation.

The email closes by asking when the updated deposit information will take effect, signaling intent to synchronize the change with an upcoming payroll cycle. The message is concise, polite, and free of contextual specifics—traits consistent with AI-generated, reusable templates designed to blend seamlessly into everyday payroll communications.

Detecting AI-Generated Payroll Fraud

AI-generated payroll fraud succeeds because it removes nearly every technical signal traditional defenses rely on. The email contains no links, malware, or suspicious attachments, and the language is polite, concise, and free of urgency cues commonly associated with financial fraud. By impersonating a senior employee and framing the request as a routine payroll change, the message blends seamlessly into normal HR workflows. With nothing overtly malicious to analyze, legacy secure email gateways have few reliable technical indicators to distinguish the request from legitimate internal correspondence.

Abnormal's platform analyzes this message using behavioral and identity-based models rather than static rules. It understands that the real employee typically sends email from a corporate internal account, not an unrelated external account, and has no prior communication history with this payroll contact. The platform flags the external sender, first-time relationship, and sensitive payroll-change request as anomalous for this user and workflow, scoring the message as high risk and remediating it before funds can be redirected.



Predictions for 2026 and Beyond

►► SaaS-to-SaaS Supply Chain Attacks Will Accelerate

Attackers will increasingly compromise one SaaS platform to pivot into another—similar to the SalesDrift and Gainsight-style incidents. As organizations stitch together dozens of interconnected applications, OAuth trust relationships, API permissions, and multi-tenant integrations become high-value targets. Expect attackers to exploit legitimate connectors, app marketplaces, and automated workflows to spread silently across environments.

►► Vendor Compromise Will Become a Top Business Risk

Organizations will see more attacks originating from trusted partners, vendors, and service providers. Compromised inboxes, stolen API tokens, and hijacked customer success platforms will deliver highly targeted, in-context attacks that bypass traditional controls. Generative AI models will amplify this risk by enabling attackers to craft hyper-contextual messages using publicly available data, compromised inboxes, and CRM exports. Trusted relationships—not malicious payloads—will be the primary attack surface.

►► AI-Native Security Solutions Will Become Increasingly Necessary

Unlike legacy systems reliant on static rules, AI-native tools excel at analyzing real-time data, detecting anomalies, and adapting to new attack vectors. These solutions are critical in addressing threats across expanding attack surfaces where speed and precision are paramount. By continuously learning and providing actionable insights, AI-native defenses empower organizations to stay ahead of cybercriminals, mitigating both known and emerging threats with agility and precision.



Defending Against New and Emerging Threats



- ▶ Research reports and stakeholder surveys are increasingly drawing the same conclusion: employees remain the most vulnerable part of an organization's cybersecurity posture. While security awareness training is a critical component of a comprehensive defense strategy, the most effective way to protect your employees from ever-more complex attacks is to prevent malicious emails from reaching them in the first place.

There is little denying that email threats will continue to increase in both volume and severity. However, these attacks can be effectively neutralized with the right solution—one that leverages AI to analyze identity, context, and content and build behavioral baselines for every identity in your cloud environment. By understanding an organization's unique patterns of communication, a robust email security platform can identify and block anomalous messages before they become a threat.

With the right technology in place, you can be confident that your employees are protected from all types of attacks—even those that have yet to be observed.





►► About Abnormal AI

Abnormal AI is the leading AI-native human behavior security platform, leveraging machine learning to stop sophisticated inbound attacks and detect compromised accounts across email and connected applications. The anomaly detection engine leverages identity and context to understand human behavior and analyze the risk of every cloud email event—detecting and stopping sophisticated, socially-engineered attacks that target the human vulnerability.

You can deploy Abnormal in minutes with an API integration for Microsoft 365 or Google Workspace and experience the full value of the platform instantly. Additional protection is available for Slack, Workday, ServiceNow, Zoom, and multiple other cloud applications. Abnormal is currently trusted by more than 3,200 organizations, including over 20% of the Fortune 500, as it continues to redefine how cybersecurity works in the age of AI.

Want to Learn More?

[Request a Demo](#) ►

