

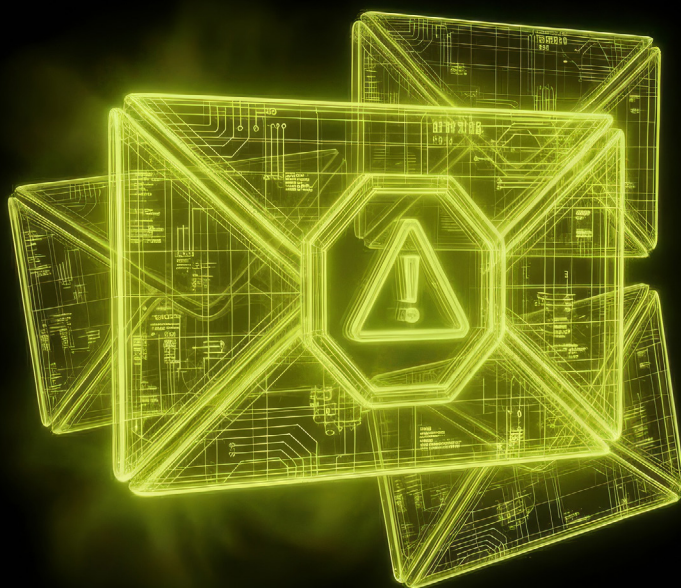
Abnormal

脅威レポート

2026年

脅威の展望：

知っておくべき5つのメール攻撃





エグゼクティブサマリー

個人事業主からグローバル企業に至るまで、あらゆる組織において電子メールは業務運営の主要なコミュニケーションチャネルであり、同時に脅威アクターにとって最も確実な侵入経路でもあります。

攻撃者にとってメールが特に有効である理由は、その普及度の高さだけでなく、「人」という要素へ直接アクセスできる点にあります。従業員は常に受信トレイを確認・処理する必要があり、その過程で心理、親近感、日常的な行動パターンを悪用する機会が無数に生まれます。一見無害に見えるやり取りを通じて攻撃が成立するのです。この構造は、正規のやり取りの「外側」ではなく「内部」に自然に溶け込むことができるサイバー犯罪者に有利に働きます。

本レポートで取り上げる5つの攻撃手法は、技術的脆弱性の悪用から、より労力を要する「アイデンティティ中心型」の脅威へとシフトしている現状を示しています。これらの攻撃は日常的な業務ワークフローを悪用する点が特徴です。マルチステージ型のQRコードフィッシングや、スレッド偽装によるベンダーなりすましは、脅威アクターが多層的な口実、捏造された文脈、そして現実味のある成果物を組み合わせることで標的を徐々に信頼させ、検知を回避する傾向が強まっていることを示しています。OAuth同意フィッシングやラテラルフィッシングは、持続性および侵害後アクセスの価値が高まっていることを示しています。これにより攻撃者は内部環境内に足場を維持しつつ、シグネチャベースのセキュリティ制御を回避できます。さらに、AI生成による給与関連詐欺は、生成AIがビジネスメール詐欺攻撃における情報収集、パーソナライズ、実行を加速させている現状を浮き彫りにしています。最小限の技術的痕跡で高い被害を生み出すことが可能になっています。

これらの脅威の強さは、通常の業務活動に違和感なく溶け込む点にあります。既知の連絡先へのなりすまし、侵害済みアカウントの悪用、信頼されたプラットフォームの武器化により、攻撃メールは構造的にも文脈的にも通常の企業内コミュニケーションと整合する内容になります。その結果、従来の侵害指標の信頼性は損なわれています。

こうした状況により、従来型のセキュアメールゲートウェイなどの従来型セキュリティソリューションでは、悪意ある活動と正規の業務活動を区別することが困難な脅威環境が生じています。既知の指標や明白に不審なシグナルを検出する設計では、日常的かつ信頼性があり、文脈上妥当であるように見せかけた攻撃を阻止できません。このギャップを埋めるには、事前定義されたルールや静的なシグナルに依存するのではなく、やり取り全体にわたるアイデンティティ、行動、文脈を理解し、自動的に異常を検知するAIネイティブな防御が必要です。



目次

2026年に企業リスクを拡大させる5つのメール攻撃 04

- 01 マルチステージ型QRコードフィッシング 05
 - 02 スレッド偽装型ベンダーなりすまし 08
 - 03 OAuth同意フィッシング 11
 - 04 ラテラルフィッシング 14
 - 05 AI生成型給与詐欺 17
-

2026年以降の予測 19

新たに出現する脅威への防御 20

Abnormalについて 21



2026年に企業リスクを拡大させる 5つのメール攻撃



- ▶ メールはビジネスコミュニケーションの基盤であり、業種や地域を問わず広く採用されています。

この広範な依存は、長年にわたりメールを魅力的な攻撃ベクターにしてきました。脅威アクターにとって、戦術を試行し、洗練させ、繰り返し実行できる安定した環境を提供するためです。

現代の攻撃者は、組織のセキュリティ対策を回避し、受信トレイを通じて標的を操作する新たな手法を常に模索しています。単一の技術に依存せず、新たな口実、形式、配信手段を試しながらアプローチを適応させ、目的を達成するまで実験を重ねます。

以下は、2025年にAbnormal AIの顧客が実際に受信した脅威の事例です。これらは、攻撃者が現在進行形でこれらの手法を活用していることを示すとともに、2026年に向けて組織が検知・防御に備えるべき戦略を浮き彫りにしています。



企業リスクを拡大させる5つのメール攻撃 \ 01

マルチステージ型 QRコードフィッシング

2010年代前半には、悪意のあるQRコードが至る所で確認され、脅威アクターはQRコードフィッシング攻撃を積極的に展開していました。当初、これらの攻撃は単純な口実に依存し、手順も最小限でした。一般的には、攻撃者が悪意あるQRコードを含む単一のメールを送信し、標的がそれをスキャンすると、MicrosoftやGoogleのポータルを装った偽のログインページへ直接誘導され、認証情報の入力を求められるという流れでした。

しかし近年、脅威アクターは悪意あるQRコードをより複雑な形で利用するようになってきました。マルチステージ型QRコードフィッシングでは、攻撃フローに追加の段階を組み込み、正当性をより強く演出するとともに検知を回避しています。従来のQRコードフィッシング攻撃が主に不正な多要素認証の有効期限通知や共有ドキュメント通知（Abnormalが検知したQRコード攻撃の約50%を一時期占めていました）を口実としていたのに対し、マルチステージ型では多様な口実が活用されています。

その結果、標的が最終的な認証情報窃取ページに到達するまでには、複数回のやり取りを通じてワークフローへの信頼が形成されます。最終ページには、企業ロゴの表示やメールアドレスの事前入力といった高度なパーソナライズが施されることが多いです。この多層的アプローチにより、標的の疑念を最小限に抑えるだけでなく、既知の悪性ドメインへの直接リダイレクトを検知するセキュリティ制御も回避されます。



単一スキャン型の単純な誘導から、複数段階のワークフローへと進化してきたこの継続的な戦略高度化は、ユーザーの警戒心および従来型の検知手法の双方を巧みに回避することを目的とした、計画的かつ高労力型のフィッシングキャンペーンへの広範な移行を反映しています。この傾向は、2026年にさらに加速すると予想されます。

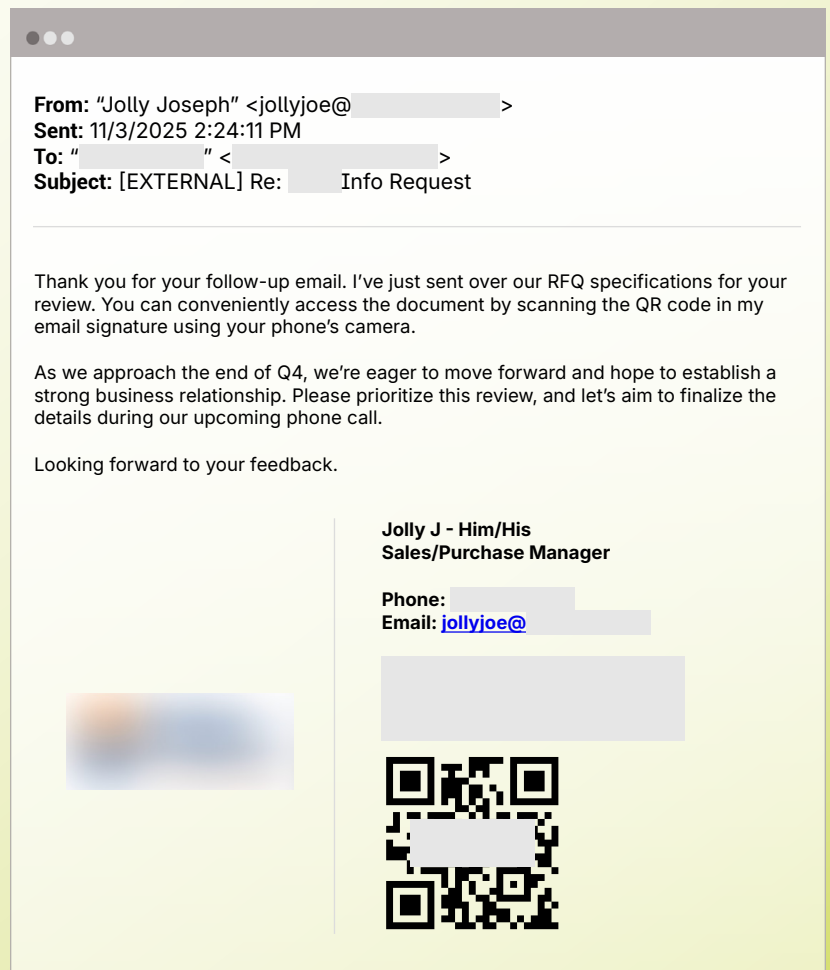
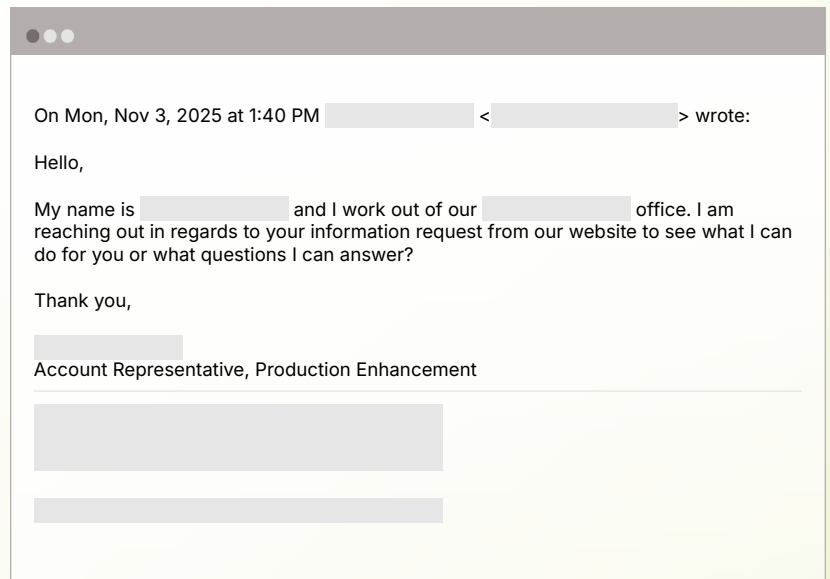


マルチステージ型 QRコードフィッシングの実例

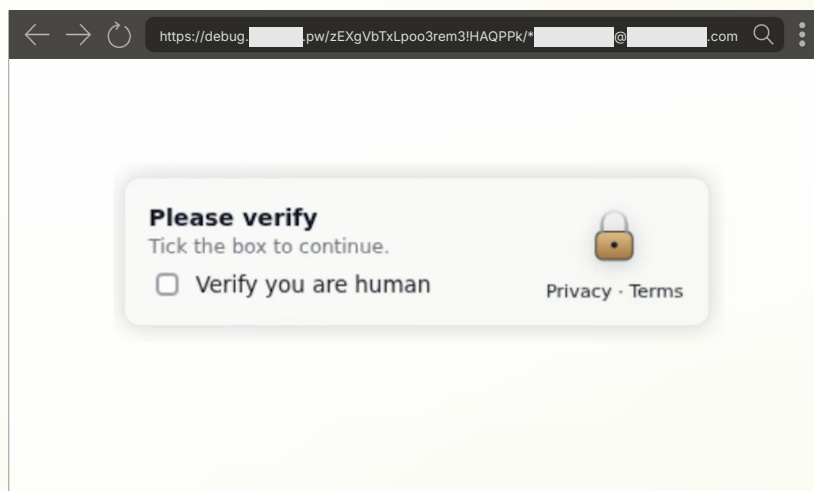
この攻撃では、脅威アクターが標的組織のWebサイト上のフォームを通じて情報提供依頼を送信することで接触を開始しています。この送信により正規の従業員が返信を行い、攻撃者が後に悪用できる実在のメールスレッドが形成されます。

返信の中で、攻撃者はベンダー側の担当者になります。メッセージでは見積依頼(RFQ)に関する口実を提示し、四半期末の緊急性を強調することで、受信者に優先対応を促します。正当性を高めるために、攻撃者は最近登録された正規ドメインに酷似したドメインを使用し、公式に見えるブランド表示、企業住所、そしてプロフェッショナルな署名要素を組み込みます。

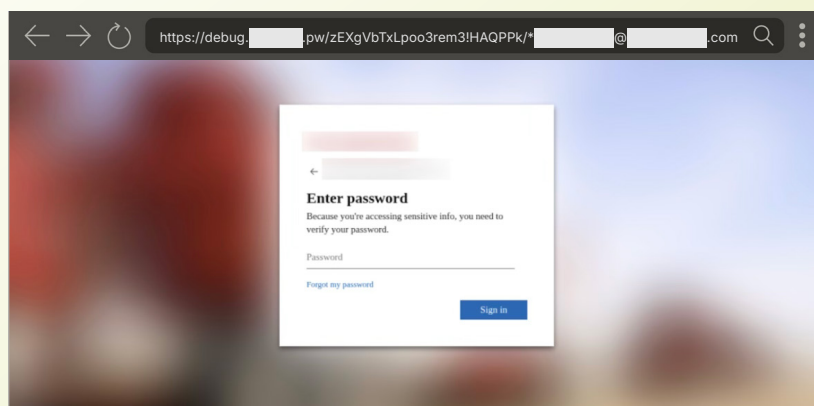
攻撃者は、「お使いのスマートフォンのカメラでスキャンしてください」と指示し、RFQ仕様書にアクセスするためにQRコードを読み取るよう誘導します。これによりやり取りはモバイル端末へと移行し、企業のメール保護の外に置かれることになります。その結果、リンクスキャンなどのセキュリティ制御を回避することが可能になります。



QRコードをスキャンすると、被害者は「あなたが人間であることを確認してください」というページへリダイレクトされます。最終的なURLには標的のメールアドレスがパスの一部として含まれており、このフローが個別にパーソナライズされ、後続のコンテンツを事前入力する目的で使用されている可能性が高いことを示しています。



この認証ステップの後、標的には自社の認証ポータルを模倣したブランド付きのログインページが表示されます。ページには正しいメールアドレスがあらかじめ入力された状態で表示され、パスワードの入力が求められます。標的が情報を送信した場合、攻撃者は認証情報を入手できるようになります。



マルチステージ型QRコードフィッシングの検知

従来型のセキュアメールゲートウェイがマルチステージ型QRコードフィッシングを検知することは困難です。その理由は、単にメールが正当に見えるからではなく、攻撃が意図的に可視性を複数の環境に分散しているためです。メール自体は既存スレッド内での通常のRFQフォローアップのように見え、プロフェッショナルなブランディングやもっともらしいベンダーの身元情報も含まれているため、従来型スキャナが分析できる要素はほとんどありません。実質的に確認可能な要素は署名内に埋め込まれたQRコードのみであり、URLベースや添付ファイルベースのスキャナでは検査対象がほとんど存在しません。さらに、QRコードがスキャンされると、やり取りは即座にモバイル端末へ移行し、リダイレクトや認証情報窃取は企業メールの境界外で実行されるため、従来型ツールでは攻撃全体のシーケンスを把握することができません。

Abnormalは、この脅威を単独のメールイベントとしてではなく、より広範なマルチステージ攻撃の最初の段階として捉えることで検知します。行動分析型AIにより、通常のベンダー関係やコミュニケーションパターンを理解し、新規送信者が最近登録された類似ドメインを使用している点や、不自然なアクションを促している点などの異常を特定できます。また、AbnormalはQRコードをデコードし、遷移先URLを解析し、それらの結果を文脈的シグナルと相関分析します。アイデンティティ、行動、ワークフロー全体にわたる逸脱を検出することで、当該メッセージを悪性と高精度で分類し、従業員が関与する前に攻撃を阻止することが可能です。

企業リスクを拡大させる5つのメール攻撃 // 02

スレッド偽装型 ベンダーなりすまし

スレッド偽装型ベンダーなりすまし攻撃では、脅威アクターが、ベンダーと社内従業員との間で行われた正規のやり取りのように見せかけた捏造メールスレッドを使用し、信頼できる口実を構築します。攻撃者は最初のメールにこの偽のメッセージチェーンを挿入し、自身の問い合わせが正当であることを裏付ける証拠として提示します。内容は通常、銀行口座情報の更新や未払い請求書の解決など、財務ワークフローに関連するものです。スレッド内では、なりすまされた従業員が、なりすまされたベンダーの要求を進める許可を与えているように見せかけます。

多くの場合、攻撃者がなりすます従業員は権限を持つ立場の人物であり、受信者に対応を促す心理的圧力を生み出します。また、架空のやり取りは、十分に具体的に関連性があるように感じられる一方で、不正を示唆する不自然な兆候を生まないよう、あえて汎用的な内容にとどめられています。正当性を装うために、攻撃者は類似ドメインを使用するほか、可能な場合には侵害された正規ドメインを悪用してブロックリストを回避します。さらに、改ざんされた請求書や加工された税務書類などの偽造文書を添付し、標的に取引を完了させるよう欺きます。



スレッド偽装型ベンダーなりすましは、金融詐欺の手法における顕著な進化を示しています。攻撃者はより多くの労力をかけて信頼性の高い口実を構築し、攻撃件数よりもリアリティを重視する傾向にあります。2026年には、このような口実重視型の攻撃がベンダーなりすましキャンペーンにおいてさらに顕著になると予想されます。これは、技術的脆弱性の悪用よりも、信頼された業務プロセスを操作する手法が優先されていることを反映しています。



スレッド偽装型 ベンダーなりすましの実例

スレッド偽装型ベンダーなりすましの本質は、典型的なソーシャルエンジニアリングです。人間の心理を悪用し、信頼された当事者を模倣し、説得力のある口実を用いて標的に有害な意思決定をさせます。以下の事例は、これらの要素がどのように機能するかを示す優れた例です。

このメールは侵害された正規のアドレスから送信されており、著名なビジネスインテリジェンスプラットフォームである ZoomInfo からの連絡であるかのように装っています。件名は転送メールを装った形式「Fw: Prospect Intelligence Platform - Discover Hidden Opportunities」となっており、自然なやり取りの一部であるかのように見せかけています。攻撃者は、「P.Z.」と記載されたCEOから未払い請求書の送付依頼を受けたと説明し、受信者に対して「詳細は以下のやり取りをご確認ください」と案内します。

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



スレッド偽装型ベンダーなりすまし攻撃に典型的に見られるように、文面は曖昧であり、即時の支払いへと誘導する内容になっています。本文では「できるだけ早く支払いが処理されるよう、本件を速やかに解決したいと考えております」や「サービスの中断を回避するため、迅速なご対応をお願いいたします」といった表現を用いて緊急性を強調しています。この攻撃では、2つの添付ファイルが含まれていました。1つはACH送金による約5万ドルの支払いを求める請求書、もう1つはベンダー確認用として提示されたW-9フォームです。

標的が請求書の処理を進めてしまった場合、知らないうちに数万ドルもの資金を脅威アクターへ直接送金してしまうことになります。

スレッド偽装型ベンダーなりすましの検知

財務ワークフローにおいては、「見慣れていること」や「承認済みに見えること」が正当性のシグナルとして扱われることが多くありますが、スレッド偽装型ベンダーなりすま시는まさにこの前提を悪用します。多くの従来型セキュリティツールは、このメールを攻撃としてフラグ付けできない可能性があります。なぜなら、このメッセージは信頼性のあるプロバイダーからの通常の請求書フォローアップのよう見え、もっともらしい件名、プロフェッショナルな署名、そしてマルウェアや明白に悪意のあるリンクを含まない請求書およびW-9が添付されているためです。さらに、このメールは正規ドメイン(govquest.com)から送信されており、経営層による実在の社内承認スレッドのよう見える内容も含まれています。そのため、シグネチャベースやレピュテーションベースのフィルタリングではリスクが低いと判断されやすく、不正な支払い依頼がそのまま受信トレイに届いてしまいます。

この攻撃を検知するには、メッセージの真正性だけでなく、業務上の意図を検証する必要があります。Abnormalは、確立された財務ワークフローの文脈でベンダーとの通信を評価し、既知の取引関係、支払い行動、承認パターンをモデル化しています。その結果、送信ドメインが当該組織と過去に取引関係を持たないこと、返信先が ZoomInfo の類似ドメイン(zoom-info-us.com)に誘導されていることが、既知のベンダー情報と整合しない点を特定します。さらに、自然言語処理およびメールヘッダー分析により、過度な緊急性、請求書関連の文言、不自然なスレッド形式などを検出し、本メールを高リスクなベンダー詐欺として分類します。これにより、支払い処理が開始される前に攻撃をブロックすることが可能になります。



OAuth同意フィッシング

従来のフィッシング攻撃がユーザー名やパスワードを盗み取ることを目的としているのに対し、OAuth同意フィッシングは、悪意のあるアプリケーションをユーザーに承認させることで攻撃を行います。これにより、攻撃者は多要素認証を回避し、継続的なアクセス権を取得します。

OAuth (Open Authorization) は、正規の認可プロトコルであり、ログイン認証情報の代わりにトークンを使用して、ユーザーがサードパーティ製アプリケーションに保護されたリソースへのアクセスを許可できる仕組みです。

攻撃者は、ユーザーが日常的に目にしているOAuthの同意画面（例：「このアプリは次の権限を要求しています…」と表示される許可画面）への慣れを悪用します。これらの画面は通常の正規プロセスと見分けが付きにくいいため、ユーザーはリスクに気づかないまま承認してしまう可能性があります。ユーザーがリクエストを承認すると、攻撃者はトークンベースでアカウントへアクセスできるようになります。付与された権限の内容によっては、メールの閲覧、メールボックス設定の変更、ユーザーになりすましてのメッセージ送信などが可能になります。

OAuthトークンは明示的に取り消されない限り有効な状態が継続します。そのため、ユーザーが後からパスワードを変更した場合でも、攻撃者はアカウントへのアクセスを維持できます。これにより、情報収集（偵察活動）、データの不正持ち出し、さらには組織内の他ユーザーに対する追加攻撃の実行が可能となる恐れがあります。



OAuth同意フィッシングは、従来の認証情報窃取とは異なる新たな攻撃手法として台頭しています。特に、多要素認証 (MFA) の導入によってパスワードベースの攻撃の効果が低下している環境において、その存在感を強めています。2026年に向けて、企業がクラウドネイティブアプリケーションへの依存をさらに深める中、この種の攻撃は今後拡大することが予想されます。本来は想定されていなかった従来型の防御対策では検知が難しく、攻撃者にとっては持続的かつ低負荷でアクセスを維持できる手段となる可能性があります。



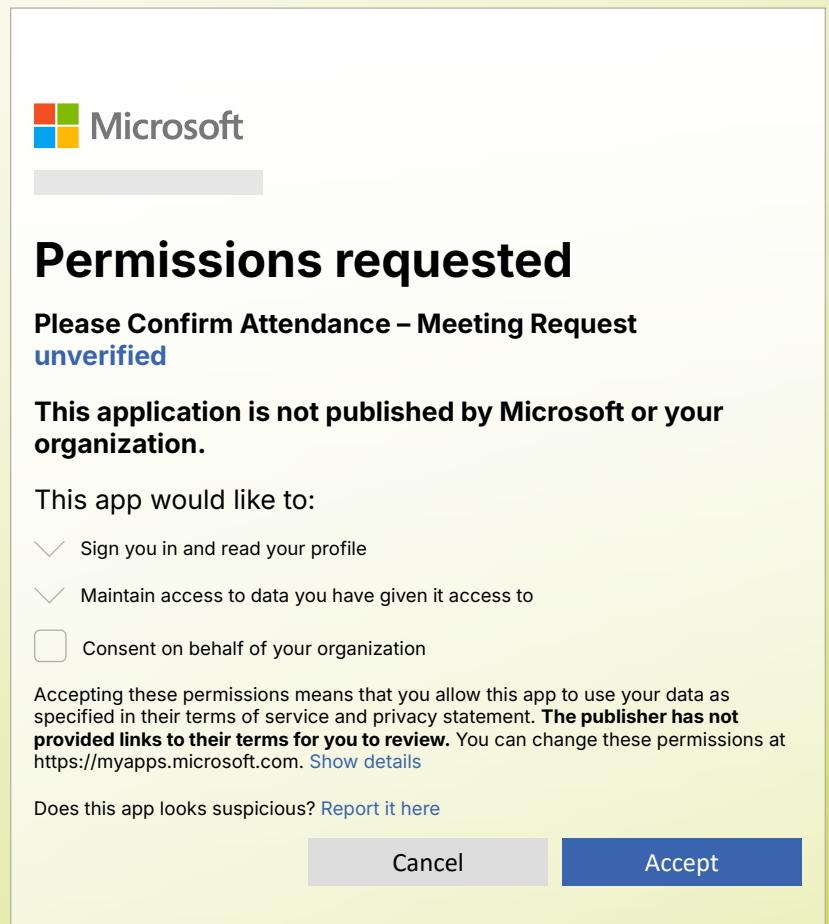
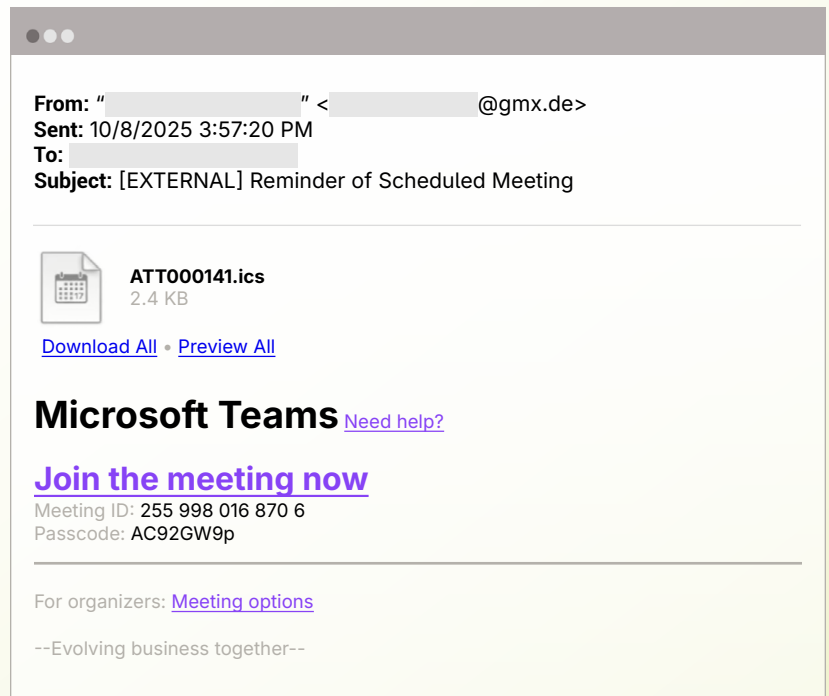
OAuth同意フィッシングの実例

会議招待メールは、企業の受信トレイにおいて最も信頼されやすいメッセージの一つです。現代の業務環境では日常的にやり取りされるものであり、通常は特に疑われることはありません。しかし、その「当たり前さ」こそが、OAuthフィッシング攻撃の格好の配信手段となります。本例はその典型的なケースです。

最初のメールは「Teams HR Customer」から送信されたように見えますが、実際にはドイツに拠点を置く無料の個人向けメールサービスである GMX Mail の外部アドレスから送信されています。メールは Microsoft Teams の会議招待を装って巧妙に作成されており、「Join the meeting now (今すぐ会議に参加)」という目立つCTA (行動喚起) ボタン、Meeting ID、Passcode、「For organizers: Meeting options (主催者向け: 会議オプション)」リンクなどが含まれています。

ターゲットが「Join the meeting now」をクリックすると、侵害された Azure Web Apps 上のWebアプリケーションに誘導されます。そこではMicrosoftブランドを装ったOAuth同意ページが表示され、「Please Confirm Attendance – Meeting Request」という未検証のアプリケーションへの認可を求められます。

この偽アプリケーションは、ユーザーとしてサインインする権限、ユーザープロフィールを読み取る権限、付与されたデータへの継続的なアクセス権限、対象組織を代表して同意を行う権限を要求します



ユーザーが同意すると、Azure Active Directory (Azure AD) は攻撃者が制御するアプリケーションに対して、認可コードを発行し、その後アクセストークンおよびリフレッシュトークンを発行します。offline_access 権限が付与されている場合、攻撃者はユーザーの操作なしにトークンを自動更新できます。これにより、最初の認可後は多要素認証 (MFA) を実質的に回避した状態でアクセスを継続することが可能となります。OAuth同意フィッシングが特に巧妙である理由は、悪意のあるアプリケーションが、信頼された正規アプリと同一の認可フローを使用して権限を要求する点にあります。OAuthプロトコルの各要素 (同意画面やトークン発行プロセスを含む) は、仕様どおり正常に機能します。決定的な違いは、ユーザーの同意後に付与された権限が悪用される点にあります。

OAuth同意フィッシングの検知

OAuth同意フィッシングは、「正規のセキュリティプロトコルは安全である」という前提により生じる盲点を悪用する攻撃です。この攻撃では、送信元として正規の個人向けメールドメインが使用され、SPF・DKIM・DMARCの検証を通過します。メール本文も標準的な Microsoft Teams の会議招待に酷似しており、添付ファイルや明らかなフィッシング文言は含まれていません。

埋め込まれたリンクは、最初に正規のMicrosoftログインエンドポイントを経由し、その後、侵害された Azure Web App へリダイレクトされます。そのため、シグネチャベースの検知やサンドボックス型のセキュリティツールからは、既知の認証情報窃取サイトやマルウェア配布サイトではなく、信頼されたクラウド基盤のみが確認されることになります。さらに、この攻撃は「認証情報の入力」ではなく「ユーザーによる同意付与」に依存しているため、従来型のセキュアメールゲートウェイでは、配信後に進行するリスクを正しく認識できない場合があります。

本攻撃を阻止するには、正規の認可フローが悪用されている兆候を検知する必要があります。Abnormal AI は、過去に関係のない個人向けメールアドレスがMicrosoft Teamsのシステム送信者を装っている点を検出します。また、メールの内容、送信タイミング、CTA (行動喚起) が受信者の通常のコミュニケーションパターンから逸脱していることを特定します。さらに、侵害されたAzure Web AppまでのURLチェーン全体や、未検証アプリケーションからの不審なOAuth同意リクエストを解析し、これらの異常を高リスクとして評価します。その上で、当該メールを自動的に隔離・修復 (自動対応) します。



企業リスクを拡大させる5つのメール攻撃 // 04

ラテラルフィッシング

ラテラルフィッシング攻撃は、攻撃者が正規の内部メールアカウントにアクセスし、そのアカウントを使って同じ組織内の他のユーザーを標的にする手法です。

攻撃者は、フィッシングの成功や、過去のデータ漏えいなどで入手した認証情報の悪用、または複数サービスで使い回されたパスワードを通じてアカウントを侵害します。その後、他の内部従業員に対して、普段のコミュニケーションスタイルに酷似したメールを送信します。これらのメールは日常的で信頼できる内容に見えるよう巧妙に作られており、受信者を騙して機密情報の提供や不正な要求の実行を促します。

ラテラルフィッシングが特に危険なのは、その巧妙さと信頼性にあります。メールが実際の内部アカウントから送信されるため、多くの従来型セキュリティ対策をすり抜け、受信者の警戒心も低下します。

この信頼を悪用した攻撃により、個別のフィッシング成功率が高まるだけでなく、攻撃者はネットワーク内を横方向に移動（ラテラルムーブメント）し、権限を昇格させ、アクセス範囲を拡大することが可能となります。

組織内部から操作が行われるため、検知や対応が難しくなり、滞留時間 (dwell time) が増加します。その結果、データ窃取、金銭詐欺、二次的なシステム侵害などのリスクが高まります。



ラテラルフィッシングは、攻撃者の焦点が初期アクセスの獲得だけでなく、持続的なアクセス維持と侵害後の権限拡大に移行していることを示しています。この傾向は2026年にさらに強まると予想されます。攻撃者は信頼された内部アカウントを巧みに利用し、滞留時間 (dwell time) を延ばすとともに、単一の侵害アカウントから生じる二次的被害を拡大させることを狙います。



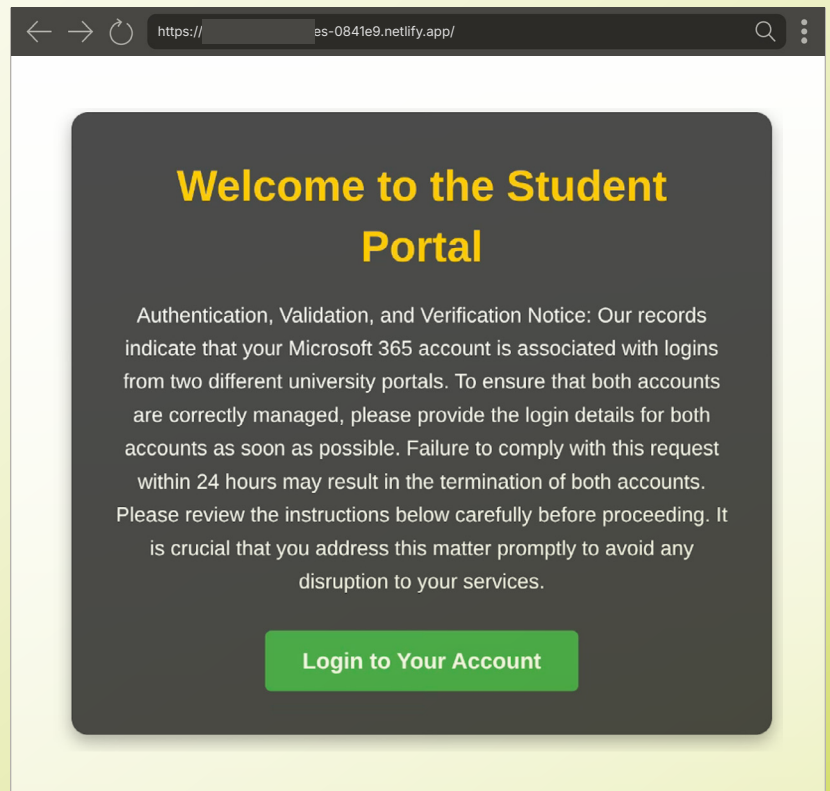
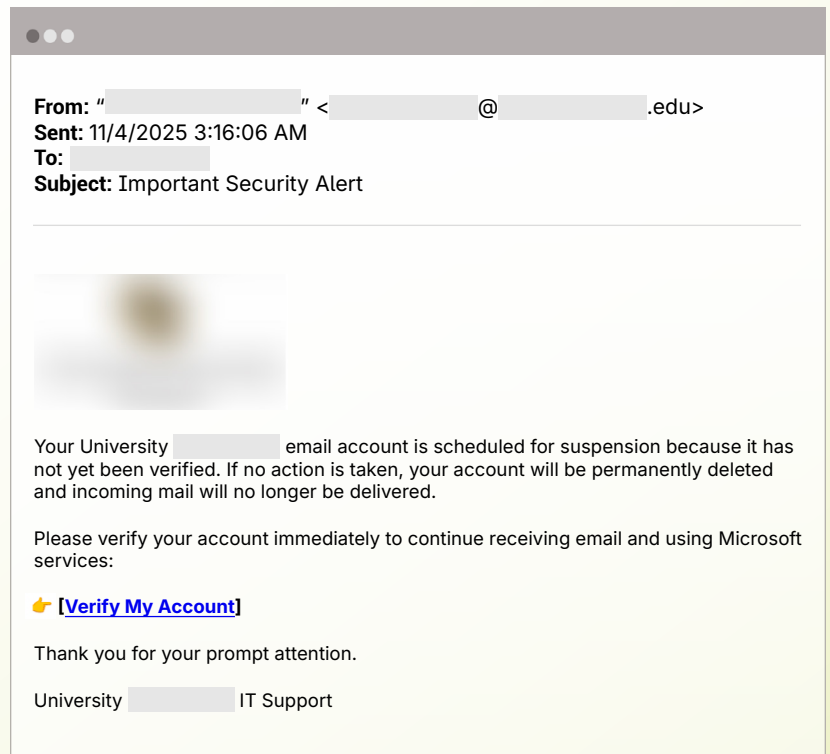
ラテラルフィッシングの実例

このラテラルフィッシング攻撃では、侵害された大学のメールアカウントが悪用され、巧妙な認証情報窃取キャンペーンが実行されました。

送信者は「[Redacted] IT Support」を装い、件名には「Important Security Alert」と記載されています。メールでは、受信者のメールアカウントが停止予定であると警告しています。また、Microsoftのサービスに言及し、大学の公式ブランド要素を使用することで正当性を強調しています。

受信者には、アカウントの永久削除や受信メールの喪失を回避するために、「Verify My Account」リンクをクリックするよう強い圧力がかけられます。

リンクをクリックすると、ユーザーは大学の公式認証フローを模倣した偽の学生ログインポータルのスプラッシュページに誘導されます。そこでは、ユーザーの Microsoft 365 アカウントが2つの大学ポータルに関連付けられていると説明され、両方のアカウントのログイン認証情報を入力するよう指示されます。さらに、必要な情報を提供しなかった場合、アカウントへのアクセスが停止される可能性があるかと警告されます。





さらに操作を進めると、ターゲットは Netlify 上にホストされたフォームに誘導されます。このフォームは、個人のメールアドレスや電話番号などの機密情報、そして特に「現行」と「過去」の大学メールアカウント用の2セットの Microsoft 365 認証情報を収集するように設計されています。

正規の .edu メールアドレス、馴染みのあるブランド、そして内容の精査をためらわせる高圧的な警告メッセージの組み合わせにより、この手口は非常に効果的です。さらに、Google フォームのようにパスワード収集を禁止しているプラットフォームではなく Netlify を使用することで、攻撃者は認証情報窃取の成功可能性を最大化しています。

攻撃が成功した場合、取得した認証情報により以下の被害が発生する可能性があります。

アカウントの乗っ取り、受信箱ルールの改ざん、新たに侵害されたアカウントを利用したラテラルフィッシングの拡大

これにより、滞留時間 (dwell time) が大幅に延長され、攻撃の影響範囲 (blast radius) が拡大し、組織全体でのデータ漏洩リスクや二次攻撃のリスクが増大します。

ラテラルフィッシングの検知

ラテラルフィッシングは、内部送信者を本質的に信頼できると扱うメールセキュリティソリューションの根本的な弱点を突きます。これらのメールは、正規の認証済みアカウントから送信されるため、SPF、DKIM、DMARC の検証を通過し、外部からの評判リスクありません。さらに、内容は馴染みのある IT 通知に酷似しており、正確なブランドやサービス名が使用されているため、静的ルールやキーワードベースの検知はほとんど効果を発揮しません。送信量が少なく、内部に限られた対象に送られることも疑念をさらに低下させ、送信者がすでに信頼されているという理由だけでレガシーツールが悪意を見逃す要因となります。

ラテラルフィッシングを検知するには、アカウントの有効性だけでなく、アイデンティティの振る舞いを評価する必要があります。Abnormal AI は、送信者の過去のメール行動および大学の通常の IT コミュニケーションをモデル化します。このユーザーは通常、大量のセキュリティ通知を送信せず、「Important Security Alert」のテンプレートもこのアカウントでは使用されることがありません。

さらに、Abnormal はメール自体を解析し、これまで使用されなかった外部 URL (Microsoft ログインページを偽装) や、以前の IT 通知とは異なる認証情報窃取を意図した文言を特定します。これらの振る舞い、コンテンツ、URL の異常を組み合わせると高リスクと判断し、メールは自動的に隔離されます。



AI生成による給与詐欺

給与詐欺は、ビジネスメール詐欺（BEC）の一種で、攻撃者が従業員になりすまし、給与振込を攻撃者が管理する銀行口座に誘導する手口です。

これらの攻撃は技術的な脆弱性を突くのではなく、人間の信頼と給与処理の仕組みを悪用します。通常、HRや給与チーム宛てに送信され、「口座情報の更新が必要」といった虚偽の主張を含むメールから始まります。メールは緊急性や機密性を装うことが多く、リンクや添付ファイルは含まれず、スペルミスや形式上の不自然さもほとんどありません。そのため、従来型のセキュリティ対策や従業員の警戒を容易にすり抜けます。

生成AIによるリスクの増大

生成AIの登場により、これらの攻撃はさらに効果的になっています。攻撃者はAIを活用して適切なターゲットの特定、なりすます従業員の選定、組織内の役割に応じたメッセージ内容の調整を自動化します。

さらに、生成AIを用いることで、個人の口調やコミュニケーションパターンを忠実に模した、高度にパーソナライズされ文法的にも完璧なメールを作成できます。

その結果、被害は迅速かつ個人的なものとなります。給与が誤って振り込まれても、支給日まで気づかれないことがあり、組織にはブランド・評判へのダメージ、社内信頼の低下、規制上の問題・法的リスクが生じます。



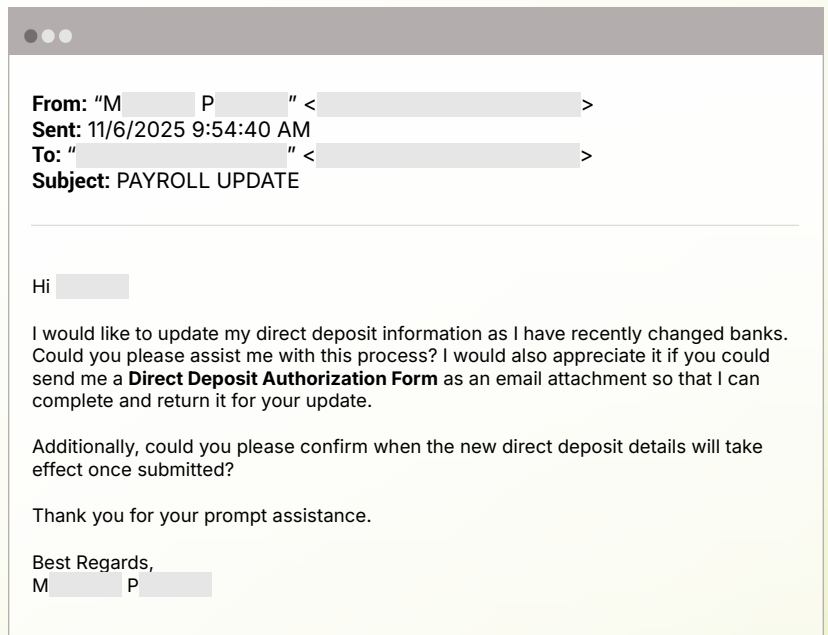
生成AIを用いた給与詐欺は、より標的を絞ったリサーチ主導型のビジネスメール詐欺（BEC）の出現を示しています。攻撃者は自動化を活用し、大規模に個人になりすます手口を精緻化しています。2026年に向けて、この攻撃はAIを活用したソーシャルエンジニアリングの拡大傾向を示しており、従来のツールでは分析が困難な形で大きな被害をもたらすことが予想されます。



AI生成による給与詐欺の実例

本攻撃は、侵害された正規のアカウントから送信されたシンプルなテキストのみのメールから始まります。送信元アドレス自体は対象組織とは無関係ですが、攻撃者は表示名（ディスプレイネーム）を標的企業のシニア・バイスプレジデントに設定することで、この不一致を隠しています。

このさりげない操作により、社内の経営幹部が給与担当に直接連絡しているかのような印象を与え、即座に権威性と親近感を確立します。件名「PAYROLL UPDATE」と、個別名を用いた挨拶（例：「Hi Ron」）は関連性を強調し、受信者が反応する可能性を高めます。



口実は日常的な事務手続きです。メールでは、銀行を最近変更したため、給与の振込先口座情報を更新してほしいと依頼します。これは給与担当者が日常的に処理している一般的な業務であり、表面的には緊急性も低く自然な依頼に見えます。

さらに詐欺を進行させるため、攻撃者は「Direct Deposit Authorization Form（口座振替承認書）」を添付で送付してほしいと求めます。これにより、手続き開始の主体を受信者側に置き、やり取りを通常の給与業務プロセス内にとどめます。同時に、標準的な書類提出を装って不正な銀行口座情報を提出するための明確な経路を確保します。

メールの最後では、更新された振込情報がいつ反映されるのかを確認しています。これは、次回の給与支給サイクルに合わせて変更を同期させる意図を示しています。

メッセージは簡潔で丁寧、具体的な背景情報は含まれていません。これらの特徴は、日常の給与関連コミュニケーションに自然に溶け込むよう設計された、AI生成の再利用可能なテンプレートに典型的なものです。

AI生成給与詐欺の検知

AI生成給与詐欺が成功する理由は、従来の防御策が依存するほぼすべての技術的シグナルを排除している点にあります。メールにはリンクやマルウェア、不審な添付ファイルはなく、文面も丁寧に簡潔、金融詐欺にありがちな緊急性を示す文言も含まれていません。さらに、管理職クラスの社員になりすまし、給与の通常の変更手続きとして依頼を装うことで、メッセージはHRの通常ワークフローに自然に溶け込みます。そのため、目立った悪意が分析対象として存在せず、従来型のセキュアメールゲートウェイでは、正規の内部連絡と区別するための信頼できる技術的指標がほとんどありません。

Abnormal AI のプラットフォームは、静的ルールではなく振る舞いとアイデンティティに基づくモデルを用いてメールを分析します。実際の従業員は通常、社内の正規アカウントから送信しており、無関係な外部アカウントから送信されることはなく、当該給与担当者との過去のコミュニケーション履歴がないことを評価の基準として使用します。

このため、外部送信者、初めての関係、機密性の高い給与変更依頼は、ユーザーとワークフローに対して異常としてフラグが立てられます。結果としてメールは高リスクと判定され、資金が誤って振り込まれる前に自動的に対応（隔離・修復）されます。



2026年以降の予測

▶▶ SaaS間サプライチェーン攻撃の加速

攻撃者は、1つのSaaSプラットフォームを侵害し、それを足がかりに別のSaaSへと横展開する手法をさらに強化していくと予想されます。

これは、SalesDriftやGainsight型のインシデントと同様の手口です。

企業が多数のアプリケーションを相互接続する中で、OAuthによる信頼関係、API権限、マルチテナント統合は高価値な攻撃対象となります。

攻撃者は、正規のコネクタ、アプリマーケットプレイス、自動化ワークフローなどを悪用し、環境全体へ静かに拡散していくことが想定されます。

▶▶ ベンダー侵害が主要な経営リスクに

組織は、信頼しているパートナー企業、ベンダー、サービスプロバイダーを起点とした攻撃の増加に直面することになります。

侵害されたメールアカウント、窃取されたAPIトークン、乗っ取られたカスタマーサクセス基盤などを通じて、従来のセキュリティ対策を回避する、高度に標的化され、文脈に沿った攻撃が実行されます。

さらに、生成AIモデルはこのリスクを一層拡大させます。公開情報、侵害済みメールボックス、CRMデータなどを活用することで、極めて高度なコンテキスト型メッセージの作成が可能になります。今後は「悪意あるペイロード」ではなく、「信頼関係そのもの」が主要な攻撃対象（アタックサーフェス）となります。

▶▶ AIネイティブ型セキュリティの必要性の高まり

静的ルールに依存する従来型システムとは異なり、AIネイティブ型のセキュリティツールは、リアルタイムデータの分析、異常検知、新たな攻撃手法への適応に優れています。

攻撃対象領域が拡大し続ける環境では、迅速性と精度が極めて重要です。AIネイティブ型ソリューションは、継続的な学習と実用的なインサイト提供を通じて、組織がサイバー犯罪者に先手を打つことを可能にします。

新たな脅威と 進化する脅威への防御



- ▶ 各種調査レポートやステークホルダー調査では、共通した結論が示されています。それは、従業員が依然として組織のサイバーセキュリティ体制における最も脆弱な部分であるという点です。
セキュリティ意識向上トレーニングは包括的な防御戦略の重要な要素ですが、ますます高度化・複雑化する攻撃から従業員を守る最も効果的な方法は、悪意のあるメールをそもそも受信させないことです。

メールを起点とする脅威は、今後も量・深刻度の両面で増加し続けると考えられます。しかし、適切なソリューションを導入することで、これらの攻撃を効果的に無効化することが可能です。

具体的には、AIを活用してアイデンティティ、コンテキスト、コンテンツを分析し、クラウド環境内のすべてのアイデンティティに対する行動ベースラインを構築する仕組みが求められます。

組織固有のコミュニケーションパターンを理解することで、堅牢なメールセキュリティプラットフォームは、脅威となる前に異常なメッセージを特定・遮断することができます。

適切なテクノロジーを導入することで、既知の攻撃はもちろん、まだ観測されていない新種の攻撃に対しても、従業員が保護されているという確信を持つことができます。





▶▶ Abnormal AIについて

Abnormal AI は、人間の行動に基づくAIネイティブ型セキュリティプラットフォームを提供するリーディングカンパニーです。機械学習を活用し、高度化する受信型サイバー攻撃を防止するとともに、メールおよび連携アプリケーションにおけるアカウント侵害を検知します。同社の異常検知エンジンは、アイデンティティ情報およびコンテキスト情報を活用して人間の行動を理解し、すべてのクラウドメールイベントのリスクを分析します。これにより、人間の脆弱性を狙った高度なソーシャルエンジニアリング攻撃を検知・阻止します。

Abnormalは、Microsoft 365 または Google Workspace とAPI連携することで、数分で導入可能です。導入後すぐにプラットフォームの価値を実感可能です。

さらに、Slack、Workday、ServiceNow、Zoom をはじめとする複数のクラウドアプリケーションに対する追加保護も利用可能です。

現在、Abnormalは3,200社以上の組織に導入されており、Fortune 500 企業の20%超が利用しています。同社はAI時代におけるサイバーセキュリティの在り方を再定義し続けています。

詳しくはこちら

デモの依頼

