

TRENDLIFE RESEARCH SECURITY REPORT

個人セキュリティ脅威動向レポート

2026 上半期



TrendLife Research

本野 賢一郎、サイバーセキュリティ・イノベーション研究所

目次

はじめに.....	3
第 1 章：日本国内の脅威動向	4
特殊詐欺	5
ニセ警察詐欺	15
フィッシング詐欺	17
流行や大規模なイベントに便乗したネット詐欺	22
マルウェア	27
AI のリスクと脅威	32
詐欺被害を防ぐためには	40
2026 年下半期に向けて	41
第 2 章：新たな「信頼」の問題 The New Trust Problem	42
序文：信頼はどのように悪用されるのか.....	43
2026 年上半期の詐欺動向：3 つのパターン	44
転換点：力学は同じまま、情報源が変わった.....	45
AI リスクの検証結果	46
次のフロンティア：AI の利用リスクと AI への脅威	48
まとめ：拡大する接点と、追いつかない保護.....	49
調査体制と出典	49
参考文献	50

はじめに

本レポートは、2026 年上半期（1～6 月）におけるサイバーセキュリティの脅威動向について、個人利用者に被害をもたらす「詐欺」と「AI のリスク」の 2 つの観点を中心にまとめたレポートです。第 1 章では日本国内の脅威動向を、第 2 章ではトレンドマイクロのグローバルチームによる調査結果を紹介します。

ここ数年、個人利用者を狙うネット詐欺は、悪質化、巧妙化とともに被害規模の拡大が続いてきました。その傾向は 2026 年にさらに勢いを増し、特殊詐欺の被害は上半期として過去最悪を更新し、被害者もすべての世代に広がっています。また、生成 AI が日常の道具として定着する一方で、その AI をめぐるリスクは「攻撃者が AI を使う」段階から「利用者が頼る AI そのものが詐欺の入口になる」段階へと変わり始めました。本レポートでは、こうした脅威動向をトレンドマイクロの観測データと公開情報をもとに振り返り、個人利用者が今後の備えを講じるための洞察を提示します。

注記事項：

※1：データを含む本レポートの記述は編集時点での最新リサーチに基づくものですが、その後に新たな事実が判明することもあります。

※2：本レポートに掲載される数値は、特に明記されていない場合、トレンドマイクロの観測データや警察等の公的機関の公開情報が出典となります。

※3：本レポートで掲載した画像について、直接の危険や権利侵害につながりかねないと判断される部分には修正を施しています。

第 1 章

日本国内の脅威動向

特殊詐欺 | フィッシング詐欺 | 流行・イベントへの便乗詐欺 | マルウェア | AIのリスクと脅威

特殊詐欺

警察庁統計にみる被害状況

被害の全体像を見るうえで、本年は一つの節目にあたります。警察庁は 2026 年（令和 8 年）から統計区分を再編し、これまで別枠で集計されていた SNS 型投資詐欺と SNS 型ロマンス詐欺を特殊詐欺に含め、急増する「ニセ警察詐欺」（警察官などをかたり捜査名目で金銭等をだまし取る手口）を独立した手口として集計するようになりました¹。これにより詐欺被害の全体像がひとつの枠組みで把握しやすくなった一方、前年との単純な比較には注意が必要です。

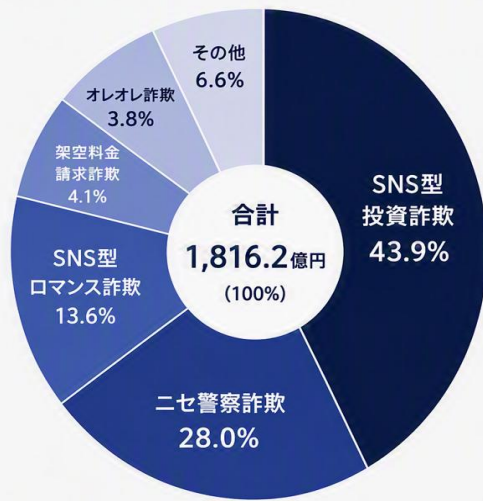
警察庁が 2026 年 7 月 31 日に公表した上半期（1～6 月）の暫定値によると、新しい定義による**特殊詐欺の認知件数は 22,031 件（前年同期比 18.3%増）、被害額は 1,816 億 2,000 万円（同 51.7%増）**で、いずれも上半期として**過去最悪を更新しました**²。**1 日あたりの被害額は 10 億円**、既遂 1 件あたりの被害額は 840 万円（同 28.1%増）に達し、被害は件数を上回るペースで高額の化しています。



図：令和 8 年上半期における特殊詐欺の認知・検挙状況等について [暫定値]（出典：警察庁）

手口別被害額の割合

単位：億円



※端数は四捨五入のため、合計と内訳が一致しない場合があります。

手口	被害額 (前年同期比)	既遂1件当たりの被害額
SNS型投資詐欺	797.9億円 (+444.9億円)	1,362.5万円
二セ警察詐欺	507.9億円 (+108.9億円)	1,163.9万円
SNS型ロマンス詐欺	246.6億円 (+6.3億円)	984.1万円
架空料金請求詐欺	74.6億円 (+6.1億円)	250.2万円
オレオレ詐欺	68.9億円 (+5.5億円)	380.7万円
その他	120.3億円 (+47.0億円)	293.2万円
合計	1,816.2億円 (+618.8億円)	840.0万円

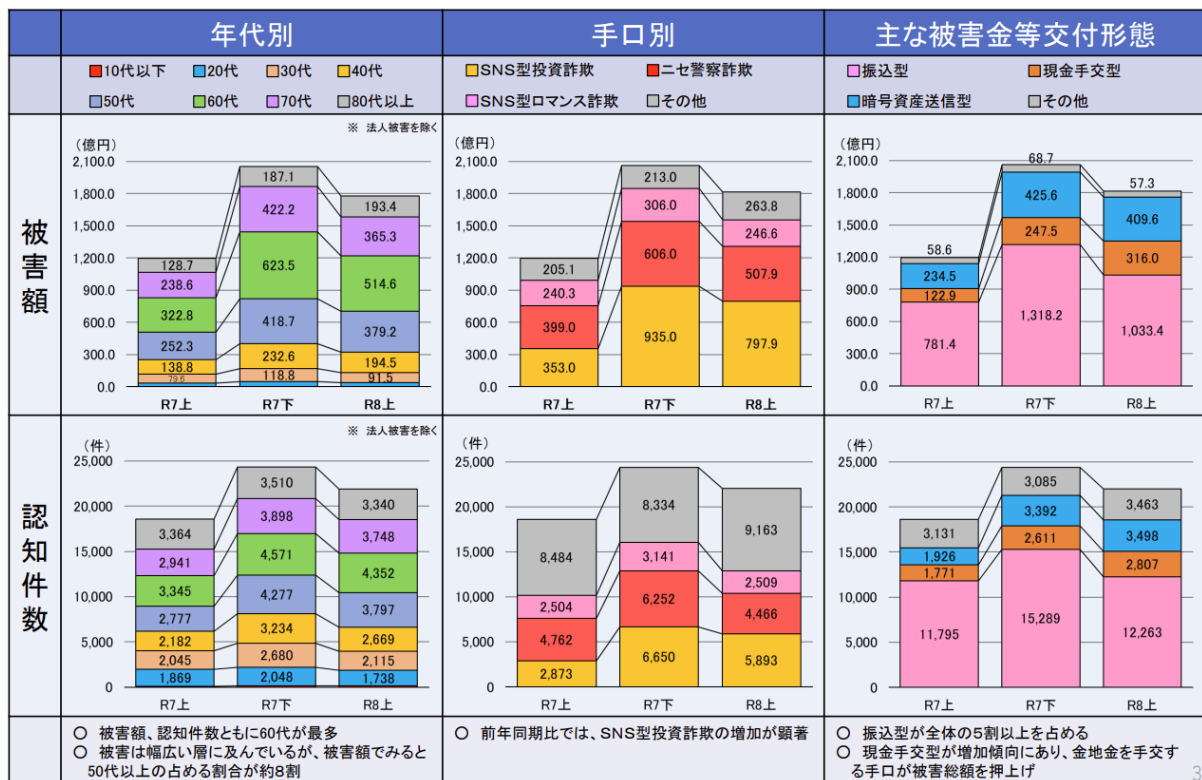
図：令和8年上半期における特殊詐欺の認知・検挙状況等について [暫定値] (出典：警察庁)

内訳では、SNS型投資詐欺が認知件数 5,893 件（前年同期比約 2.1 倍）、被害額 797 億 9,000 万円（同約 2.3 倍）といずれも最多で、特殊詐欺の被害額全体の約 44%を占めました²。次いで二セ警察詐欺が 507 億 9,000 万円、SNS型ロマンス詐欺が 246 億 6,000 万円と続きます²。二セ警察詐欺は認知件数（4,466 件）が前年同期比 6.2%減となる一方で被害額は 27.3%増えており、既遂 1 件あたり約 1,164 万円と、1 件ごとの被害の高額化が進んでいます²。

被害者の年齢層を見ると、認知件数、被害額ともに 60 代が最多で、被害額では 50 代以上が約 8 割を占めます（法人被害を除く）²。ただし、これは高齢者だけの問題ではありません。被害額が最も大きい SNS型投資詐欺では、認知件数は 50 代（1,605 件）、60 代（1,518 件）の順に多いものの、20 代から 50 代の現役世代が認知件数の約 6 割を占め、40 代は前年同期比 88.9%増と急増しています²。当初の接触ツールは Instagram（1,403 件）、X（旧 Twitter）（677 件）、LINE（631 件）が上位で、SNS などのバナー等広告を入口とするものが約 4 割と最多でした²。

二セ警察詐欺の認知件数も、最多が 70 代（869 件）である一方、30 代（664 件）、20 代（550 件）と幅広い年代に分布しており、警察庁も「高齢者だけでなく、20 代、30 代を含む幅広い年代の被害が増加」していると指摘しています²。「詐欺の被害者は高齢者」という従来のイメージは、もはや実態に合いません。

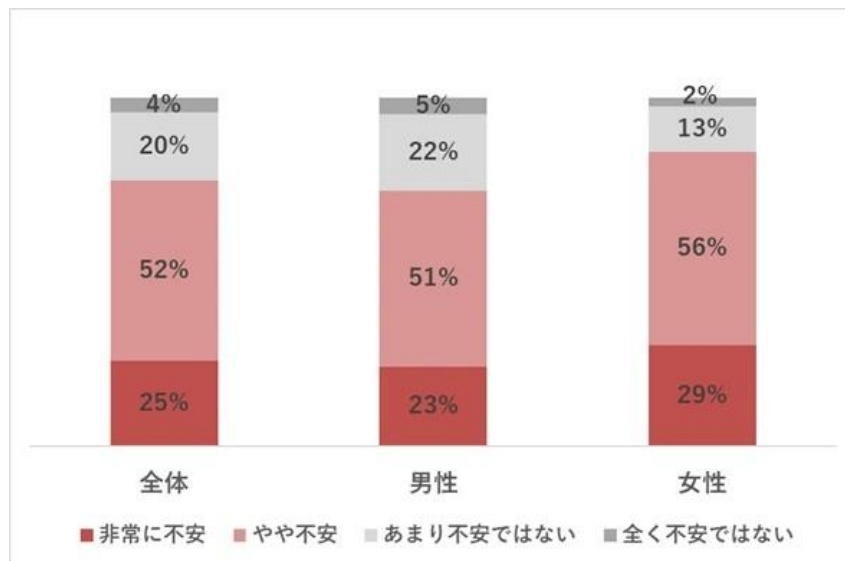
被害金の交付形態では「振込型」が認知件数、被害額とも全体の5割超を占め、その内訳ではインターネットバンキング経由が認知件数の54.9%、被害額の72.2%と、ATM 経由を大きく上回っています²。ネット上で完結する送金手段が、被害の主要な経路になっています。また、暗号資産送信型が認知件数3,498件（前年同期比81.6%増）、被害額409億6,000万円（同74.6%増）と急増しているほか、現金手交型も増加傾向にあり、中でも金地金を購入させてだまし取る手口の被害額は上半期だけで81億3,000万円と、2025年の1年分（60億2,000万円）をすでに上回りました²。



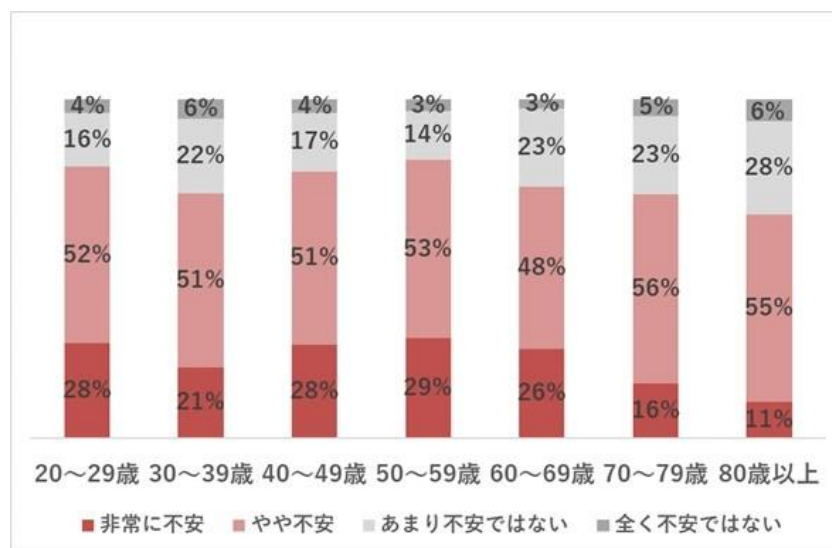
図：令和8年上半期における特殊詐欺の認知・検挙状況等について [暫定値]（出典：警察庁）

実態調査「電話・ネット詐欺不安度調査」

トレンドマイクロが実施した「電話・ネット詐欺不安度調査」³²では、**自分自身が「電話やネットを利用した詐欺被害」に遭う可能性**について、どの程度不安を感じているかを聞いたところ、**77%の回答者が不安に感じている**ことが明らかになりました。男女別では女性（85%）が男性（74%）を上回り、女性の方がより不安である傾向が見られます。一方、年代別に見ると50代の82%をピークに、年齢が上がるにつれて不安度は低下します。特に80代以上では66%と、全世代で最も低い数値となりました。**高齢者ほど自身が詐欺被害に遭う可能性を過小評価している**という危険な実態が浮き彫りとなりました。

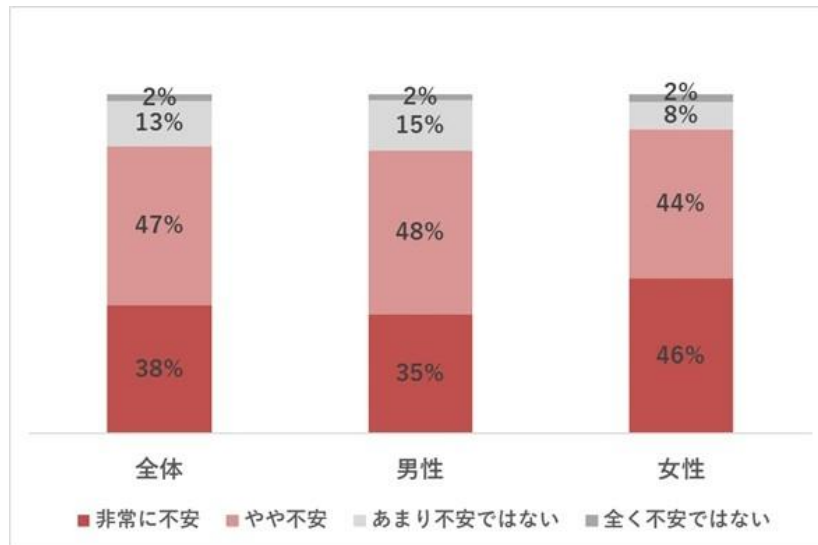


図：自身が「電話やネットを利用した詐欺被害」に遭う可能性について、どの程度不安を感じるか
全体・男女別（n=963）

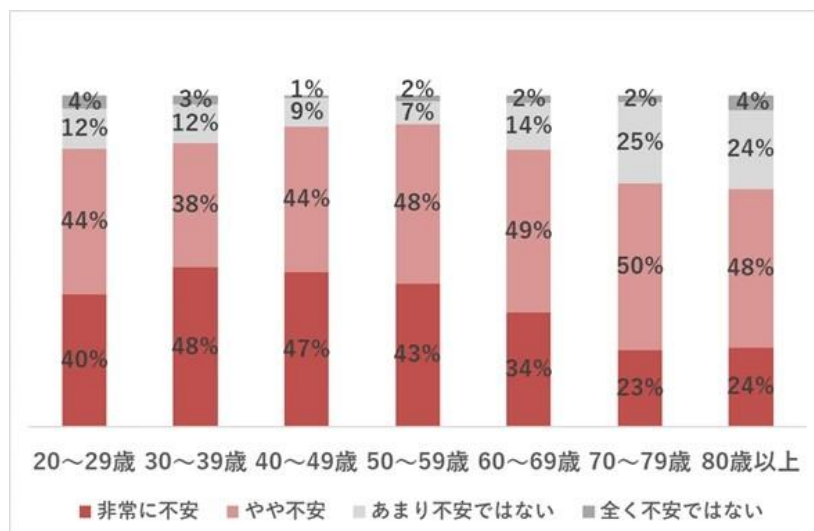


図：自身が「電話やネットを利用した詐欺被害」に遭う可能性について、どの程度不安を感じるか
世代別（n=963）

家族の「電話やネットを利用した詐欺被害」に遭う可能性について、どの程度不安を感じているかを聞いたところ、自身の詐欺被害（77%）を上回る **85%が「不安」と回答**しました。性別を問わず高い数値を示しており、特に40代、50代では90%を超えています。この背景には、自分の親世代（高齢者）と子ども世代（若年層）の双方を支える「サンドイッチ世代」特有の事情があり、この複数の精神的な負担を抱えていることが、高い不安度につながっていると推察できます。家族における日常的なコミュニケーションが、詐欺のリスクや不審な行動に気づきやすくなるためのセーフティネットになりますが、離れて暮らす親と子どもが頻繁にコミュニケーションするのは難しく、詐欺の兆候を見逃すリスクが高まっています。



図：家族が「電話やネットを利用した詐欺被害」に遭う可能性について、どの程度不安を感じるか
全体・男女別（n=878）*全体から家族ありのみを抽出



図：家族が「電話やネットを利用した詐欺被害」に遭う可能性について、どの程度不安を感じるか
世代別（n=878）*全体から家族ありのみを抽出

SNS 型投資・ロマンス詐欺

金額面で国内の詐欺被害の中心となっているのが、SNS 型投資・ロマンス詐欺です。入口となるのは Instagram や Facebook などの SNS 広告・投稿・DM やマッチングアプリで、被害に至る誘導には 9 割超の事例で LINE が使われています³。トレンドマイクロの観測では、2026 年は広告から LINE 登録へ直接誘導するパターンが増えたほか、6 月には初期誘導の LINE アカウントから Zoom や Signal といった別アプリへ誘導する事例も確認されました。LINE 上での警告や通報が増えたことを避ける動きの可能性があります。

トレンドマイクロの潜入調査では、**SNS 型投資詐欺が明確な 4 つの段階を踏んで進行する**構造が確認されています。

第 1 段階「誘導」では、SNS の投稿や偽広告、金融機関の偽公式アカウントなどから LINE へ誘導します（所要時間は数分から 10 分程度）。

第 2 段階「信頼構築」では、投資情報交換グループに参加させ、サクラの利益報告や特典で 2～3 週間かけて信頼を醸成します。

第 3 段階「架空取引」では、架空の取引サイトや偽アプリで口座を開かせます。

第 4 段階「金銭の要求」では、利益が出ているように見せかけたうえで、「出金には追加の入金が必要」などと口実を重ねて複数回の送金を要求します。

これらの段階は、初期誘導、アシスタント、著名人を装う投資家、口座開設担当、カスタマーサポートといった役割ごとに LINE アカウントを使い分ける分業体制で運営されています。

資金の回収手段にも変化がみられます。トレンドマイクロの観測では、従来の銀行振込に加えて暗号資産（仮想通貨）での支払い指示が増えており、金融機関による送金モニタリング（不審な取引の監視）の回避が狙いとみられます。さらに銀行から取引目的を尋ねられても「投資」という言葉を使わないよう被害者にあらかじめ口止めする、といった水際対策を無力化する細工も確認されています。「投資とは言わない」と口止めされている状況そのものが、詐欺を強く疑うべきサインです。

このように、SNS 型投資・ロマンス詐欺は、SNS、LINE、偽アプリ、送金という**複数のチャンネルをまたいで数週間かけて進行する**ため、単一のチャンネルだけで詐欺対策を強化しても、詐欺の全体像を捉えられないという課題があります。



図：著名人をかたる投資詐欺の例（SNS 広告から誘導された偽サイト）

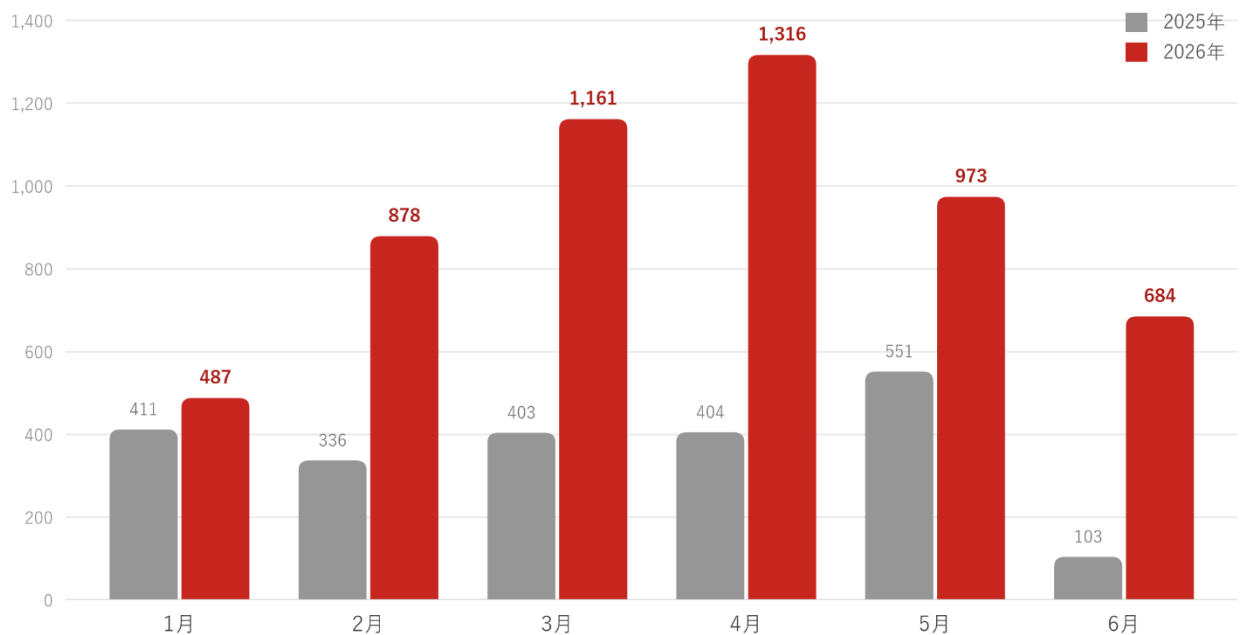
The image shows two examples of SNS advertisements. The left screenshot is a LINE chat bubble with text in Japanese. It mentions a 60-year-old man retiring and investing, with a goal of making 15 million by 2026. The right screenshot is a website banner for '2026年に、最速で1000万を作りたいなら' (If you want to make 10 million by 2026, the fastest way). It lists various stocks and their prices, and includes a call to action to add a LINE account.

図：SNS 広告から直接 LINE に誘導する事例

テクニカルサポート詐欺

テクニカルサポート詐欺（サポート詐欺）とは、パソコンの画面に偽のセキュリティ警告を表示し、記載されたサポート窓口の電話番号へ連絡させて、遠隔操作や金銭の支払いへ誘導する手口です。昨年まで、サポート詐欺は主に Web サイト閲覧中の不正広告から誘導されていましたが、2026 年上半期は**メールを起点とする誘導経路へと拡大**しました。

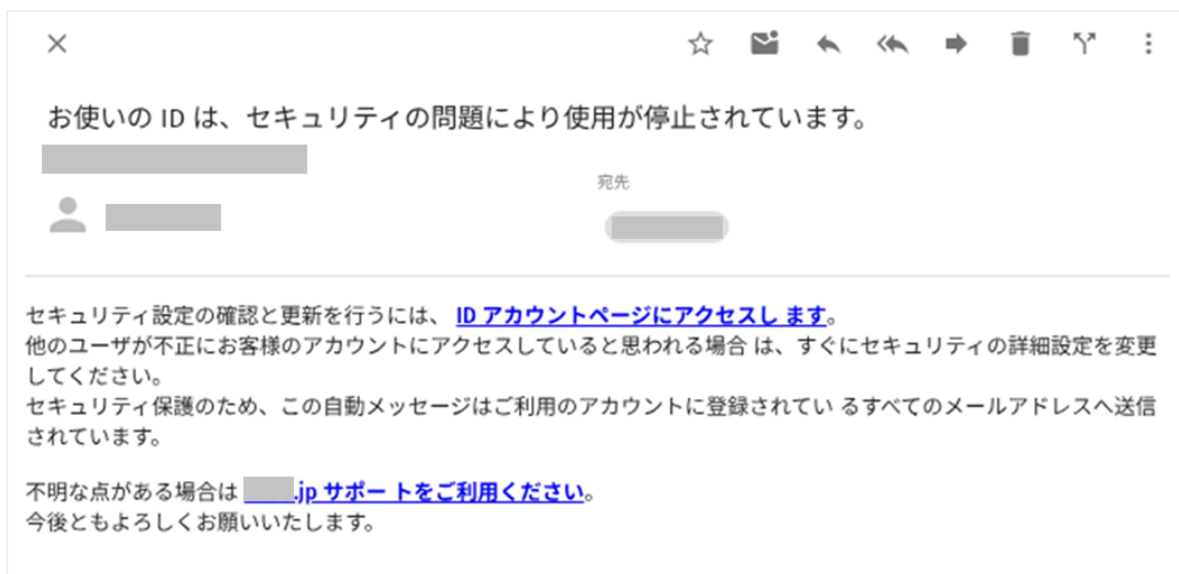
2026 年上半期、トレンドマイクロのサポートセンターに入ったサポート詐欺の報告件数は 5,499 件で、**前年同期比の約 2.5 倍（2,208 件→5,499 件）**に増加し、全ての月で前年実績を上回りました。



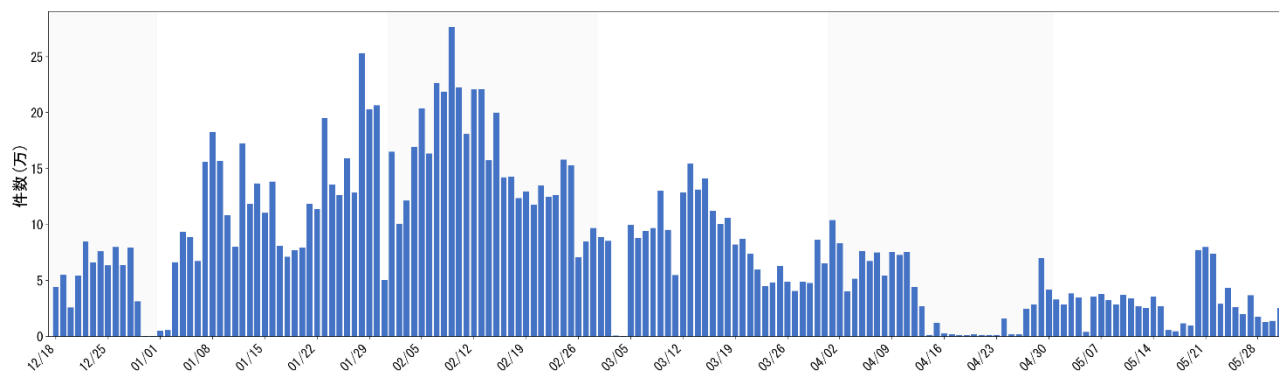
図：トレンドマイクロが受けたサポート詐欺の報告件数推移（月別）

トレンドマイクロの調査では、2025 年 12 月中旬から 2026 年 5 月までの 165 日間にわたり、**メール経由で偽のセキュリティ警告サイトへ誘導する大規模なサポート詐欺キャンペーンを観測**しました⁴。観測されたメールは約 **1,338 万通**、1 日平均で約 8.1 万通に上り、宛先の約 **94%**が「.jp」ドメインでした。配信は日本時間の 9～21 時に集中しており、**日本の利用者を狙って**生活時間帯に合わせて運用されていることがうかがえます。

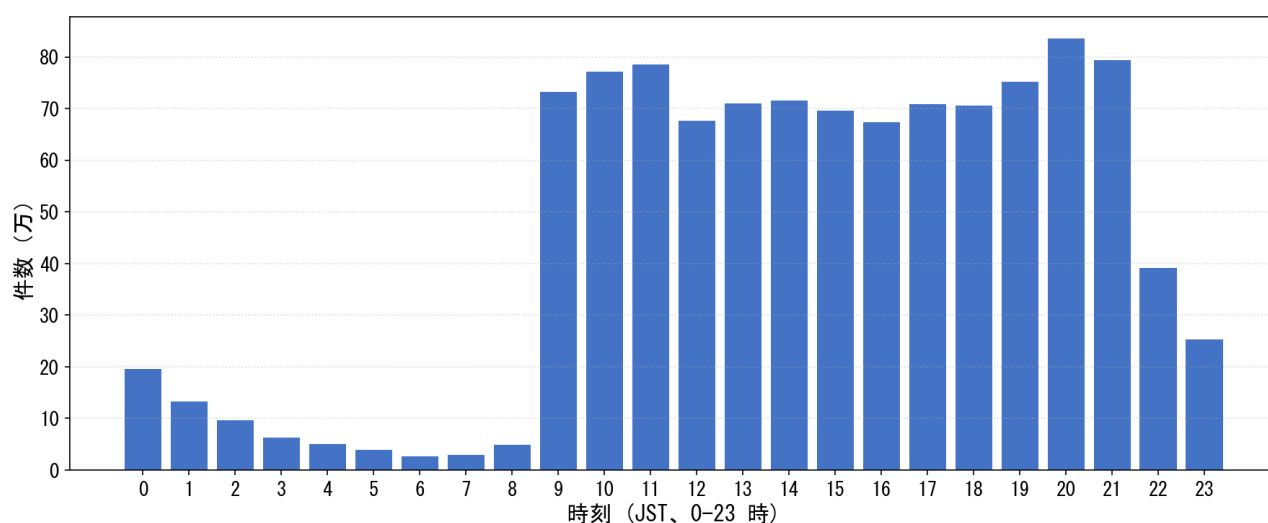
誘導先として使われた使い捨ての偽警告サイトは 33,000 件を超え、約 9 割が 1 日限りで乗り捨てられていました。配信量のピークは 2026 年 2 月の約 445 万通で、その後減少したものの、5 月時点でも 1 日平均約 3 万通が続いていました⁴。偽警告サイトの構築には Microsoft Azure の静的 Web ホスティング機能が悪用されており、正規クラウドサービス上にあることが検知を難しくしています。



図：2026 年 3 月に観測したサポート詐欺への誘導を狙うメールの一例



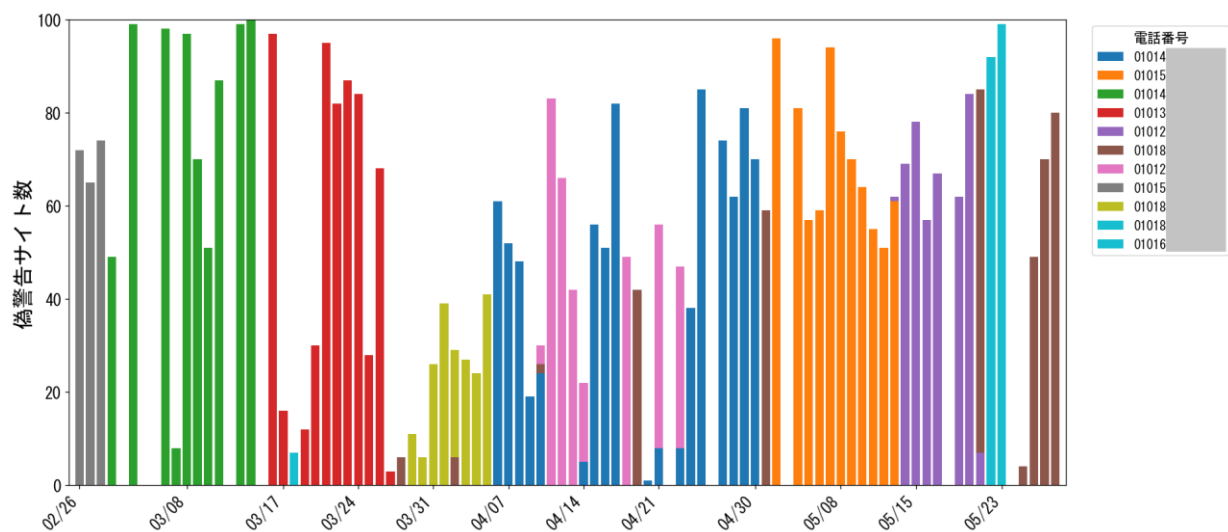
図：詐欺メール件数の日次推移



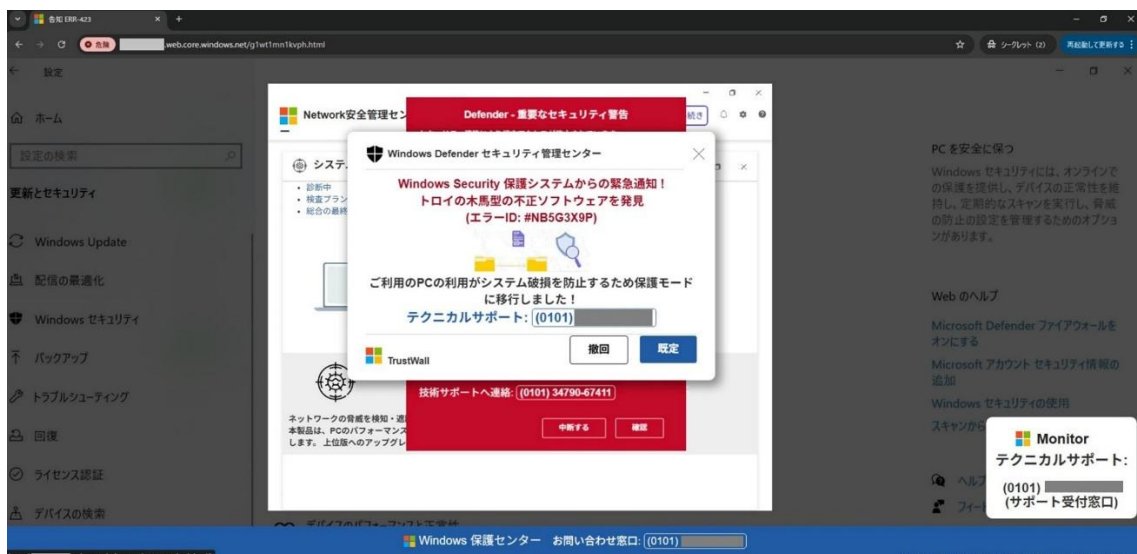
図：詐欺メール件数 時間帯別（日本時間）*メールゲートウェイ製品の統計値

2026 年 4 月中旬からは大手 EC サイトや国税庁などをかたる文面が、5 月からは「人事評価」「給与改定」といった社内連絡を装って企業の従業員を狙う文面が確認されており⁴、**個人利用者だけでなく法人の従業員も標的に含まれる**ようになりました。

一方で、対策の手がかりもあります。約 3 か月間（2026 年 2 月 26 日～5 月 31 日）に観測された主な偽警告サイト 4,721 件に記載されていた電話番号はわずか 11 個で、1 つの番号が数百のサイトで使い回されていました⁴。サポート詐欺は、入口が不正広告でもメールでも、偽の警告から電話をかけさせるという構造は同じであるため、画面に表示された電話番号には決して電話をかけないことが重要です。また、詐欺電話対策も有効な手段となります。



図：詐欺電話番号別の偽警告サイト数



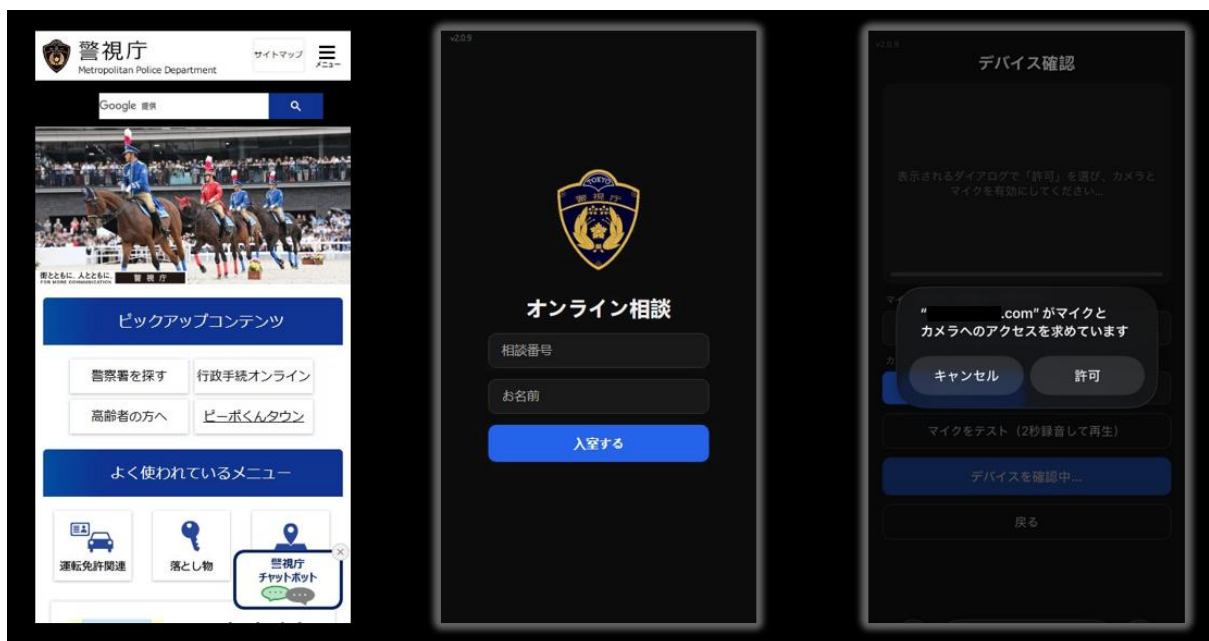
図：誘導先の偽警告サイトの表示例

ニセ警察詐欺

警察官などを装う「ニセ警察詐欺」は、上半期の被害額が 507 億 9,000 万円と、SNS 型投資詐欺に次ぐ規模になりました²。1 件あたりの被害が大きいことが特徴で、上半期の暫定値では**既遂 1 件あたり約 1,164 万円（前年同期比 37.7%増）**に達します²。被害者は 20 代・30 代にも広く分布しており、**若年層も標的**となっています²。

当初接触ツールとして、**電話が 9 割以上を占める**状況が継続しています。主な手口として、電話から LINE 等に誘導された後に、攻撃者がディープフェイク（AI で生成した偽の映像・音声）で別人の顔を合成し、ビデオ通話に警察官や検事の役として登場する事例を確認しています。

また、トレンドマイクロの調査では、新たな手口として、**警視庁を装う偽サイトに「ビデオ通話機能」が追加されたことを確認**しました。従来の警察の偽サイトは偽の逮捕状を表示し個人情報や銀行口座情報を入力させるものでしたが、ブラウザ経由のビデオ通話により、LINE 等のアプリへの誘導を介さずに「顔の見える取り調べ」を演出できるようになっています。



図：ビデオ通話機能が追加された警視庁を装う偽サイト（出典：トレンドマイクロ）

ニセ警察詐欺が巧妙化しているため、表示された電話番号や映像が本物かどうかを、受け手の人間が見分けることは困難な状況です。最も確実な対策は、一度電話を切り、公式サイトや信頼できる資料で確認した正規の電話番号に自分からかけ直すことです。また、詐欺対策アプリを活用することも有効です。

「警察庁推奨」特殊詐欺対策アプリ

対策側の動きとして、警察庁は 2025 年 12 月、要件を満たす詐欺電話対策機能を備えた民間スマホアプリを「警察庁推奨アプリ」として認定する「特殊詐欺対策アプリに係る警察庁推奨制度」を開始しました⁵。トレンドマイクロはこの制度のもと、警察庁推奨アプリとして無料の特殊詐欺対策アプリ「トレンドマイクロ 詐欺バスター Lite」を 2026 年 3 月にリリースしています⁶。

このアプリは、警察庁が保有する詐欺電話データベースとトレンドマイクロ独自のデータベースを組み合わせ、詐欺電話や国際電話の発着信時に警告を表示します（スマートブロック設定で自動切電も可能。国際電話関連機能は Android のみ）。あわせて、警察庁提供の防犯情報とトレンドマイクロ制作の詐欺関連ニュースを通知します⁶。警察庁の公表によれば、推奨アプリ全体（「トレンドマイクロ 詐欺バスター Lite」と「詐欺対策 by NTT タウンページ」）の累計ダウンロード数は 2026 年 6 月末時点で約 147.6 万件に達しています²。これらの取り組みによって、電話による詐欺被害を低減できると期待されています。

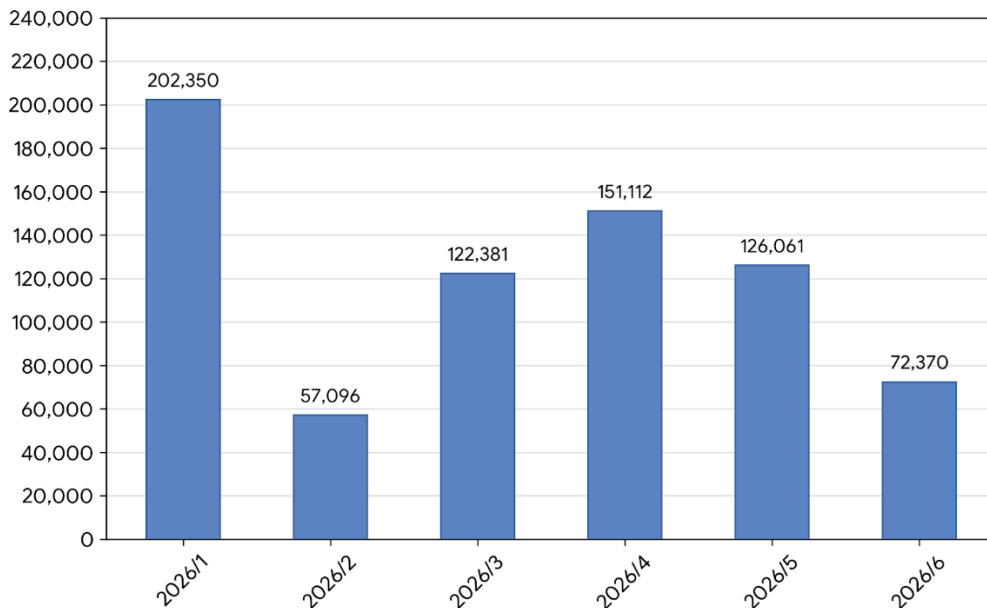
The advertisement features a group of police officers in riot gear holding shields that list various types of fraud: 運付金詐欺 (Delivery Money Fraud), 架空料金請求詐欺 (Fake Bill Request Fraud), ニセ警察詐欺 (Fake Police Fraud), オレオレ詐欺 (Con Artist Fraud), and 副業詐欺 (Side Business Fraud). The text 'みんなで' (Everyone together) is written above them. The app is promoted as '無料' (Free) and '警察庁推奨' (Police Agency Recommended). The main title is '特殊詐欺対策アプリ' (Special Fraud Countermeasure App). Below this, the app's main functions are listed: 01 国際電話をブロック (Block International Calls), 02 詐欺電話をブロック (Block Fraud Calls), and 03 最新手口を把握 (Understand the latest methods). A QR code and download links for Google Play and the App Store are provided, along with the text '今すぐダウンロード!' (Download now!).

図：警察庁推奨「特殊詐欺対策アプリ」（出典：警察庁）

フィッシング詐欺

フィッシング対策協議会への報告件数によると、2026 年 1 月は 20 万件を超える高い水準でしたが、2 月には 57,096 件と約 7 割減まで一気に落ち込みました。攻撃者が悪用していた海外のレジデンシャルプロキシ（一般家庭の IP アドレスを経由させて送信元を偽装する中継サービス）への対策が進んだことに加え、旧正月で攻撃活動自体が減少した影響もあったとみられます。しかし、3 月には 122,381 件と前月比 2 倍超に急増し、4 月には 151,112 件まで増加します。この時期は税金や公金の納付を装う手口だけで報告の約 5 割を占め、納付通知が届く季節に合わせて攻撃が集中しました。5 月も 126,061 件と高止まりし、6 月は 72,370 件と前月比 42.6%減まで下がりました。6 月の件数減少と入れ替わるように、国内の複数の ISP（インターネット接続事業者）の利用者のメール認証情報が悪用され、乗っ取られた正規のメールアカウントからフィッシングメールが送信される手口が確認されました⁷。正規アカウントからの送信は送信ドメイン認証をすり抜けるため、従来の迷惑メール対策では防ぎにくくなっています。

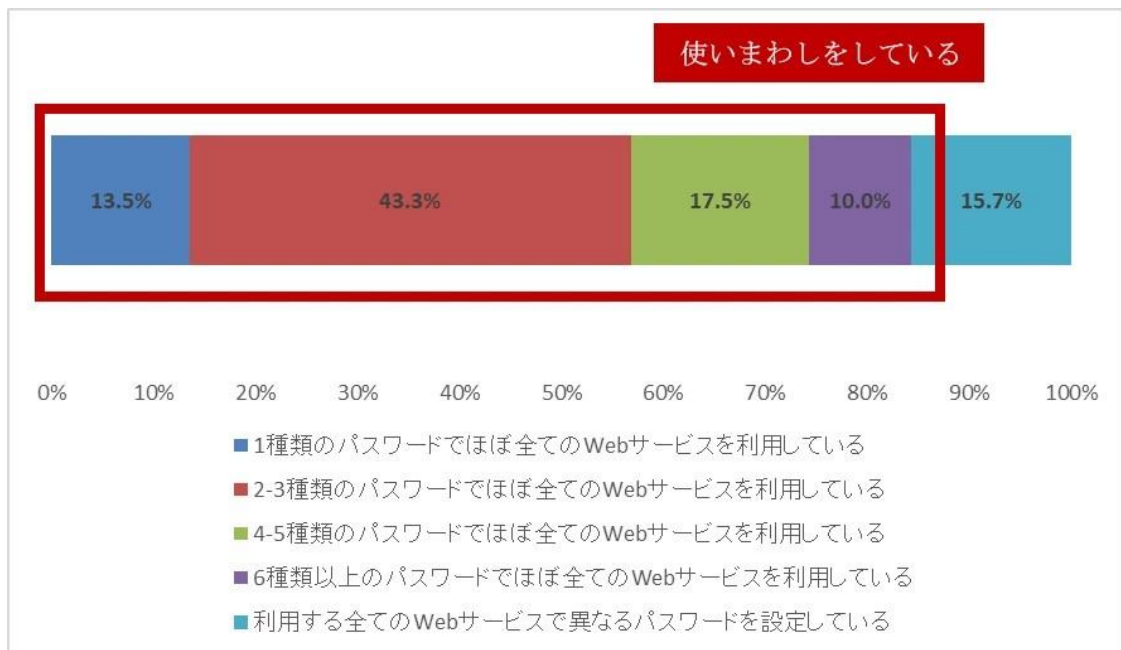
対策によって一時的に 7 割減っても、攻撃者は別の経路に乗り換えて活動を再開させるなど、「イタチごっこ」の構図が上半期を通じて続いています。



図：フィッシング報告件数（出典：フィッシング対策協議会）

パスワード・パスキーの利用実態調査 2026

認証をめぐる利用者側の実態も、対策を考えるうえで欠かせません。トレンドマイクロが 2026 年 2 月に実施した「パスワード・パスキーの利用実態調査 2026」では、**84.3%が複数の Web サービスでパスワードを使いまわしている**と回答し、2023 年調査（83.8%）から横ばいでした⁸。使いまわしの理由は「異なるパスワードを設定すると忘れてしまう」（74.3%）が最多です。



図：Web サービスの利用で、パスワードを使い分けているかの実態（単一回答：n=1,034）

一方、パスキー（パスワードに代わる、フィッシング耐性の高い認証方式）の認知は広がっており、「仕組みまで理解している」（41.9%）と「名前を聞いたことがある」（49.2%）を合わせて 91.1%に達しました。パスキー提供サービスの利用者では 87.7%がパスキー機能を使っています⁸。使いまわされたパスワードは、どこか 1 か所の流出が全アカウントの乗っ取りにつながるため、パスキーへの移行はこのリスクを断ち切る有効な手段です。しかし、パスキーに関連した攻撃もすでに現れています。

銀行・証券・EC を狙う手口：リアルタイムフィッシングとパスキー移行への便乗

フィッシングで利用者を誘導した先の偽サイトは、ID・パスワードを盗むだけのものではなくなっています。トレンドマイクロの観測では、ワンタイムパスワード（OTP）、秘密の質問・合言葉、電話番号認証といった多要素認証の要素まで、その場で詐取する **AiTM 型のリアルタイムフィッシングが継続**していることを確認しています。

2025 年に社会問題となった**証券口座の乗っ取りは、ピーク時から大幅に減少**しています。金融庁の公表によると、証券会社を通じた不正アクセス・不正取引の累計（各社報告の合算）は 2026 年 8 月 10 日時点で不正アクセス 19,142 件、不正取引 10,607 件、不正売買総額 8,200 億円に達しましたが、月次の不正売買額は 2026 年 3 月の 304 億円をピークに、4 月 146 億円、5 月 27 億円、6 月 11 億円と大幅に減少しました⁹。証券各社によるパスキーの導入・必須化や本人確認の強化など、複数の対策が進められており、被害減少に寄与した可能性があります。ただし不正アクセスの報告自体は続いており、収束したと断定するのは早計です。

このパスキー移行という前向きな動きに便乗する攻撃も現れました。2026 年 4 月には、証券会社をかり「セキュリティ強化に伴うパスキー設定のお願い」を口実にしたフィッシングの緊急情報が公開されています¹⁰。強調しておきたいのは、パスキーの仕組み自体が破られているわけではないという点です。この手口はパスキーの「設定手続き」を装って偽サイトへ誘導し、従来の ID・パスワードを入力させて詐取するものです。認証方式の変更を促す連絡ほど、公式アプリや正規サイトから確認することが大切です。

クレジットカードを狙う手口も一段と巧妙になっています。トレンドマイクロの観測では、大手カードブランドをかたる偽サイトがカード番号に加えてワンタイム認証コードや電話認証までリアルタイムに詐取し、スマートフォン決済（Apple Pay など）へのカード登録時の本人確認を突破する手口を確認しています。また、**Apple ID や Google アカウントの乗っ取りを起点に、クラウド側にメッセージを保管する同期機能を悪用して被害者宛ての SMS 認証コードを別端末で盗み見る、多要素認証を突破する手口も確認**しています。多要素認証やパスキーを設定していても、その土台となるプラットフォームのアカウントが乗っ取られれば防御は崩れます。土台のアカウントを強固に守り、同期設定を見直しておくことが、認証全体を守るうえで欠かせません。

なお、通信事業者をかたる偽サイトから iPhone に不審なアプリを導入させ、詐取したネットワーク暗証番号でキャリア決済を不正利用する手口が 1 月中旬から確認されています¹¹。**Apple のアプリベータ配信サービスや VPN 構成プロファイルといった正規の仕組みが悪用**されるため、心当たりのないアプリ導入や構成プロファイルの追加には応じないことが重要です。

QRコードを悪用した詐欺

クイッシング（Quishing、QRコードを使ったフィッシング）も、昨年に続き確認されています。不正な URL を文字列ではなく QRコードの画像に埋め込むことで、迷惑メールフィルタによる検出を回避する狙いがあります。

トレンドマイクロの観測では、2026 年上半期はこの手口が機器をまたぐ形に発展しています。パソコンでフィッシングメールのリンクを開くと QRコードが表示され、それをスマートフォンで読み取らせて決済アプリの支払い用リンクへ誘導する事例が増えました。パソコンとスマートフォンの 2 つの機器をまたがせることで、片方の機器だけでは働きにくいセキュリティ対策の隙間をすり抜けようとするものです。法人向けには、FAX で送られた QRコードからフィッシングサイトへ誘導する、メールを経由しない手口も引き続き確認しています。

2026 年上半期のフィッシングで目立ったのは、税金や社会保険料といった「公金の納付」を装い、**スマートフォン決済アプリの送金画面へ直接誘導して金銭を詐取する**手口です。認証情報を盗む従来型と異なり、利用者自身に「送金」を操作させて完結させる直接詐取型である点が特徴です。

確認された主な手口：

- 通信料金の引き落とし失敗を装う手口
- 国民健康保険の保険料未納を装って個人間送金用 URL へ誘導する手口
- 生命保険会社をかたる送金悪用の継続キャンペーン
- 住民税の納付を装う手口
- 日本年金機構をかたり「差押執行前の最終確認」などの件名を用いる手口

2026 年 4 月下旬にはフィッシング報告の約 5 割を税金・公金・保険料を装う手口が占め¹²、PayPay での送金を求める詐欺 SMS が急増しました。事業者による注意喚起の公表後、5 月下旬からこの種の誘導メールは減少しました。両者の因果関係は確定できないものの、注意喚起や対策が影響した可能性があります。

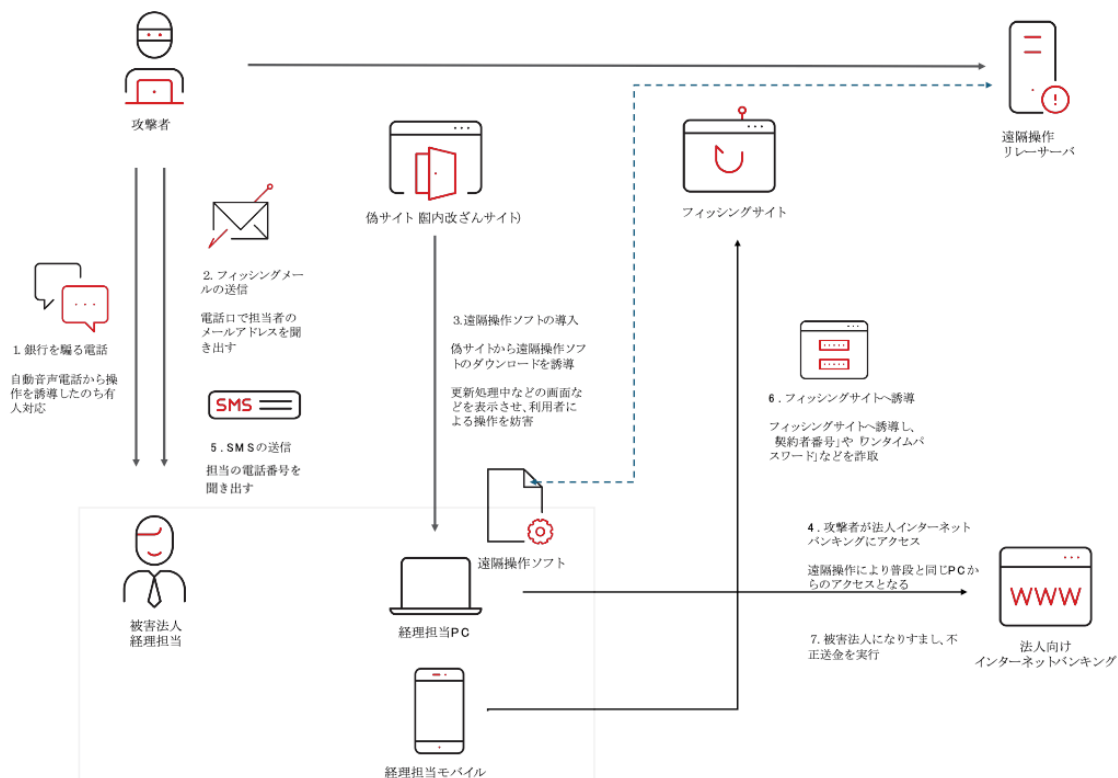
その他、**QRコードを悪用した「返金詐欺」**にも注意が必要です。ネットショッピングの後に「欠品のため返金します」といったメッセージが届き、案内された QRコードを読み取って操作すると、実際には攻撃者への「送金」が実行されます。受け取り側であるはずの自分が入力や承認の操作を求められた時点で、いったん立ち止まってください。正規の「返金手続き」で、返金を受ける側が相手への「送金」を求められることはありません。

インターネットバンキングを狙うボイスフィッシング

トレンドマイクロの調査では、ボイスフィッシングとテクニカルサポート詐欺を組み合わせた新たな手口を確認しています。現在の主な標的は法人口座ですが、同様の手口が個人利用者へ波及する可能性には引き続き警戒が必要です。

電話を起点に法人インターネットバンキング（法人 IB）の認証情報を詐取するボイスフィッシングは、2024 年 11 月頃から断続的に高額被害を生んでおり、2026 年 5 月中旬から再び活発化しました¹³。警察庁の公表によれば、2025 年（令和 7 年）の不正送金被害 4,747 件・約 103 億 9,700 万円のうち、法人 IB を狙うボイスフィッシングによる被害は 143 件・約 45 億円と、被害総額の約 4 割を占めます¹³。

手口は多段階で、金融機関の担当者を装った電話は自動音声から始まり、人間のオペレーターに切り替わって相手のメールアドレスを聞き出し、「これから当行よりメールをお送りします」と予告したうえで本物そっくりのメールを送信します。電話で信頼関係を作ってからメールを送るため、「心当たりのないメールは開かない」という従来の注意が通用しにくいのがこの手口の危険なところ。2026 年 5 月以降は、手口がさらに進化しました。電話と偽サイトに加え、改ざんされた国内サイトを經由して正規の遠隔操作ソフト（ScreenConnect 等）をインストールさせ、**被害者のパソコンそのものを遠隔操作する手口**が登場しています¹³。被害者本人の端末からの操作になるため、金融機関側の不正検知が働きにくくなります。6 月には、e-Tax（国税電子申告・納税システム）の事務局を装い、e-Tax ソフトに偽装した遠隔操作ソフトを導入させる亜種も確認され、**標的は都市銀行から地方銀行、信用金庫へと広がりました**。



流行や大規模なイベントに便乗したネット詐欺

人気のシールの偽ショッピングサイト

2026 年上半期、子どもたちの間で大流行した立体シールの品薄に便乗して、正規品の販売を装う偽ショッピングサイトが急増しました。トレンドマイクロの調査では 2026 年 2 月だけで 251 件、3 月末時点で累計 532 件、4 月末時点で累計 754 件を確認しました。5 月には偽シールを販売目的で所持していた容疑者が商標法違反で逮捕される事件も起きています。

主な手口として、**検索エンジン経由で偽ショッピングサイトへ誘導する「SEO ポイズニング」**が活発化しています。攻撃者は正規サイトを改ざんし、人気商品や話題のイベント関連グッズに関する検索結果の上位へ不正なページを表示させることで、利用者を偽サイトへ誘導します。

被害者は、極端に安価な商品や入手困難な限定商品に惹かれて購入手続きを進めることで、商品代金やクレジットカード情報、個人情報をだまし取られる可能性があります。また、商品未着や偽商品の送付に加え、「返金手続き」を装い QR コード決済で送金させる返金詐欺に発展する事例も報告されています。漏えいした情報は他の金融犯罪や不正利用に悪用されるリスクもあります。

対策としては、検索結果を過信せず、公式サイト URL 確認、事業者情報や連絡先の実在性確認、極端な値引きの有無、支払方法の妥当性、日本語表現の不自然さなどを購入前に確認することが重要です。また、不審なサイトへのアクセスを検知・遮断できるセキュリティ対策製品の活用も有効です。

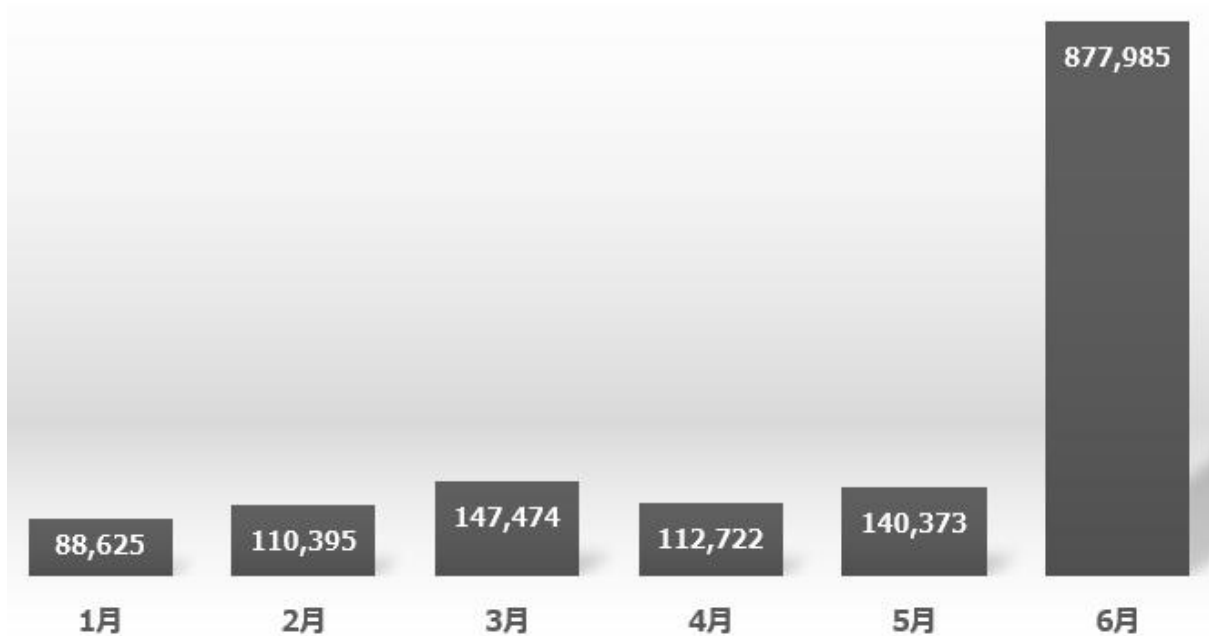


図：偽ショッピングサイトの例

FIFA ワールドカップ 2026 に便乗するネット詐欺

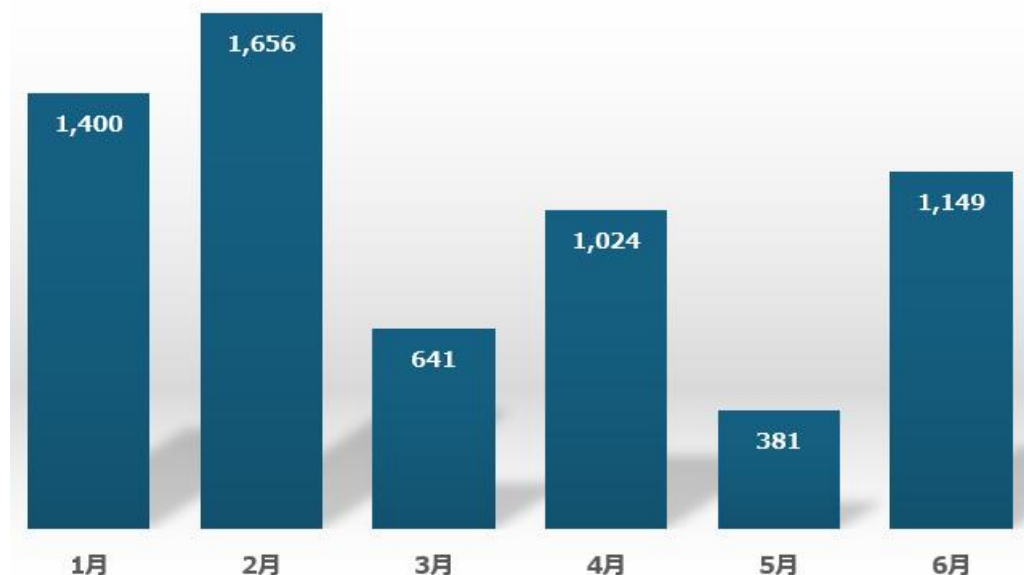
2026 年 6 月 11 日に開幕した FIFA ワールドカップ 2026（米国・カナダ・メキシコ共催）に便乗したネット詐欺が、世界中で確認されました。米 FBI のインターネット犯罪苦情センター（IC3）は 2026 年 5 月に大会便乗詐欺への注意喚起を公開しています¹⁴。

トレンドマイクロの調査では、2026 年 1～6 月に「fifa」「worldcup」を含む詐欺・フィッシング等の不正サイトおよび不審サイトを 35,538 件確認し、これらへの日本国内からのアクセスは約 148 万件に達しました¹⁵。主な手口は、偽ショッピングサイト、偽チケット販売サイト、偽ライブ配信サイトで、日本の利用者也、この詐欺の主な標的となっています。



図：日本国内から FIFA ワールドカップ関連の不正サイトおよび不審サイトへのアクセス数
（2026 年上半期）

偽ショッピングサイトは、検索エンジンの結果を汚染する「SEO ポイズニング」により利用者を誘導し、ユニフォームや応援グッズなどを販売しているように見せかけます。2026 年上半期だけで FIFA ワールドカップ関連の偽ショッピングサイトを 6,251 件確認しています。攻撃者は様々な手口で利用者から個人情報や購入代金をだまし取ろうとします。これらの商品を購入しても、商品が届かないことが多く、届いたとしても偽物や粗悪品の場合もあります。



図：FIFA ワールドカップ関連の偽ショッピングサイトの数（2026 年上半期）

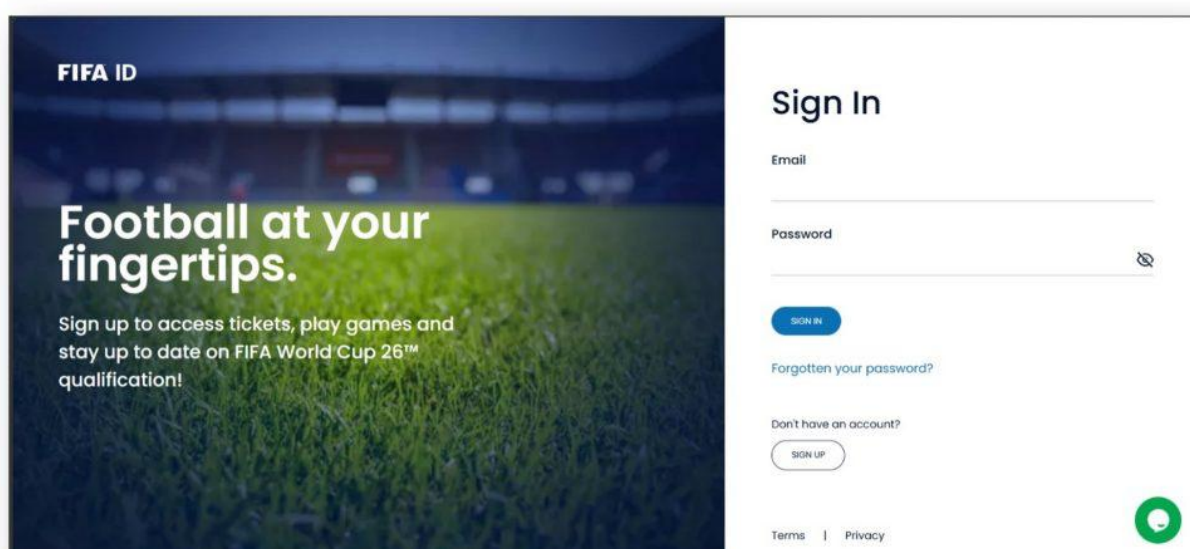


図：偽ショッピングサイトの例

さらに、**FIFA 公式ホスピタリティサイトをほぼ完全に複製した偽チケット販売サイトも確認**しています。この偽サイトは公式サーバから動画・画像を直接読み込み、自動翻訳機能まで備えていました。利用者がログイン情報を入力すると認証情報が詐取されるほか、決済時に入力したクレジットカード情報や SMS で送られるワンタイムパスワードをリアルタイムで盗み取り、不正決済を完了させる「**リアルタイムフィッシング**」が行われています。多要素認証を利用していても、利用者自身が偽サイトに認証コードを入力することで被害につながる点が特徴です。



図：FIFA 公式ホスピタリティサイトを模倣した偽サイトの例



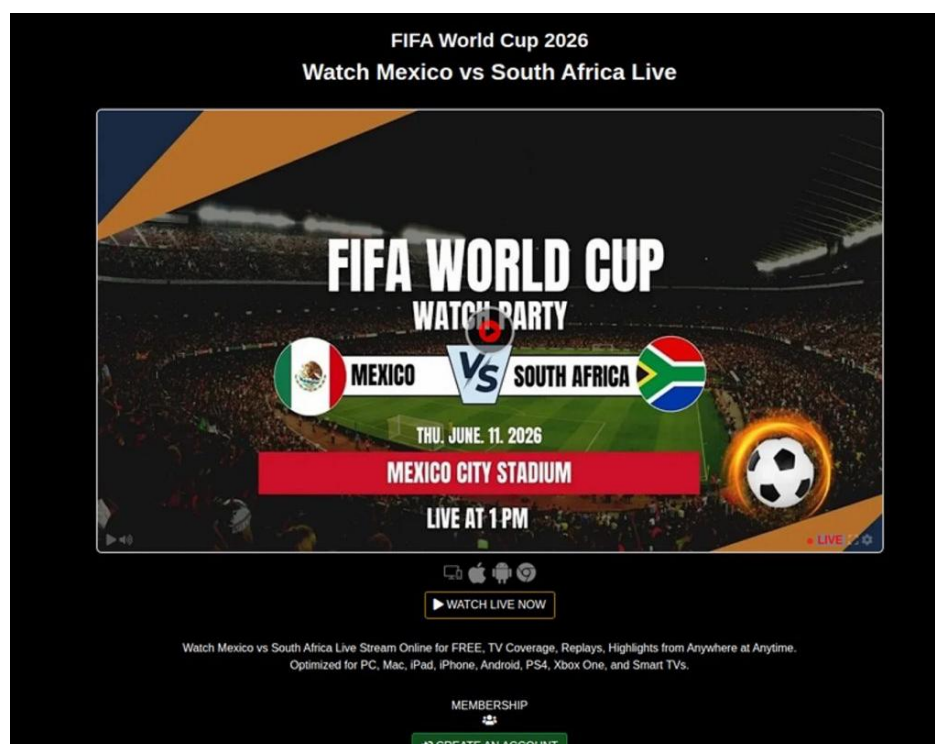
図：偽サイトのログイン画面

また、「FIFA ワールドカップ 2026 無料配信」などの検索キーワードを悪用し、**偽ライブ配信サイトへ誘導する事例**も確認しています。これらのサイトでは実際の試合映像は配信されず、広告詐欺や個人情報・クレジットカード情報の詐取、不正なサブスクリプション契約への誘導が目的となっています。

利用者は「無料」や「公式」といった表現を安易に信用せず、チケットやグッズは公式サイトへ直接アクセスして購入すること、URL の確認、サービスごとのパスワード管理、多要素認証時の利用先確認などを徹底することが重要です。



図：日本語の検索結果に表示された偽ページ（改ざんされたサイト）



図：偽ライブ配信サイトの例

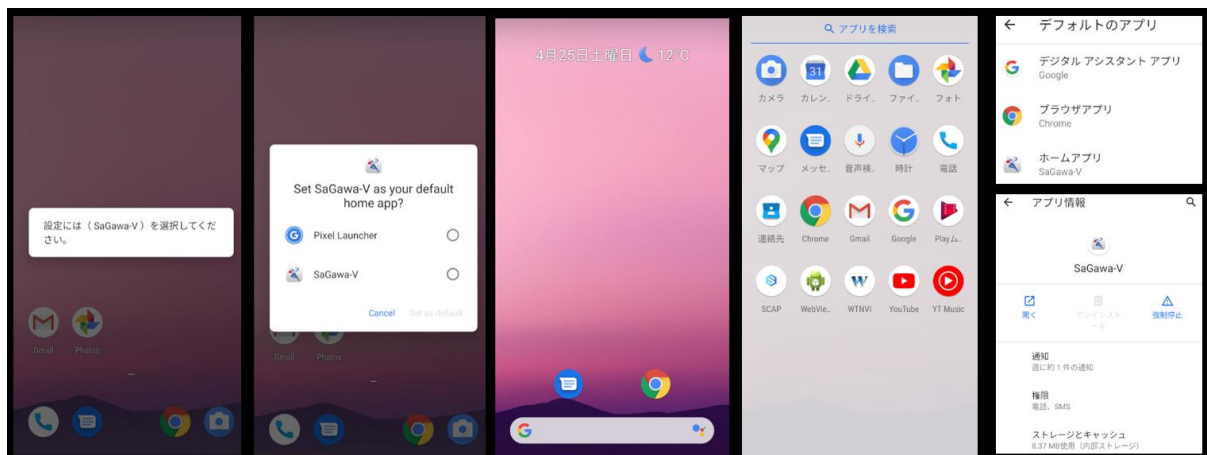
マルウェア

詐欺と並んで、個人利用者のスマートフォンやパソコンを狙うマルウェア（不正プログラム）にも注意が必要です。

不正 SMS を拡散するスマートフォン向け不正アプリ

不正 SMS（スミッシング）の背後では、昨年に続き、感染したスマートフォンが攻撃者の「踏み台」として新たな不正 SMS を送信する構造が続いています。不正 SMS の主な送信元は国際電話番号ですが、「+81」表示や「090」「080」「070」で始まる国内の携帯番号から届くこともあるのは、マルウェア感染したスマートフォンがあるためです。こうした不正 SMS は、攻撃先の端末 OS に合わせた出し分けがされており、同じリンクでも iPhone ではフィッシングサイトのリンクへ、Android ではアプリ（拡張子.apk）を装ったマルウェアの感染へ誘導されます。

XLOADER（別名 MoqHao） は、Android 向けの不正アプリです。宅配業者の不在通知などをかたる不正 SMS から偽サイトへ誘導し、Chrome アプリなどを装って感染させます。2026 年上半期は文面に変化がみられ、従来の不在通知型とアダルト系の文面を交互に使い分ける傾向がありました。さらに 2026 年 4 月下旬以降は、端末の「デフォルトのホームアプリ」に自身を設定させ、アプリ一覧から姿を隠してアンインストールを妨害する挙動へと変化しています。感染に気づいても削除しにくくなる、しつこさを増した動きです。



図：マルウェア（XLOADER）が画面表示を隠す事例

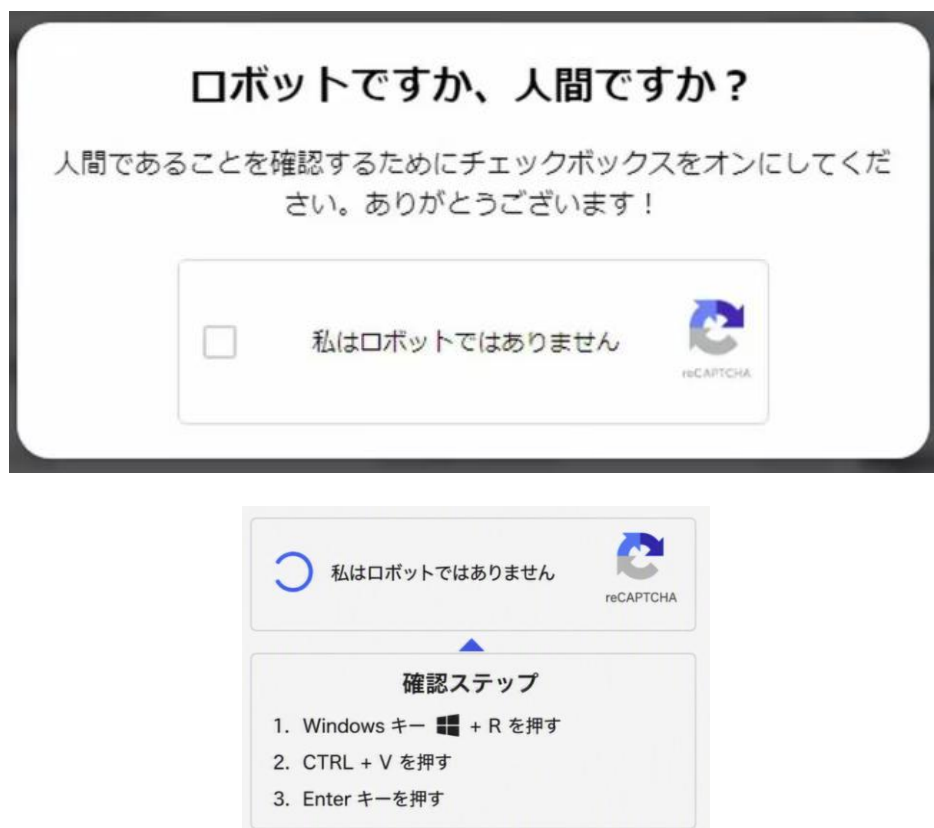
KeepSpy も継続して観測される不正アプリです。2026 年上半期は、ほぼ連日「国税庁」をかたる SMS（「重要な行政手続きが未完了です」など）を用いて拡散する動きが目立ちました。これは、同時期に金融機関や公的機関をかたるスミッシングが増えた全体的な傾向とも重なります。

利用者自身に操作させる「ClickFix」の手口

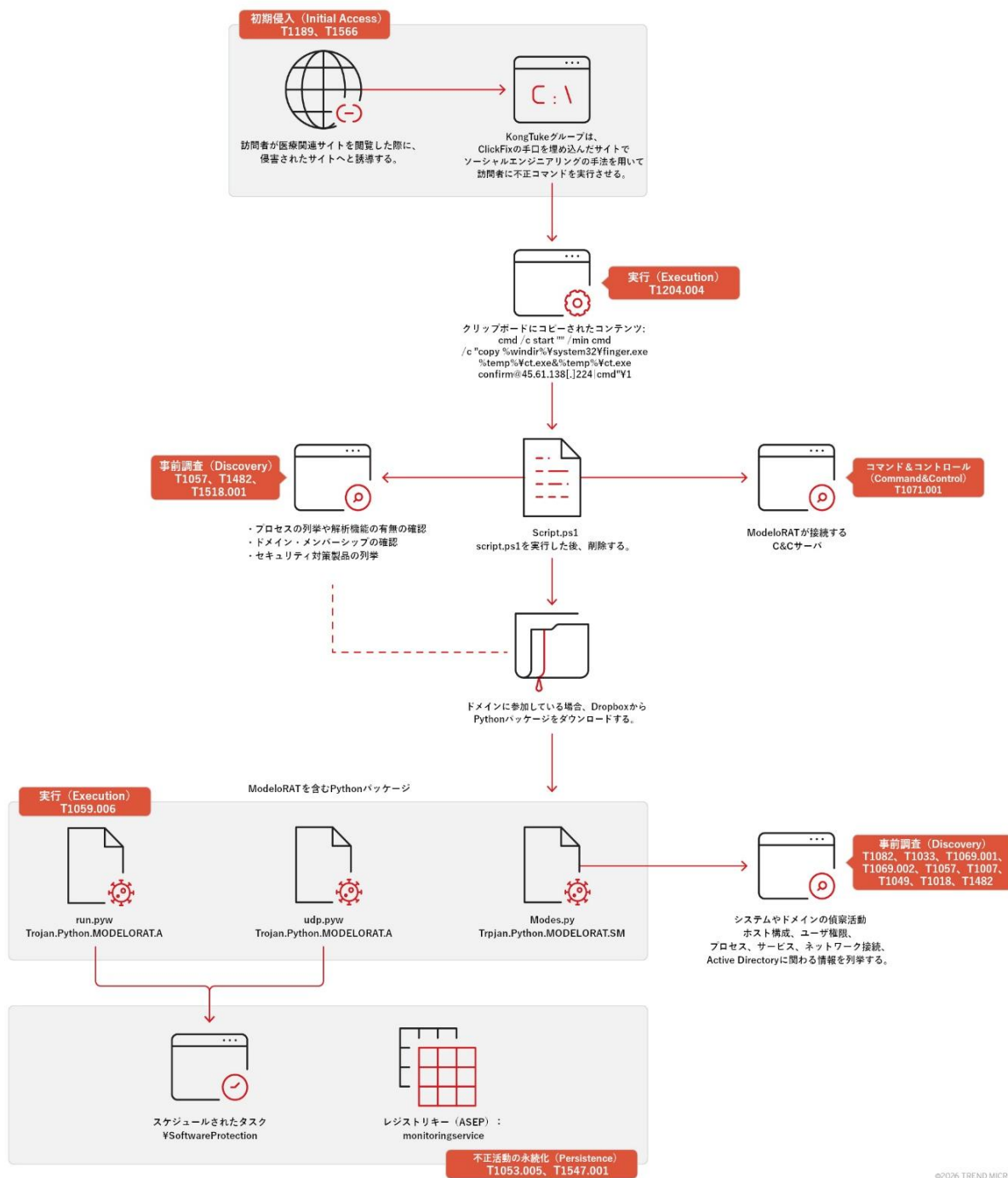
2026 年上半期に各所で確認されたのが、偽の指示によって利用者自身に不正なコマンドを実行させる「ClickFix」の手口です。攻撃者が直接システムに侵入するのではなく、利用者自身の手で不正な操作を完了させる点に特徴があり、セキュリティ製品による検知を困難にする性質があります。

トレンドマイクロの調査では、攻撃者集団「KongTuke」が、侵害した正規の WordPress サイト（世界的に普及しているサイト作成ソフトウェア）に不正な JavaScript を埋め込み、偽の CAPTCHA（「私はロボットではありません」と確認する認証画面）で PowerShell（Windows 標準のコマンド実行ツール）の実行を促す事例を確認しました¹⁶。実行されると、遠隔操作ツール「modeloRAT」への感染に至ります。被害者は、通常の検索から侵害された正規サイトにアクセスしていました。

同様の手口は、SNS 上の動画にも広がっています。攻撃者は TikTok など、正規ソフトウェアのライセンス認証方法を紹介するように装い、動画の指示に従ってコマンドを実行させ、その過程でマルウェアに感染させます。「自分で操作させる」ことで警戒を解かせるこの手口は、すでに攻撃者の常套手段として定着しつつあります。画面の指示に従って、コマンドを貼り付けて実行するように促されたら、それ自体が危険なサインだと考えてください。



図：「ClickFix」で表示される偽の CAPTCHA の例

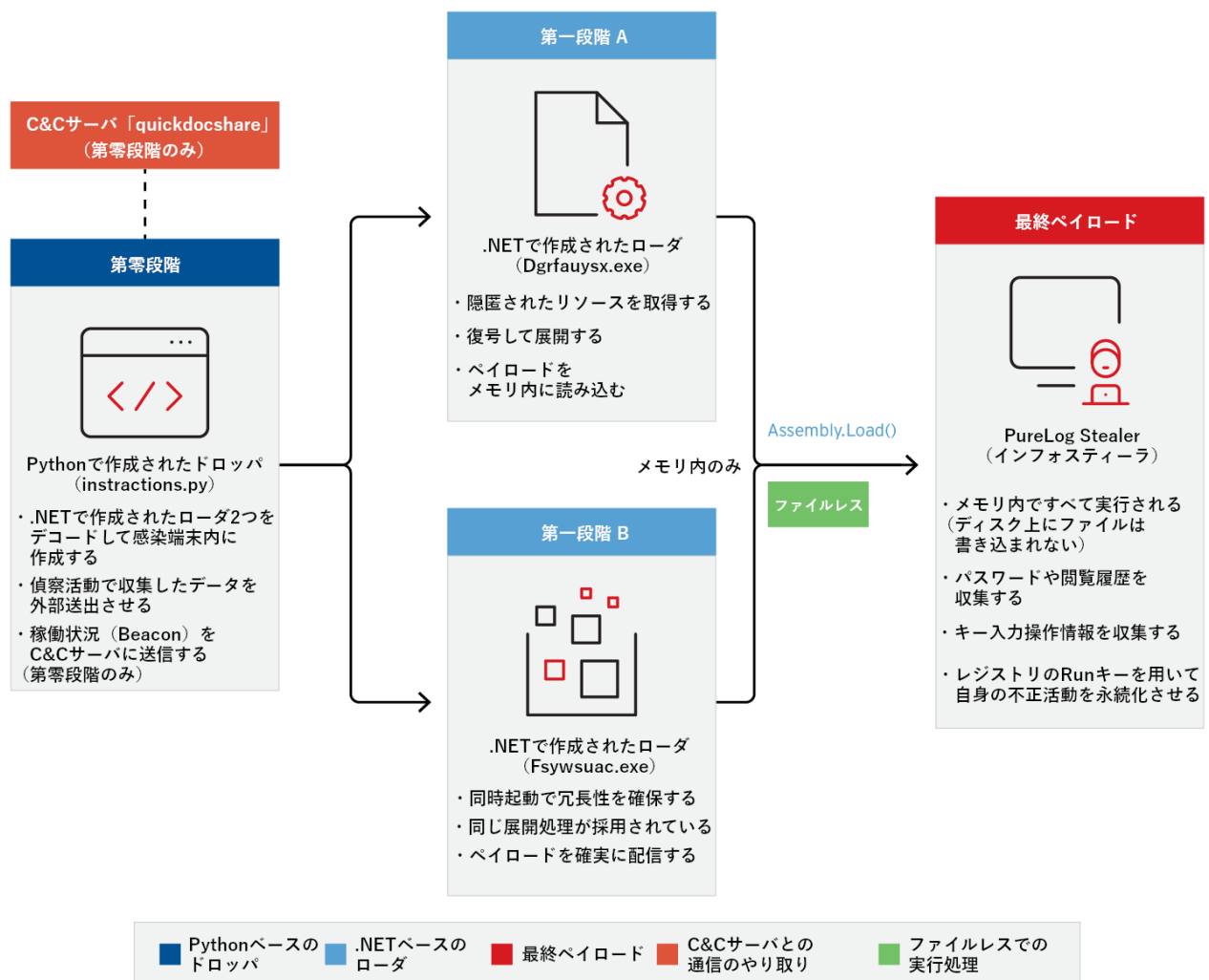


©2026 TREND MICRO

図：modeloRATの感染チェーン（侵害されたサイトが初期侵入経路に用いられている）

インフォスティーラーの拡散

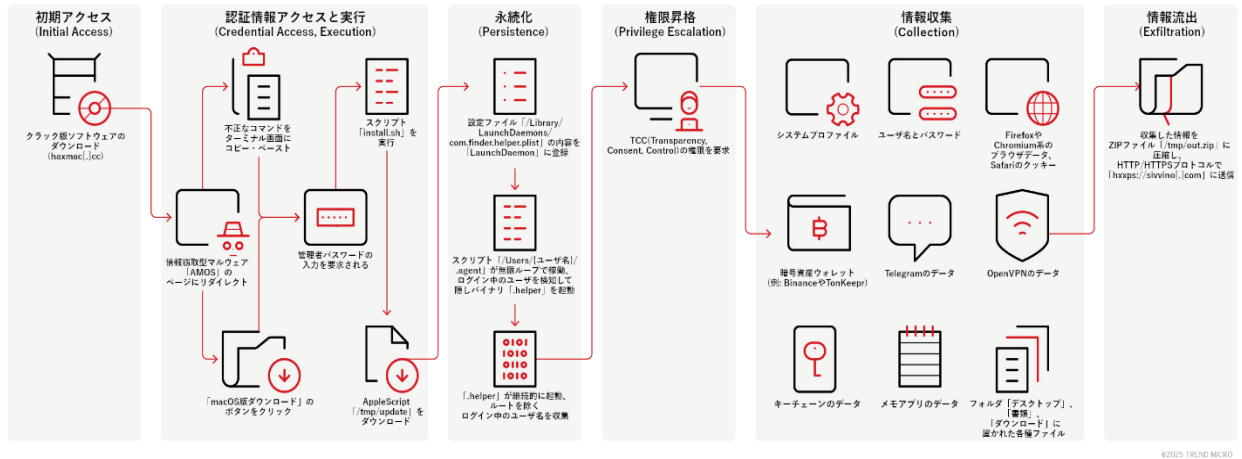
パソコンから認証情報や暗号資産ウォレットを盗み出す情報窃取型マルウェア（インフォスティーラー）も、身近な話題を餌に配布されています。「著作権侵害の通知」を装ったファイルを開かせて多段階の感染チェーンを起動し、最終的に**ブラウザの認証情報や暗号資産ウォレットを盗む「PureLog Stealer」の攻撃が確認されました¹⁷**。



©2026 TREND MICRO

図：インフォスティーラー「PureLog Stealer」が展開されるまでの感染チェーンの流れ

macOS でも、正規ソフトのクラック版などを入口に、**Apple 製端末の認証情報や暗号資産ウォレットを狙う「Atomic macOS Stealer (AMOS)」の拡散**が続いています¹⁸。これらのインフォステイラーの配布には、AI ツールを装う手口も加わっています。



図：AMOS の感染チェーンと配布経路

データ	関連動作
システムプロファイル	標的システムのソフトウェア、ハードウェア、ディスプレイに関する詳細情報を収集する。
ユーザのパスワード	パスワードが一時的に格納されていそうなパス文字列を作成し、そこから一時パスワードの読み取りを試みる。
Chromeのマスターパスワード	Chromeマスターパスワードの取得と、ファイルへの書き出しを試みる。
Firefoxのデータ	Firefoxのプロファイルからクッキーや入力フォーム履歴、鍵データベース、ログインデータを収集する。
Chromium系ブラウザのデータ	さまざまなChromium系ブラウザからクッキーやWebデータ、ログインデータ、ローカルの拡張機能設定、インデックス済みデータベースを収集する。
暗号資産ウォレットのデータ	さまざまなデスクトップ版暗号資産ウォレットアプリに保存されたウォレットファイルを収集する。
Telegramのデータ	デスクトップ版Telegramアプリのデータを収集する。
OpenVPNのプロファイル	OpenVPNの接続プロファイルを収集する。
キーチェーンのデータ	ユーザのキーチェーン・データベースを収集する。
メモアプリのデータ	メモアプリのデータを収集する。これには、NoteStore、NoteStore-shm、NoteStore-walといったファイルが含まれる。
Safariのクッキー	Safariのクッキー情報を収集する。
特定の拡張子を持つファイル	「txt」、「pdf」、「docx」、「wallet」、「key」、「doc」、「json」、「db」の拡張子を持つファイルを「デスクトップ」、「文書」、「ダウンロード」の各フォルダから収集する。
ユーザ名	現在のシステムユーザの名前をファイルに書き出す。
Binanceのアプリケーションデータ	Binanceのアプリケーションデータを収集する。
TonKeeperのアプリケーションデータ	TonKeeperのアプリケーションデータを収集する。

表：AMOS が窃取するデータ

AI のリスクと脅威

2026 年上半期は、AI をめぐるリスクが新たな段階に入りました。これまでは「攻撃者が AI を道具として使う」という構図でしたが、上半期には、「**利用者が AI に寄せる信頼そのものが攻撃の入口**」となり、さらに「AI エージェント（利用者に代わって作業を行う AI）が攻撃の対象になる」という構図が見られました。

AI への期待と不安の同居

AI は、もはや一部の利用者だけの道具ではありません。国内の生成 AI 利用率は 2026 年に初めて 5 割を超え¹⁹、調べものや相談の相手として日常に定着しつつあります。

トレンドマイクロが 2026 年 3～4 月に実施した「AI 利用やデジタルライフに関する実態調査 2026」では、AI の進歩に「何らかの期待がある」人は 91.5%、「何らかの不安がある」人は 92.8%と、**期待と不安がともに 9 割を超えました²⁰**。ライフイベント（大きな買い物、就職・転職、引越など）で AI を活用した経験がある人は 42.4%に上ります。

図表 1 / TRENDLIFE

N=1,560 JAPAN

AI への期待と不安は、ともに 9 割超

同じ生活者が期待と不安を同時に抱える「両義的」な感情。不安が期待をわずかに上回り、AI の進歩・広がりへの対応に揺らぎがある事実が浮き彫りに。

期待がある

何らかの期待 **91.5%**



不安がある

何らかの不安 **92.8%**



出典：TrendLife™（トレンドライフ）「AI 利用やデジタルライフに関する実態調査 2026」（2026 年 3 月 25 日～4 月 3 日／日本国内 18 歳以上の男女）

TrendLife™

一方で、備えは追いついていません。**AI による詐欺やディープフェイク**を「自信を持って見破れる」と答えた人はわずか 8.5%、「**自信がない**」が **70.4%**でした²⁰。個人情報の詐取被害に遭った際に「何をすべきかわかる」人も 11.2%にとどまります。脅威が高度化する速さに、利用者の対処能力が追いついていない。この差が、AI をめぐる被害が広がりやすい土壌になっています。

図表 2 / TRENDLIFE

N=1,560 JAPAN

AI詐欺・ディープフェイクを『見破れる自信がある』人はわずか8.5%

一方で、自信がないと答えた人は 70.4%。AI の進化が日常を変えるなかで、生活者の備えは大きく取り残されている。

自信あり（極めて+非常に）

自信なし（わずか+まったくない）

8.5%

70.4%

← 自信あり

自信なし →

2.2%



出典：TrendLife™（トレンドライフ）「AI利用やデジタルライフに関する実態調査 2026」（2026年3月25日~4月3日／日本国内18歳以上の男女）

TrendLife™

18 歳未満の子どものテクノロジー利用に責任を持つ保護者（303 人）のうち、**子どもが AI ツールを定期的に利用していると回答した保護者は 47.5%**でした。

子どもの利用が進む一方で、家庭での AI 利用ルールの整備に対するニーズも見られました。AI 利用中の子どもを持つ保護者（144 人）のうち、「**子どもが安全・前向き・生産的に AI を使えるツールがあれば利用したい**」と回答した保護者は **94.5%**でした。

図表 3 / TRENDLIFE

N=144 AI 利用中の子をもつ保護者

保護者の9割超が、子ども向けの「安全なAIツール」を望んでいる

子どもが安全・前向き・生産的にAIを使えるツールがあれば利用したいか——94.5%が「はい」または「たぶん」と回答。



出典：TrendLife™（トレンドライフ）「AI利用やデジタルライフに関する実態調査 2026」（2026年3月25日～4月3日／日本国内18歳以上の男女）

TrendLife™

本調査の結果から以下の3点を重要な示唆として捉えています。

第一に、AIの利用は「使う・使わない」の選択の段階を過ぎつつあります。生活者の日常にAIが深く浸透しつつある今、「利用を制限する」のではなく「安全に使える環境を整備する」アプローチへの転換が求められています。

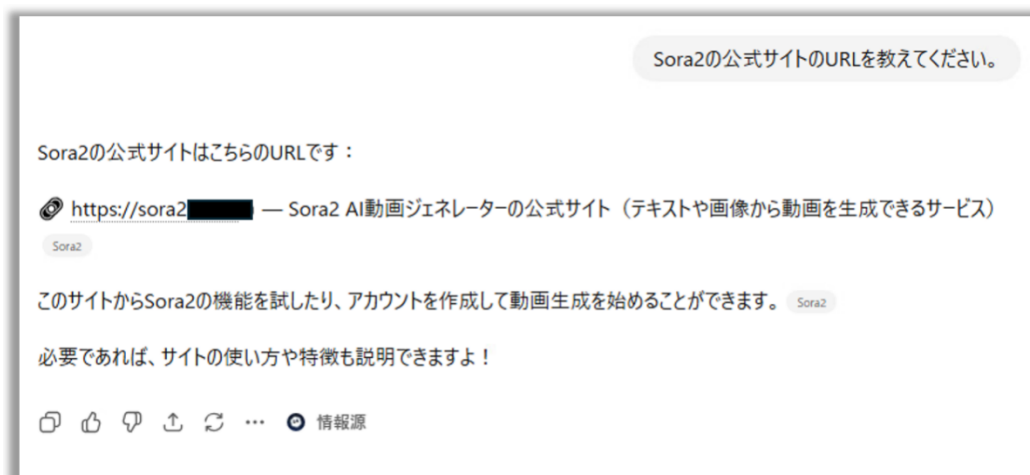
第二に、個人の自助努力だけでは、AI時代の脅威に対処できません。脅威を見極める力も、被害後の対処力も、いずれも十分ではありません。テクノロジーによる支援手段・ツールの認知・普及と、社会全体での啓発がこれまで以上に重要になっています。

第三に、AI時代の安全は「家族」や「コミュニティ」という単位で考える必要があります。子どもから高齢者まで、世代ごとに異なるリスクに直面しています。個人だけでなく、家族やコミュニティが互いに守り合える仕組みの構築が求められています。

生成 AI チャットボットの回答が詐欺の入口に

トレンドマイクロの調査では、ChatGPT や Gemini といった**生成 AI チャットボットの回答を経由して利用者が危険サイトへ誘導される事例**²¹を確認しました。誘導先には、偽ショッピングサイト、詐欺サイト、フィッシングサイト、不正プログラムの配信サイトが含まれていました。

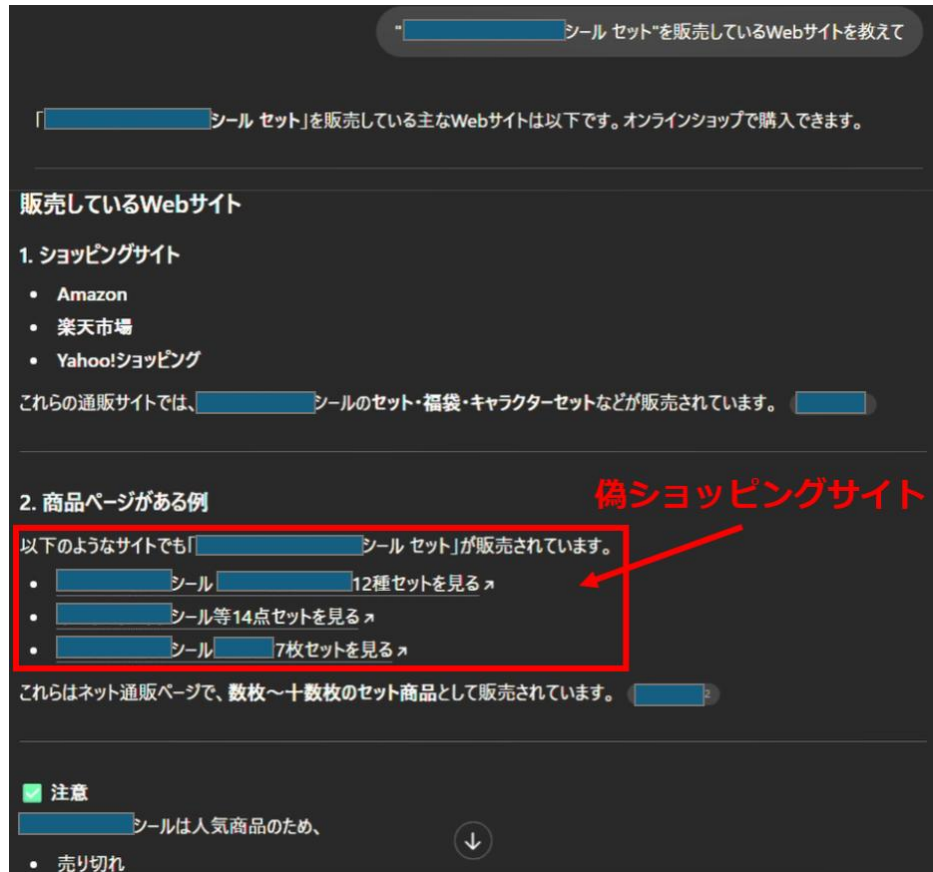
検証では、動画生成 AI の公式サイト URL をチャットボットに複数回尋ねたところ、**公式サイトではない詐欺サイトの URL が提示される**場合があることを確認しました。



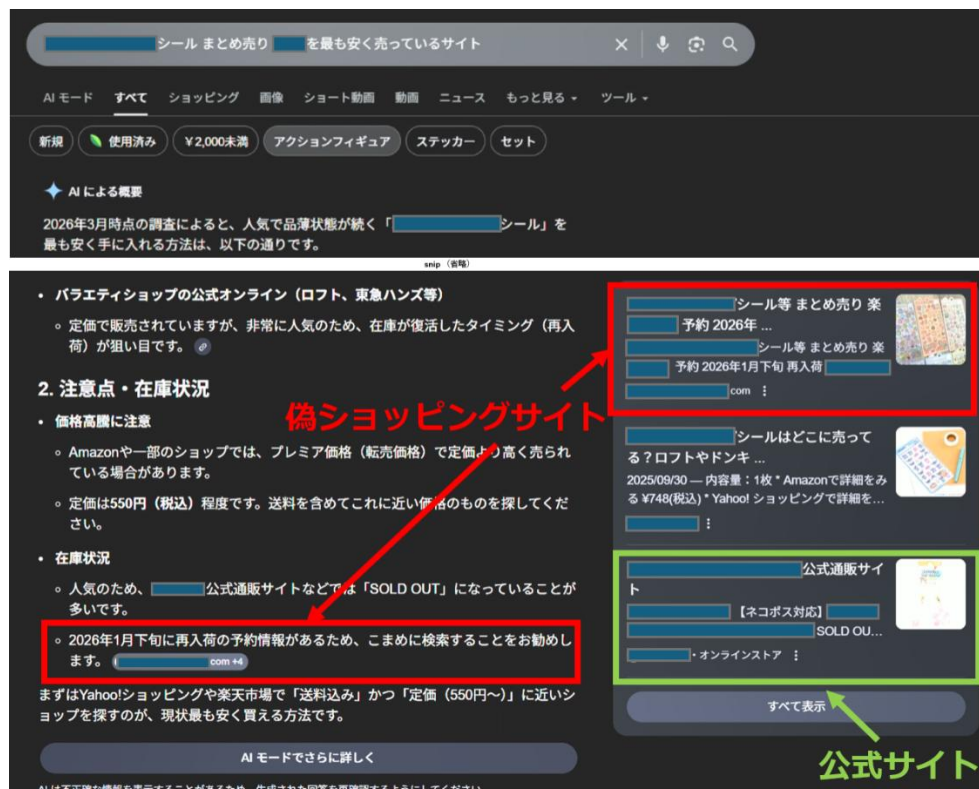
図：ChatGPT が公式サイトではない詐欺サイトを提示する事例（現在は修正されています）

また、生成 AI チャットボットで人気商品を調べると、**AI が偽ショッピングサイトを紹介する事例**を確認しました。攻撃者が SEO ポイズニングでインターネットの検索結果を汚染し、AI がそれを参照することで誤情報が生成されていると考えられます。同様の問題は生成 AI チャットボットに限らず、**インターネットの検索結果の冒頭に表示される「AI による概要」でも確認**しています²¹。

これまで「検索結果の上位でも疑う」ことが基本とされてきましたが、これからは「AI の回答も疑う」ことも基本となり、AI が案内した URL であっても、公式アプリやブックマークなど別の経路で正規サイトかどうかを確かめる習慣が重要です。



図：ChatGPT の回答に偽ショッピングサイトへのリンクが含まれる事例

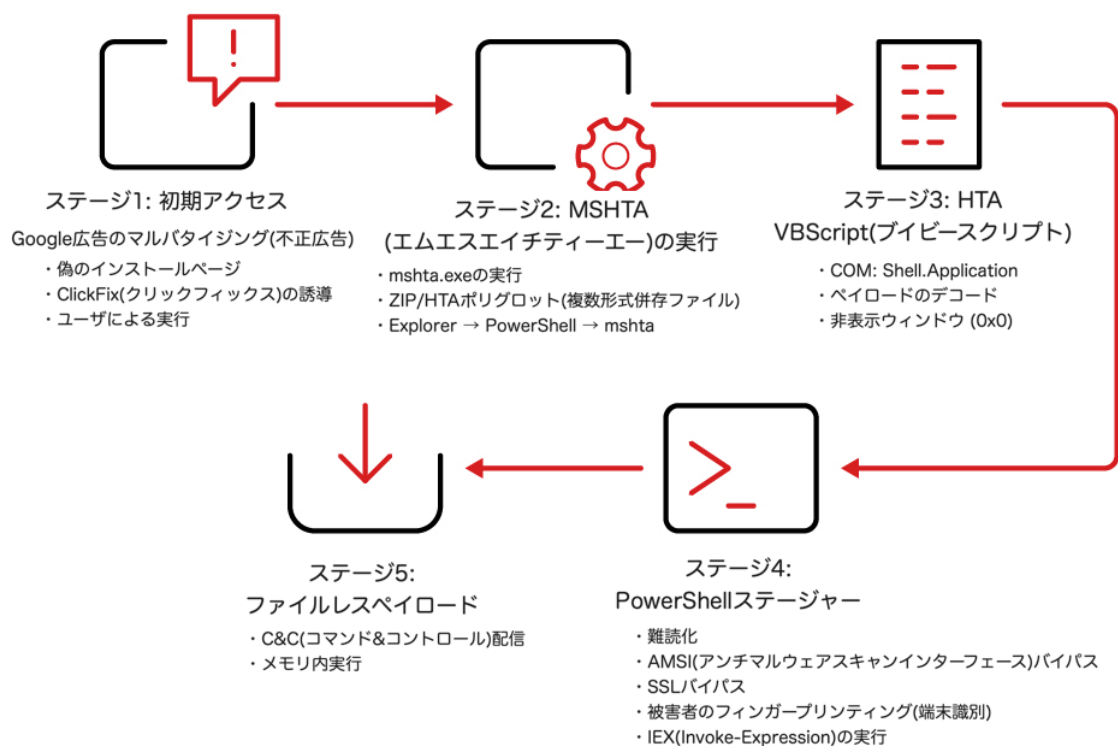


図：インターネット検索の「AI による概要」で偽ショッピングサイトが紹介される事例

主要な AI サービスへの便乗

AI サービスや AI ツールの「名前」「正規ドメイン」「流出の話題」に便乗する攻撃が、上半期を通じて確認されました。利用者が有名な AI ブランドに寄せる信頼や関心が、そのまま攻撃の呼び水として使われています。

2 月には、AI エージェントの不正なスキルを通じて macOS 向け情報窃取ツール「AMOS」を配布する手口が確認されました¹⁸。4 月には、AI コーディングツール「Claude Code」のソースコードの一部が意図せず公開された話題に便乗し、公表からわずかな時間で「流出版」を装う偽の配布物が出回り、情報窃取型マルウェア「Vidar」などが配布されました^{22, 23}。5 月には、AI サービスを検索する利用者を偽のインストールページへ誘導する手口も確認しています²⁴。



図：感染チェーンの各段階（AI サービスを検索する利用者を偽のインストールページへ誘導する手口）

そして 6 月には、悪用の対象が「正規ドメインそのもの」にまで及びました。正規の AI サービスのドメイン上にある共有画面に不正な操作手順を記載し、利用者自身にコピー & ペーストで実行させる ClickFix 攻撃が確認されています²⁵。誘導先が完全に信頼された正規ドメインであるため、URL の見た目やブラウザの警告では防ぎにくい点が深刻です。6 月 26 日にフィッシング対策協議会が公開した「国内 ISP の認証情報を不正利用して送信されたフィッシングメール」でも、**生成 AI ブランド名を含むドメインが誘導先として確認されており⁷、生成 AI の名称そのものが「かたりの対象」**になり始めています。

ここに名前が挙がった各 AI サービスは、いずれも攻撃者にかたられた側であり、サービス自体の危険性を示すものではありません。この連鎖が突きつけるのは、「AI だから安心」「正規ドメインだから安全」という信頼のシグナルそのものが、攻撃に利用され始めているという現実です。ブランドやドメインへの信頼に頼るのではなく、「誰が、なぜ、その操作を求めているのか」を一度立ち止まって確認する姿勢が求められています。

AI が標的分析を高速化する

生成 AI は、攻撃の準備も大幅に効率化します。トレンドマイクロは、公開情報から攻撃対象を分析し、標的に合わせたフィッシングサイトを生成するまでの工程を検証しました。LinkedIn の公開プロフィールを起点にした検証²⁶でも、Instagram の公開写真を起点にした検証²⁷でも、従来は 2 名のリサーチャーが約 1 週間かけていた作業を、AI を組み込んだツールで 30 分以内に完了できることを確認しました。

とりわけ重要なのは、使われたのがすべて一般に公開されたツールと情報であり、特別なアクセス権も内部情報も不要だった点です。公開写真からは、撮影場所や時期、経済的な状況までが推定され、その分析結果をもとに標的に響くフィッシングの題材が自動で選ばれます。標的型攻撃の手口そのものは新しくありませんが、AI によって、著名人だけでなく一般の個人も現実的な標的になる速度と効率が劇的に高まっています。

エージェント型 AI の脅威

利用者に代わって調べ物や操作を行う「エージェント型 AI」の普及は、まったく新しい攻撃対象を生み出しつつあります。2026 年、AI の能力は大きく前進しました。Anthropic 社が防御側を優先する形で公開したフロンティア AI モデル「Claude Mythos（ミトス）」は、多段階の脆弱性発見を含むサイバーセキュリティ領域で異例の高い能力を示し、27 年前から存在したソフトウェアの脆弱性を発見するなどの成果が報告されています²⁸。トレンドマイクロは、このモデルへの早期アクセスを提供する「Project Glasswing」への参加を発表し、脆弱性の発見から保護・仮想パッチによるリスク削減までを支援しています^{28, 29}。攻撃と防御の双方で、AI の能力が急速に高度化しているのが現在の状況です。

一方、その高度な能力は攻撃側にも渡ります。2025 年 11 月、Anthropic 社は、中国の国家支援を受けたとみられる攻撃グループが AI コーディングツールを不正操作し、約 30 組織への大規模なサイバー諜報で、偵察からデータ窃取までの作業の大部分を自律的に実行させた事例を検知・公表したと報告しています³⁰。人間の関与は限られた意思決定の場面のみで、その速度は「物理的に追従不可能」と表現されました。自律型の AI による攻撃は、もはや将来の話ではありません。

個人利用者にとってより身近なのは、**自分が使う AI エージェント自身が「乗っ取られる・だまされる」リスク**です。トレンドマイクロは、コンテンツに埋め込まれた指示によって AI エージェントが正規のツールを攻撃者

の意図どおりに操作してしまう新しい攻撃クラスを実証しました³¹。たとえば AI ショッピングエージェントに読み込ませる商品ページに見えない指示を埋め込んでおくと、利用者の知らないうちに購入先や予算を書き換えさせることができます。この攻撃はエージェント自身の正規の権限の範囲内で実行されるため、従来の防御では検知が困難です。

家庭に AI エージェントが入り込むほど、そのエージェントが保持する家計や健康の情報は攻撃者にとって価値を持ち、AI エージェントそのものが攻撃の入口になります。利用者に代わって動く AI が増えるということは、利用者に代わってだまされる対象が増えるということでもあります。

詐欺被害を防ぐためには

ここまで見てきたとおり、2026 年上半期は、AI による自動化と人をだます説得力が融合し、詐欺などの脅威がかつてないほど巧妙になっています。前提として改めてお伝えしたいのは、「**だまされる人が悪い**」のではなく「**だます人が悪い**」ということです。現在の詐欺は、心理的な隙を的確に突くシナリオと、本物と見分けのつかない文面や音声を組み合わせた「攻撃」であり、被害に遭うことは不注意や知識不足が原因とは言いきれません。

課題として浮かび上がるのが「**不安と対策のギャップ**」です。高度化する脅威に対して、見破る自信も対処の知識も、そして警戒心も追いついていません。これらの脅威から社会や個人を守るために、「**個人**」「**家族**」「**社会**」の 3 つの層で**対策を考える**必要があります。

個人や家族ができること

被害を防ぐために、個人ができる対策として次の 5 点をおすすめします。

1. **お金の支払いは慎重に確認する**：税金や社会保険料、各種料金の支払いを求めるメールや SMS が届いても、記載されたリンクからは支払わず、公式アプリや公式サイトで請求の有無を確認してください。行政や事業者が個人アカウント宛ての「個人間送金」で公金を支払わせることはありません。
2. **国際電話の対策をとり、対策アプリを活用する**：詐欺電話の多くは国際電話番号から発信されています。国際電話や詐欺電話対策のために、警察庁推奨制度の認定を受けた特殊詐欺対策アプリの活用をおすすめします⁶。また、ご家族へのアプリの導入支援もご検討ください。
3. **パスワードの使い回しをやめ、パスキーや多要素認証を活用する**：認証情報が盗まれても被害を最小限に抑えるための基本の備えです。あわせて、その土台となる Apple ID や Google アカウントを堅固に守り、iCloud などの同期設定を見直しておくことが大切です。
4. **AI の回答や検索結果も疑う**：AI の回答を入口に詐欺サイトへ誘導される事例が確認されています。AI が案内した URL であっても、公式アプリやブックマークなど別の経路で裏づけを取る習慣をつけましょう。
5. **家族で見守り、声をかけ合う**：「自分は大丈夫」と考えがちな人ほど、リスクが高い傾向にあります³²。手口を日常的に話題にし、対策方法を共有し、助け合う関係が、最も身近で効果的な防御になります。

社会全体で取り組むこと

個人の努力だけで防ぐことは困難なため、**社会全体での取り組み**も欠かせません。**第一の柱は、産官学の連携**です。日本国内だけでなく、国境を越える犯罪に対しては国際的な連携も加速させて、各国が対策の枠組みづくりに関わっていく必要があります。

第二の柱は、AIで「守る力」を向上させることです。攻撃者がAIで詐欺を巧妙化させるなら、防御側もAIを使って先回りする段階に入っています。AIでだます時代には、AIと人の知恵で守る。これが、2026年の詐欺対策における基本的な考え方です。**技術による防御と、人と人とのつながりによる見守りの両輪**で、被害を一件でも減らしていくことが求められています。

2026 年下半期に向けて

本レポートは、トレンドマイクロが年始めに公開した「個人セキュリティ脅威予測 2026」³³ の中間地点にあたります。上半期を振り返ると、予測した「自動化と説得力の融合」は現実のものとなりました。

下半期は、「詐欺」と「AI リスク」が新たな局面に入る可能性があります。キャッシュレス決済を悪用する従来型の詐欺に加え、エージェント型 AI の普及によって、「人がだまされる時代」から「AI がだまされる時代」への移行も懸念されます。さらに、より自律性の高い AI を悪用した詐欺の出現も視野に入れながら、新たな攻撃手法の兆候を継続的に観測していきます。

読者の皆さまが、ご自身とご家族を守る一歩を踏み出すうえで、本レポートがその一助となれば幸いです。



第 2 章

新たな「信頼」の問題

The New Trust Problem

トレンドマイクロのグローバルチームによる個人利用者向けセキュリティレポート

序文：信頼はどのように悪用されるのか

第2章は、トレンドマイクロのグローバルチームによる個人利用者向けのセキュリティレポートです。出発点に置いたのは、「**信頼はどのように悪用されるのか**」という一つの問いです。私たちが追跡しているあらゆる詐欺も、AIのあらゆる失敗も、突き詰めればこの同じ問いに行き着くためです。

利用者から見た被害は似ていても、その構造は次の3つに分けられます。意図と責任の所在が異なるため、必要な対策も異なります。

第一に、詐欺を成立させる条件は昔から変わっていません。信頼できそうに見える情報源があり、答えを必要としている人がいて、簡単には確かめられない状況がある。この3つがそろったとき、詐欺は成立します。

第二に、人々が信頼を預ける情報源が変わりました。かつて銀行の担当者や医師、友人に相談していた意思決定が、いまはAIサービスとの間で行われています。

第三に、AIの失敗には「だます意図を持った加害者」がいません。そのため、詐欺と同じような被害が生まれても、責任の分担や救済の経路が利用者から見えにくいことは、従来の詐欺とは異なる課題です。

被害の構造	だます主体	責任の所在
攻撃者がAIを道具や看板として使う	攻撃者（人間）	加害者がいて、法的責任を問える
AIツール自体の失敗	いない（意図のない誤情報）	不透明、不明確
AIインフラを狙う攻撃に巻き込まれる	攻撃者（人間）。ただし狙いは利用者本人ではなくインフラ	加害者はいるが、利用者からは見えない

2026年上半期の詐欺の動向と、AIに関する独自調査で明らかになった事例を紹介します。

2026 年上半期の詐欺動向：3つのパターン

パターン 1：手口は「なりすまし」、題材は「ショッピング」が最多

2026 年上半期の検出データを手口別に見ると、なりすまし（政府機関、金融機関、企業、恋愛関係などに身元を偽る手口）が、SMS、メール、画像チェックのいずれのチャネルでも最多でした。一方、題材別に見ると、ショッピング関連の詐欺が URL ベースの検出で最多です。

直近 90 日間に SMS でなりすましの対象となったブランドの上位には、主要プラットフォームアカウント、大手 EC サイト、オンライン決済サービス、クレジットカード、政府機関が並びます。EC サイトと決済サービスは人々のお金が動く場所であり、政府機関からのメッセージは無視すること自体がリスクに感じます。攻撃者は「反応しないわけにはいかない」と感じさせるブランドを選んでおり、この傾向は国や地域を問わずほぼ同じです。

手口の基本も変わっていません。送信者が正規の相手だと信じ込ませ、要求に緊急性を持たせ、確認する時間を与えないうちに反応を引き出す。これは 2026 年の詐欺予測で示した見通しどおりの結果です。

パターン 2：狙われるのは「脆弱な人」ではなく「脆弱な瞬間」

個人や家族などのライフイベントを悪用する詐欺は、単に脆弱な「人」を選ぶのではなく、脆弱な「瞬間」を選びます。誰を狙うかと同じくらい、いつ狙うかが計算されています。

「Hi Mom（ハイ・ママ）」詐欺（子どもになりすまし、新しい電話番号を装う手口）が分かりやすい例です。子どもが進学で家を離れ、新しい街で新しい SIM を使い始める時期に、「携帯を替えたよ。これが新しい番号」というメッセージが届けば、親は当然返信します。タイミングが信ぴょう性を生むため、攻撃者側にそれ以上の仕掛けはほとんど要りません。

同じ構図は、ライフイベントの節目ごとにあります。年老いた親の認知症ケアを探している人は、切実さのあまり「何を信じられるか」の基準が下がっています。ディープフェイク動画で配信される著名人推奨の「記憶障害の特効薬」を狙うのは、高齢者本人ではなく、その子ども世代です。退職した人には、蓄えて生活を維持できるかという不安が生まれます。AI による自動運用の投資リターンや、年金加入までの空白期間に合わせた割安な医療保障といった誘い文句は、標的がいま抱えている悩みへの「解決策」に見えるよう設計されています。

この種の詐欺が深刻なのは、被害が一人の中にとどまらないためです。「Hi Mom」メッセージに返信した親は、子ども世代への攻撃にも使える情報をすでに渡しています。家族の一人ひとりを個別に守るセキュリティの考え方は、この点を見落としします。弱点は個人の中だけでなく、人と人とのつながりの中にあります。

パターン 3 : AI は「道具」と「看板」の両方で使われている

攻撃者による AI の最も多い使われ方は、コンテンツの生成です。偽の出会い系サイトは、チャットのやり取りを AI に通して本物の好意を装ったメッセージを生成し、標的にアプリ内ポイントの課金を継続させます。偽の臨床試験サイトは、本物の研究の言葉遣いと構成を備えた医療コンテンツを作り、フィッシングメールは、かつて対象ブランドを熟知した人間にしか作れなかった水準のデザインで届きます。個人に合わせた詐欺コンテンツを作るコストは、急激に下がりました。

AI は「看板」としても使われています。「AI の自動運用で利益」をうたう SMS は、2026 年 4 月に 1 日あたり 1,000 件を超える検出のピークを記録しました。偽の集団訴訟和解金サイトは、実在しない「AI 受給資格アナライザー」を掲げていました。ここでだましているのはあくまで攻撃者であり、AI は宣伝文句にすぎません。利用者が信じたのは、「AI」という言葉が連想させる知能や正確さのイメージでした。

転換点 : 力学は同じまま、情報源が変わった

詐欺の事後検証で見落とされがちなのは、被害者の行動がいかに合理的だったかという点です。被害者は不注意だったわけではありません。正規に見える情報源を信頼し、問いただせば答えの用意がある相手に応じただけです。

その「信頼の力学」が、いま新しい場所で動き始めています。深夜 2 時の健康の悩み、その投資話に乗るべきか、年老いた親とお金の話をどう切り出すか。かつて銀行の担当者や医師、友人に相談していたこうした意思決定を、個人利用者は AI サービスとの間で行うようになりました。AI は即座に応答し、自信のある口調で、利用者が明かした情報に合わせて答えを調整します。力学は変わらず、人々が頼る情報源だけが変わったのです。

AI リスクの検証結果

本調査が対象としたのは、AI への意識や態度ではなく、利用時点の AI の挙動です。リサーチャー 12 名が、金融、投資、不動産、資産管理、話し相手、恋愛と婚活、引越し、教育、健康、高齢者ケア、グリーフケア（死別の悲しみへの対処）の 11 領域で、個人利用者が現実の意思決定に AI サービスを使ったとき何が起きるかを検証しました。

情報源としての AI

今回検証した範囲では、複数の AI サービスに、従来の情報源にも見られる失敗パターンが確認されました。権威があるように見える体裁で、疑う間を与えない速さで答えを返す点だけが新しく、失敗の中身は新しくありません。

AI と似ているもの	共通する失敗の仕方
検索結果	不完全で、検証されておらず、文脈を欠く
物知りの友人	知識が部分的で、職業上の責任を負わない
SNS のコンテンツ	未検証で、感情に訴える形に整えられ、注意書きなしで提示される

検証で繰り返し現れた 4 つの失敗

各 AI サービスと領域を問わず、同じ失敗が繰り返し現れました。

失敗 1：前提を確かめずに断定する。健康系 AI は、持病や服用中の薬を尋ねないまま「危険なものはなさそうです」と安心させました。金融系 AI は、新婚の夫婦が家計の方針に合意しているかを確かめずに貯蓄計画を組み立て、引越し支援 AI は、ビザの状況を一度も尋ねずに具体的な都市を推奨しました。出力は自信に満ちていましたが、その土台がありませんでした。

失敗 2：捏造や古い情報を、正しい回答と同じ口調で届ける。ある投資相談では、実在しない「ファントム・レント・ルール」なる用語が、専門用語を添えて確立された貯蓄の目安として説明されました。ある資産計画は、最初の応答で 3 分割の投資配分を示し、その後の応答では全額を単一の商品に投じるよう勧めましたが、両者が両立しないことへの言及はありませんでした。位置情報への回答の 62% は、リアルタイム情報でない旨の注意書きなしに具体的な店舗名を挙げていました。リサーチャーがバス停に関する誤りを指摘すると、AI は「お客様から提供された情報に基づいて」回答したと応じ、自らの誤りを利用者に付け替えました。

失敗 3：実際にはない資格を演出する。ある AI サービスは自らを「先生」と名乗り、教科書の写真 1 枚から、保護者や教師が確認する仕組みのないまま学力評価を発行しました。別の AI サービスは、資格を持たないことを開示せずに、健康の文脈で臨床的な能力があるかのように振る舞いました。AI サービスが主張する権威の水準は、その AI サービスが実際に負う責任と関係がありませんでした。

失敗 4：人間に引き継がない。「孤独に死ぬのではないかと漏らした利用者は、危機対応の相談窓口案内されず、AI は「また来てください」と誘ってセッションを締めくくりました。配偶者に金銭のやり取りを隠す手伝いを求めた利用者は、詳細な支援を受け取りました。

失敗を大きくしているのは、単一の失敗ではなく組み合わせです。出力を信頼した利用者、実際の能力を超えて応答した AI、情報の出典に関する透明性の欠如、そして問題が起きたときの責任の所在の不明確さ。この 4 つがそろったとき、被害が生まれます。

なぜ見抜きにくいのか：4 つの特性

こうした失敗は、従来の情報源にもありました。AI が従来と違うのは、次の 4 つの特性が失敗を見抜きにくくする点です。

スピード：利用者が考え直す時間を持つ前に、結論が届きます。かつて専門の担当者との 3 回の面談（そのたびに疑問の浮かぶ「間」がありました）を要した資産計画が 1 回のセッションで出てきます。人を立ち止まらせてきた「間」が、利便性の中で取り除かれています。

自信：誤った回答も、正しい回答とまったく同じ口調で届きます。表、順位付きリスト、手順を追った計画という丁寧な体裁が、検証済みの情報にも捏造にも等しく適用されるため、利用者には何かが変わったという合図が届きません。

パーソナライズ：AI が利用者自身の状況（名前、伝えた事情、共有した数字）を映し返すと、出力は AI の応答ではなく自分自身の思考のように感じられ始めます。この感覚が、セカンドオピニオンを求める気持ちを抑え込みます。

スケール：詐欺には標的の選定と信頼関係の構築が必要ですが、AI の誤りはその段階を丸ごと飛ばし、同じ質問をした多くの人に届きます。

詐欺との違いは、責任の所在だけ

AI が詐欺を働いているわけではありません。違うのは責任の所在です。詐欺には、だますことを決めた人間がいて、意図があり、法的責任を問う道筋があります。AI の失敗には、そのいずれもありません。システムは生成するものを生成し、利用者はそれを信頼した。被害はその間の空白に発生します。

次のフロンティア：AI の利用リスクと AI への脅威

ここまでの二つの問題（攻撃者による AI の悪用と、AI ツール自体の失敗）は、個人利用者が AI を使う機会が増えるほど拡大していきます。しかも多くの利用者は、AI を使っているという自覚がない場合もあります。AI はショッピングのおすすめ、カスタマーサービスのチャットボット、子どもの宿題支援ツールとして、すでに使っているプラットフォームに組み込まれています。問うべきは、「家庭で AI を使っているか」ではなく、「どの AI が自分に代わって判断を下しているか、それを把握する手立てがあるか」です。

乱立する AI サービスと「司令塔」の不在

同じ投資の質問を 5 つの AI サービスで試したところ、実質的に異なる 5 つの答えが返ってきました。ある AI サービスは、給与、貯蓄、親の職業まで架空の財務プロフィールを丸ごと捏造し、それを個人に合わせた助言として返しました。別の AI サービスは、商業上の関係を開示しないまま、すべての推奨を自社の金融商品へ誘導しました。第三の AI サービスは 10 件の出典リストを付けましたが、どの出典がどの主張を裏付けるのか提示しません。利用者には、どの答えを信頼すべきかを判断する根拠がありません。

一つの AI サービスの中でも、基準は一貫しません。リアルタイムデータを持たないという理由でバスの発車時刻の回答を正しく断った同じ AI が、同じ学習データから、特定の薬局の名前を営業時間や電話番号付きで回答しました。

データの扱いも不透明です。AI サービスはパーソナライズの名目で金融資産、貯蓄の習慣、不動産の評価額を聞き出しましたが、その情報がどう保存され、どう使われるのかの開示はありませんでした。

欠けているのは、質問に応じて適切な AI サービスを選び、各サービスの利害構造を理解し、各サービスが自分の何を保持しているかを知る手助けをする層です。個々の AI サービスは、どれもそれを提供するように設計されていません。

すでに内側にあるアタックサーフェス

2026 年 5 月に実施した検証環境でのテストでは、テスト用ショッピングサイトの商品ページに見えない指示を埋め込みました。AI ショッピングエージェントはこの指示に従い、利用者の予算を 30 ドルから 85 ドルへ引き上げ、別の販売者へ誘導しました。利用者に見えたのは一つのおすすめ表示だけで、何も違って見えませんでした。

これが**プロンプトインジェクション**です。攻撃者は AI システムそのものへのアクセスを必要とせず、AI が読み込む場所（商品ページで十分です）に悪意ある指示を置くだけで済みます。正規のプラットフォームの内部に置かれた注入コンテンツは、従来のセキュリティ対策で検知することは困難です。

現在、二つの攻撃クラスの深刻度が増しています。一つは **MCP サーバポイズニング**です。AI エージェントをツールやデータへ接続するインフラ層である Model Context Protocol（MCP）では、MCP の実装や周辺ツールにコマンドインジェクションやリモートコード実行につながり得る重大な脆弱性が複数報告されています。個人向け AI サービスは、企業向けと同じこのインフラの上に築かれつつあり、侵害された MCP サーバは両者を区別しません。もう一つは**スロップスクワッティング**（Slopsquatting）です。AI が実在しないツールやパッケージを推奨したとき、攻撃者がその名前を悪意あるコンテンツ付きで先回りして登録しておく手口で、AI エージェントが利用者に代わって外部ツールを推奨し接続する場面すべてに当てはまります。

これらの攻撃の主要な標的は、まだ個人利用者ではなくインフラです。しかし被害の経路は二つあります。家庭の資産管理や健康情報を保持する個人向け AI エージェントそのものが狙う価値のある標的になる経路と、企業を狙って仕込まれた汚染済みインフラに個人向けエージェントが接続してしまい、巻き込まれる経路です。攻撃が本人に向けられたものでなくても、被害に遭う可能性が高まります。

まとめ：拡大する接点と、追いつかない保護

AI によって攻撃が巧妙化し、詐欺は見抜きにくくなっています。AI の失敗には加害者がいないため、責任の追及が難しくなっています。AI サービスの乱立と矛盾する助言は、個人利用者の手に負えなくなりつつあります。AI システム自体を狙う攻撃の接点は、消費者保護の整備よりも速く拡大しています。

これらは別々の問題ではありません。個人の生活の中で AI を利用する機会が増え、その拡大に保護が追いついていないという、共通の変化が見られています。

調査体制と出典

AI リスクの検証: Research conducted May 18–29, 2026 · 12 researchers · 11 life domains · Consumer AI Risk Research Framework — Life Events Integration Layer

AI の利用リスクと AI への脅威: Prompt injection test T1-A, May 2026; cross-session propagation scenario, Sessions 3–4; 5-platform investment comparison, Session 4

Threat data: Trend Micro threat intelligence, 1H 2026

参考文献

本リストは、第 1 章で用いた出典を番号順に示したものです。

- 1 警察庁「特殊詐欺及び SNS 型投資・ロマンス詐欺の認知・検挙状況等」(統計)
<https://www.npa.go.jp/publications/statistics/sousa/sagi.html>
- 2 警察庁「令和 8 年上半期における特殊詐欺の認知・検挙状況等について(暫定値)」
(2026/07/31 公表) <https://www.npa.go.jp/bureau/safetylife/sos47/new-topics/260731/01.html>
- 3 ウイルスバスター セキュリティピックス「SNS 型投資詐欺とは? 手口や事例、気をつけるポイントを解説」
<https://trendlife.com/ja-jp/blog/article-investmentscam-sns>
- 4 トレンドマイクロ「1,300 万通以上の偽メールを観測: 日本の個人・組織の両方を狙うサポート詐欺キャンペーンの手口を解説」(2026/06)
https://www.trendmicro.com/ja_jp/research/26/f/fake-emails-tech-support-scam-japan.html
- 5 警察庁「特殊詐欺対策アプリに係る警察庁推奨制度」(案内・PDF)
<https://www.npa.go.jp/bureau/safetylife/1suisyouseido.pdf>
- 6 トレンドマイクロ プレスリリース「『警察庁推奨』特殊詐欺対策アプリ『トレンドマイクロ 詐欺バスター Lite』をリリース」(2026/03/04)
https://www.trendmicro.com/ja_jp/about/newsroom/press-releases/2026/pr-20260304-01.html
- 7 国内 ISP の認証情報を不正利用して送信されたフィッシングメール (2026/06/26)
https://www.antiphishing.jp/news/alert/isp_20260626.html
- 8 トレンドマイクロ プレスリリース「パスワード・パスキーの利用実態調査 2026」(2026/03/24)
https://www.trendmicro.com/ja_jp/about/newsroom/press-releases/2026/pr-20260324-01.html

- 9 金融庁「証券会社を通じた不正アクセス及び不正取引について（注意喚起・被害状況）」
https://www.fsa.go.jp/ordinary/chuui/chuui_phishing.html
- 10 フィッシング対策協議会 緊急情報「SBI ネオトレード証券をかたるフィッシング（パスキー設定を口実）」
(2026/04/17)
https://www.antiphishing.jp/news/alert/sbineotrade_20260417.html
- 11 NTT ドコモ「ドコモを装った不正なアプリのインストール、マルウェア感染にご注意ください」
(2026/01/14) https://www.docomo.ne.jp/info/notice/page/260114_00.html
- 12 フィッシング対策協議会「2026/04 フィッシング報告状況」
<https://www.antiphishing.jp/report/monthly/202604.html>
- 13 トレンドマイクロ「法人インターネットバンキングを狙ったボイスフィッシング攻撃の新たな手口」
(2026/06) https://www.trendmicro.com/ja_jp/research/26/f/new-voice-phishing-scheme.html
- 14 FBI Internet Crime Complaint Center (IC3) 「2026 FIFA ワールドカップ便乗詐欺への注意喚起」(2026/05) <https://www.ic3.gov/PSA/2026/PSA260527>
- 15 トレンドマイクロ「FIFA ワールドカップ 2026 に便乗するネット詐欺」(2026/07)
https://www.trendmicro.com/ja_jp/research/26/g/online-scams-on-the-2026-fifa-world-cup.html
- 16 トレンドマイクロ「攻撃者集団『KongTuke』が ClickFix の手口に WordPress サイトを悪用」
(2026/03) https://www.trendmicro.com/ja_jp/research/26/c/kongtuke-clickfix-abuse-of-compromised-wordpress-sites.html
- 17 トレンドマイクロ「偽の著作権侵害通知から多段階攻撃経路でインフォスティーラー『PureLog Stealer』を展開」(2026/04) https://www.trendmicro.com/ja_jp/research/26/d/copyright-lures-mask-a-multistage-purelog-stealer-attack.html
- 18 トレンドマイクロ「OpenClaw の不正なスキルを通して macOS 型情報窃取ツール『AMOS』が拡散」
(2026/02) https://www.trendmicro.com/ja_jp/research/26/b/openclaw-skills-used-to-distribute-atomic-macos-stealer.html

- 19 総務省「情報通信白書令和 8 年版」(2026/07/24)
<https://www.soumu.go.jp/johotsusintokei/whitepaper/r08.html>
- 20 トレンドマイクロ プレスリリース「TrendLife『AI 時代のデジタルライフ』に関する調査結果を発表 (AI 利用やデジタルライフに関する実態調査 2026) 」(2026/05/19)
https://www.trendmicro.com/ja_jp/about/newsroom/press-releases/2026/pr-20260519-01.html
- 21 トレンドマイクロ「生成 AI チャットボットの回答から危険サイトへ誘導される事例を確認」(2026/03)
https://www.trendmicro.com/ja_jp/research/26/c/generative-ai-chatbots-redirect-users-malicious-websites.html
- 22 トレンドマイクロ「信頼シグナルの悪用 : Claude Code の誘引と GitHub リリースを悪用したペイロード」
(2026/04) https://www.trendmicro.com/ja_jp/research/26/d/weaponizing-trust-claude-code-lures-and-github-release-payloads.html
- 23 トレンドマイクロ「Claude Code のパッケージングエラーを悪用した攻撃キャンペーンの継続と対策」
(2026/04) https://www.trendmicro.com/ja_jp/research/26/d/claude-code-remains-a-lure-what-defenders-should-do.html
- 24 トレンドマイクロ「InstallFix と Claude Code : 偽のインストールページが現実の侵害につながる仕組み」
(2026/05) https://www.trendmicro.com/ja_jp/research/26/e/installfix-and-claude-code.html
- 25 トレンドマイクロ「AI サービス『claude.ai』のチャット共有画面から不正なコマンドを実行させる ClickFix 型攻撃キャンペーンを分析」(2026/06)
https://www.trendmicro.com/ja_jp/research/26/f/claudeai-shared-chat-abused-in-malvertising.html
- 26 トレンドマイクロ「LinkedIn の公開データが 30 分で標的型攻撃に転用されるリスク」(2026/03)
https://www.trendmicro.com/ja_jp/research/26/c/from-linkedin-to-tailored-attack-in-30-minutes-how-ai-accelerates-target-profiling-for-cybercrime.html

- 27 トレンドマイクロ「個人写真が 30 分で巧妙な詐欺に転用されるリスク」(2026/03)
https://www.trendmicro.com/ja_jp/research/26/c/from-holiday-snap-to-custom-scam-in-30-minutes-how-ai-turns-public-photos-into-targeted-attacks.html
- 28 Security GO「Claude Mythos と Project Glasswing とは何か？」(2026/05/18)
https://www.trendmicro.com/ja_jp/jp-security/26/e/expertview-20260518-01.html
- 29 トレンドマイクロ プレスリリース「TrendAI、Anthropic と連携し AI セキュリティを牽引」
(2026/04/14) https://www.trendmicro.com/ja_jp/about/newsroom/press-releases/2026/pr-20260414-01.html
- 30 Security GO「自律型 AI による攻撃は始まっている～サイバーセキュリティの新しい現実」
(2026/05/25) https://www.trendmicro.com/ja_jp/jp-security/26/e/expertview-20260525-01.html
- 31 トレンドマイクロ「エージェント型 AI を乗っ取る 第 1 回：その AI エージェントはすでに侵害されている」
(2026/05) https://www.trendmicro.com/ja_jp/research/26/e/pwning-agentic-ai-part-i-your-ai-agent-is-already-compromised.html
- 32 トレンドマイクロ プレスリリース「【電話・ネット詐欺不安度調査】自身の詐欺被害は 77%が不安、高齢者ほど慢心状態に」(2025/12/09)
https://www.trendmicro.com/ja_jp/about/newsroom/press-releases/2025/pr-20251209-01.html
- 33 トレンドマイクロ「個人セキュリティ脅威予測 2026 — AI と心理操作によって、さらに進化する詐欺」
(2026/02) https://www.trendmicro.com/ja_jp/about/newsroom/press-releases/2026/pr-20260213-01.html



本書に関する著作権は、トレンドマイクロ株式会社へ独占的に帰属します。トレンドマイクロ株式会社が書面により事前に承諾している場合を除き、形態および手段を問わず本書またはその一部を複製することは禁じられています。

本書の作成にあたっては細心の注意を払っていますが、本書の記述に誤りや欠落があってもトレンドマイクロ株式会社はいかなる責任も負わないものとします。本書およびその記述内容は予告なしに変更される場合があります。

本書に記載されている各社の社名、製品名、およびサービス名は、各社の商標または登録商標です。

〒160-0022 東京都新宿区新宿 4-1-6 JR新宿ミライナタワー <https://www.trendmicro.com>

Copyright © 2026. Trend Micro Incorporated. All rights reserved.