

金融サービス企業に 害を及ぼす 内在的脆弱性

内在的脆弱性を悪用する攻撃の概要です。

攻撃の内容、発生理由、防御方法、従来の方策が効かない理由を説明します。



内在的脆弱性は、多くの場合、顧客に以下のような機能を提供することで発生します。

- オンラインアカウントの作成とそのアカウントへのログイン
- 請求書の支払い
- 振込先の追加と送金
- アカウント間での送金
- クレジットの申請
- アカウント、ローン、その他の金融商品の詳細を記載した文書へのアクセス
- 住所の変更
- パスワードのリセット
- カスタマサービスとのチャット

金融サービス企業は、世界で最も標的とされる企業の1つであり、最も洗練され、十分な資源を持つ攻撃者に狙われています。 熟練したセキュリティチームが、他の点は完璧でも、一歩間違えただけで、数百万ドルの損失、批判、高額な罰金、数十年にわたるブランド毀損を引き起こす可能性があります。対策としては、セキュリティチームが [OWASP Top Ten](#) に記載されているような脆弱性パッチを適用する、レッドチームとブルーチームを結成する、堅牢なバグバウンティプログラムを実施する方法もありますが、これらだけでは不十分です。これらの対策はすべて、その方向性は正しいのですが、たとえ完璧に実行されたとしても、組織を永続的に保護できるまでには至りません。

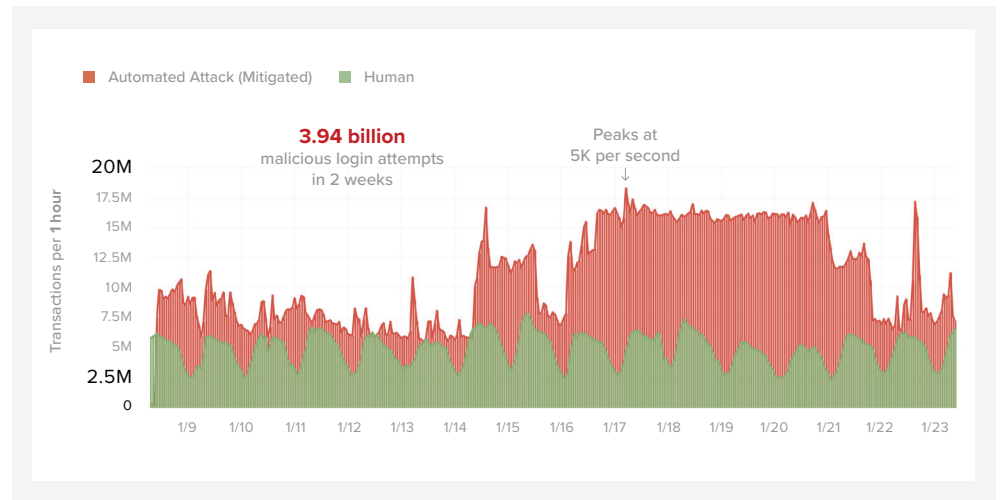
セキュリティチームだけでなく、組織全体が協力して、内在的脆弱性を悪用する攻撃を特定し、防御しなければなりません。これらの脆弱性は、コーディングエラーや不適切な設定などのミスではなく、有効であり、重要であることが多いビジネス要件に起因するため、従来のようなパッチを適用することはできません。ビジネスロジック攻撃やアプリケーションの悪用として知られているように、攻撃者は意図しない方法でこのような機能を悪用できます。

F5 は、金融サービスの顧客をオンライン不正行為や悪用から守っています。この保護の副産物として、毎日 F5 Network を通過する Web およびモバイルアプリケーションからの何十億ものトランザクションを可視化できます。この可視化により、F5 は、金融サービス業界全体の内在的脆弱性を狙った攻撃について、他に類を見ないインサイトを得ることができます。重要なことは、セキュリティチームが、これらの攻撃について、攻撃が起こる理由、防御方法、そして従来の対策だけでは効果がない理由を理解することです。

クレデンシャル検証攻撃（別名：クレデンシャルスタッフィング）

不正アクターは、ダークウェブを介して、またはセキュリティが不十分な組織への侵害により、何十万、何百万、あるいは何十億ものユーザー名とパスワードのペアを手に入れます。そして、まったく別の組織（侵害した元の組織の一部ではない組織）のログインアプリケーションに対して、自動化機能を利用してユーザー名とパスワードのペアを検証します。消費者はパスワードを使い回す習慣があるため、これらの攻撃は通常、0.1% から 3% の確率で有効な認証情報を特定できます。この成功は、多くの場合、アカウント乗っ取り（ATO）につながり、その後、侵害されたアカウント内のあらゆる資産が収益化され、不正行為による損失、顧客からの信頼の失墜、ブランド毀損につながります。

図 1：米国トップクラスの金融サービス企業
に対するクレデンシャルスタッフィング攻撃



金融サービス会社の資産には以下の
ようなものがあります。

- 通貨
- クレジット
- 他の悪意のある活動（マネーロンダリング、合成 ID、バストアウト詐欺など）のためのアカウントへのアクセス
- ローンやその他の金融商品に関する詳細
- 個人を特定できる情報（PII）

フィンテックを介した攻撃

攻撃者は、ログインアプリケーションに対して直接クレデンシャルスタッフィング攻撃を仕掛けるとは限りません。攻撃者は、ロイヤルティポイントまたは金融アグリゲータ、あるいは請求書支払いサービスなどのフィンテックを介して攻撃を仕掛けることがあります。たとえば、Mint、Prism、AwardWallet などです。

これらの無料サービスを利用するには、加入者はまず、それぞれのアカウントのユーザー名とパスワードをフィンテックに提供して、銀行、小売店、ホテル、航空会社、電気通信プロバイダなどでこれらのアカウントを「リンク」させる必要があります。その後、フィンテックは、プログラムによって各アカウントへのログインを試みます。加入者が正しいユーザー名とパスワードを提供していれば、リンクプロセスは継続します。アカウントのリンクが完了すると、フィンテックは自動化機能を使って、アカウントに繰り返しログインし、コンテンツをスクレイピングします。これは 1 日数千回行われることもあります。

不正アクターは、このリンクプロセスをクレデンシャルスタッフィングに悪用することを学びました。不正アクターは、フィンテックでアカウントを作成し、検証するユーザー名とパスワードをフィンテックに渡すだけです。リンクプロセスが続けば、ユーザー名とパスワードが正しいことを知ることができます。フィンテックは、このような行動を認識していますが、多くの場合、これを阻止するためにほとんど何もしていません。

図2は、米国トップ3の銀行におけるログイントラフィックを示しています。フィンテックのトラフィックにわずかな異常なスパイクがあることに注目してください。

図2：米国のトップ3の銀行におけるログイントラフィック（緑は人間によるログイン、赤はフィンテックによるログイン）

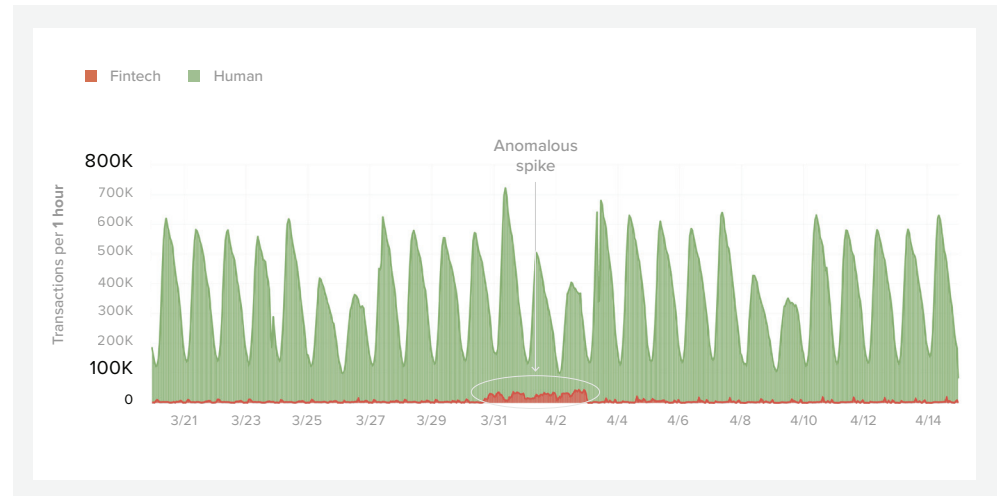
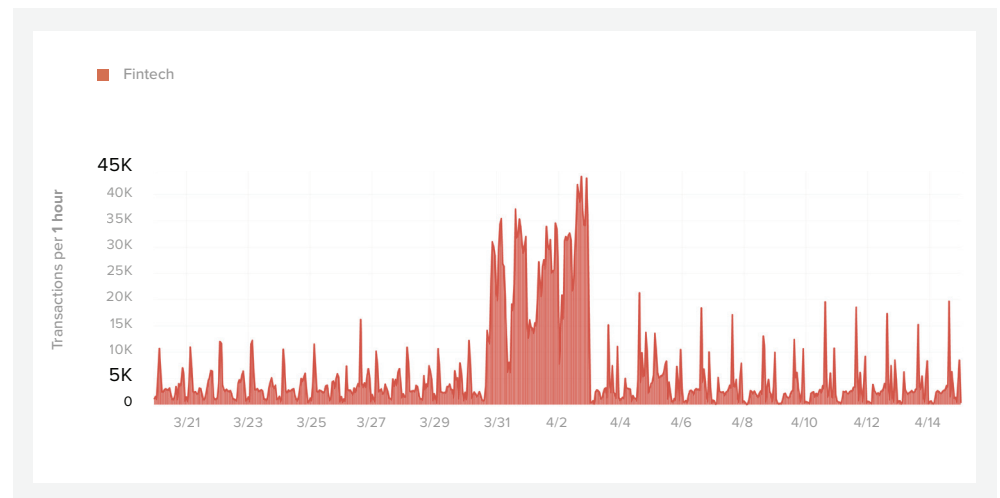


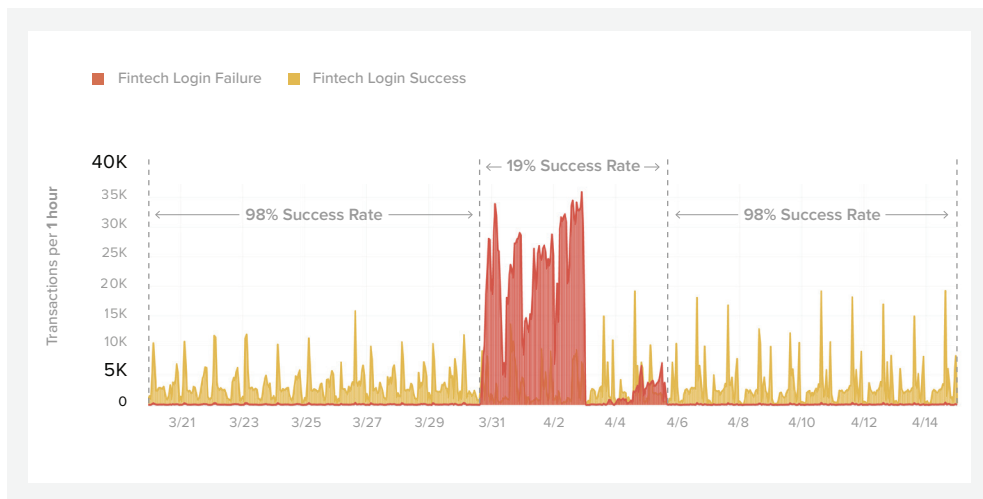
図3は、フィンテックのトラフィックのみを示していて、異常なスパイクがより鮮明になっています。

図3：米国のトップ3の銀行におけるフィンテックのトラフィック



最後に、図4は、フィンテックのトラフィックのログイン成功率を示しています。フィンテックは、プログラムでログインし、ユーザー名とパスワードを保存しているため、ログイン成功率は通常98%～99%です。しかし、フィンテックのトラフィックのスパイクがある期間は、ログイン成功率は19%まで低下しています。この率は、正規のフィンテックのトラフィックのログイン成功率98%と、クレデンシャルスタッフィング攻撃のログイン成功率0.1～3%を合わせた率に相当します。これがフィンテックによるクレデンシャルスタッフィング攻撃であることのさらなる証拠として、スパイク時に試行されたアカウントの71%（54K）がこれまでに試行されたことのないアカウントであったことが挙げられます。ログイン成功率が低いこと、新しいアカウントが試行されたこと、フィンテックが悪意のある自動化を阻止できないこと、そしてF5の顧客からのその後の立証はすべて、これが金融アグリゲータを介したクレデンシャルスタッフィング攻撃であったことの証拠です。

図 4：フィンテックのトラフィックのログイン成功率



F5 では、フィンテックがプログラムによって 1 日に何千回も顧客アカウントにログインし、スクレイピングしていることを確認しています。

フィンテックが顧客アカウントに安全にアクセスする方法の 1 つは、OAUTH 2.0 を使用することですが、残念ながら、多くのフィンテックはその採用を拒否しています。

重要なことです、攻撃者に狙われるログインアプリケーションは消費者ログインだけではありません。一般的に、攻撃者は、従業員、パートナ、サプライヤ、請負業者のログインも標的にします。このようなアカウントの乗っ取りは、請求書 / 支払い詐欺、知的財産の窃盗、情報セキュリティの侵害、またはその他の複雑なマルチステップスキームの進行につながるがよくあります。これらのスキームは、一般的に、標的となる企業のビジネスプロセスや業界特有の用語に精通し、十分な資源を持つ犯罪組織によって実行され、ソーシャルエンジニアリングに大きく依存し、重大な損失、あるいは桁外れの損失につながるが多々あります。

たとえば、米国のある犯罪組織は、却下された住宅ローン緩和申請書に記載された情報を利用して、精巧なソーシャルエンジニアリングスキームにより、申請者を標的としました。犯罪者は、被害者に電話をかけ、申請書の情報を利用して申請者を信頼させ、その信頼を悪用して、申請者に対し、住宅ローンの支払いを停止して代わりに別の住所に低額の「試験的な」住宅ローンを支払うよう指示しました。この犯罪者は、被害者が 3 回の試験的な支払いに成功すれば、住宅ローン緩和を恒久的に受けられると約束しました。これにより、数千人が数百万ドル騙し取られ、多くの人差し押さえによって家を失いました。個人識別情報 (PII) やローンなどの金融商品に関する「内部」情報の紛失は、巨額の詐欺、集団訴訟、批判につながる可能性があります。

上記のように、クレデンシャルスタッフィング攻撃での一般的なログイン成功率は 0.1% ～ 3.0% です。攻撃者は、これが異常であり、セキュリティ担当者に攻撃が気付かれる可能性があることを知っているため、ログイン成功率を上げる方法を探します。このための一般的な手法は、カナリアアカウントを使用すること、またはクレデンシャルスタッフィング攻撃の前にアカウント列挙攻撃を仕掛けることです。

カナリアアカウント

カナリアアカウントとは、攻撃者が確実に制御できるアカウント、つまり正しいユーザー名とパスワードを知っているアカウントのことです。攻撃中、サイバー犯罪者は、盗んだ、または購入した認証情報のリストのユーザー名とパスワードを試しますが、そのログイン成功率は 0.1% ～ 3.0% です。その後、同じインフラストラクチャを使用して、1 つまたは複数のカナリアアカウントへのログインを数回成功させます。このときのログイン成功率は 100% です。この結果、同じインフラストラクチャからのすべてのトランザクションの全体的なログイン成功率は上がります。

カナリアアカウントを使用する攻撃者には、さらなる利点があります。攻撃中にカナリアアカウントにログインしようとして、サーバーから「ログインに失敗しました。ユーザー名またはパスワードが正しくありません。」などのメッセージを受け取った場合、攻撃が軽減されていることを知ることができます。このようなフィードバックを受け取った攻撃者はすぐに、セキュリティ防御を回避するためにシステムを再構築する必要があるとわかります。この目標は、顧客アカウントへのログインを軽減しつつ、不正アクターに提供される情報を最小限にするためにカナリアアカウントへのログインを軽減しないようにすることです。

何千回もログインに成功しているにもかかわらず、正規の顧客活動がほとんど、またはまったくないアカウントは、カナリアアカウントである可能性があります。カナリアアカウントの作成に関連するメタデータから、攻撃者の正体が明らかになることもあります。アカウントの作成は通常、無害な活動であるため、不正アクターは、実際の攻撃時に使用するのと同じ運用上のセキュリティ (OpSec) を使用するとは限りません。

アカウント列挙攻撃

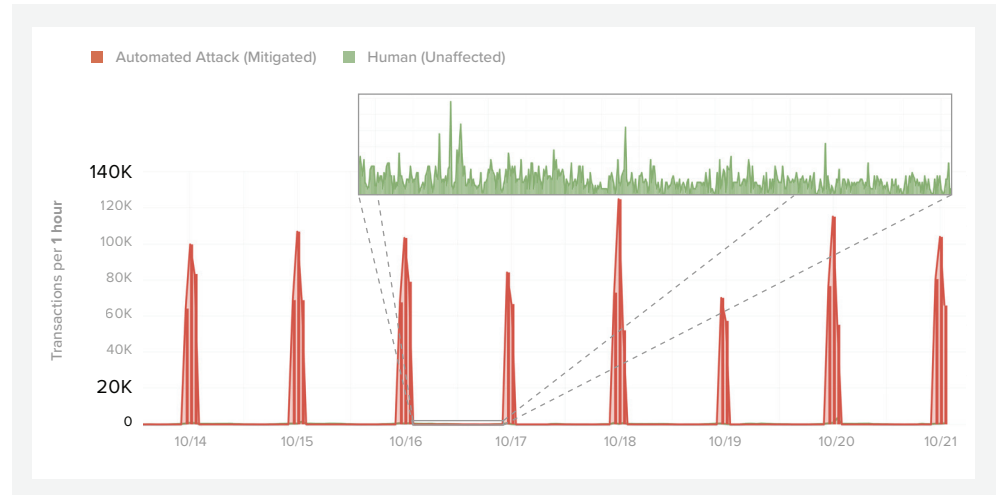
不正アクターがログイン成功率を上げるために利用するもう 1 つの手法は、まずアプリケーションに対して攻撃を仕掛け、有効なアカウントにユーザー名が対応しているかどうかを確認することです。有効なアカウントにユーザー名が対応していない場合、パスワードを試してもログインが失敗するため、攻撃者は、パスワードを試しません。アカウント列挙攻撃でよく狙われるアプリケーションは、「忘れてしまったユーザー名、ID またはパスワードの回復」および「アカウント作成」のためのアプリケーションです。

忘れてしまったユーザー名、ID またはパスワードの回復

これらのアプリケーションがアカウント列挙攻撃に対して脆弱かどうかは、ユーザーが存在しないアカウントのユーザー名、ID またはパスワードを回復しようとしたときに提供されるフィードバッ

クに依存します。「残念ですが、そのメールアドレスの記録はありません！」などのフィードバックが提供される場合、そのアプリケーションは脆弱であり、攻撃を受ける可能性があります。「メールアドレスの記録があれば、手順を記したメールを送ります。」などのより一般的なフィードバックが提供される場合、そのアプリケーションは脆弱ではありません。

図 5: Global 2000 F5 の顧客が忘れてしまったユーザー ID の回復アプリケーションへの攻撃



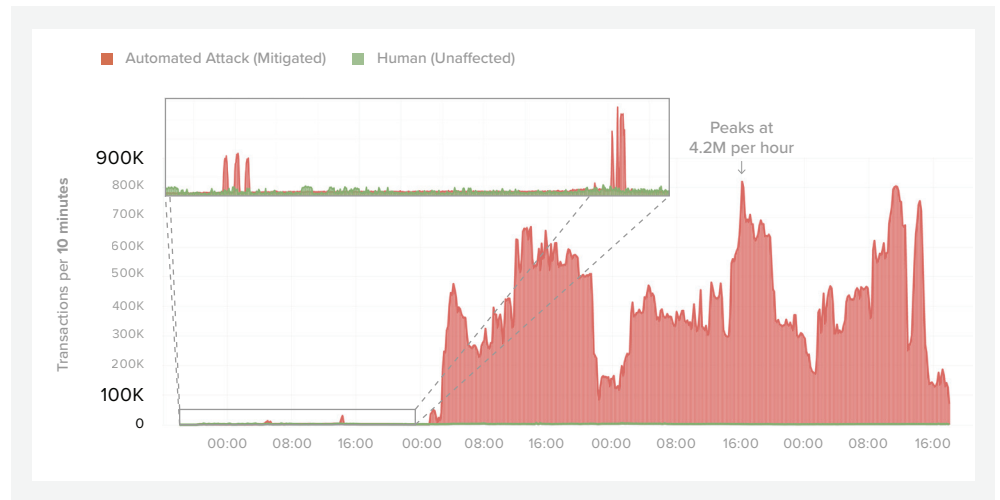
重要なことは、アカウントが存在するかどうかを確認できるフィードバックを提供しないことです。複数の電子メールアドレスを利用して、アカウント作成時に使用した電子メールアドレスを覚えていない顧客にとって、これは多少の摩擦となる可能性があります。しかし、自動化された攻撃を阻止する確実な手段がないので、攻撃者に貴重なフィードバックを提供しないように、フィードバックは一般的なものにしておくべきです。

「アカウント作成」または「忘れてしまったパスワードの回復」のためのアプリケーションに対する攻撃は、クレデンシャルスタッフィング攻撃の前兆となることがよくあります。

アカウントの作成

アカウント列挙攻撃のもう1つのタイプは、攻撃者が自動化を利用して、後続のクレデンシャルスタッフィング攻撃で使用するクレデンシャルリストにあるものと同じユーザー名を使用して、アカウントを作成することです。攻撃者がアカウントの作成に成功した場合、そのアカウントは以前には存在しなかったことを意味します。攻撃者は、特定のユーザー名とパスワードのペアがログインに確実に失敗することがわかったので、このペアをクレデンシャルスタッフィング攻撃から除外します。攻撃者が「残念ですが、そのアカウントはすでに存在します。」のようなフィードバックを受け取ると、攻撃者はそのアカウントが存在することを知り、関連するパスワードを後続のクレデンシャルスタッフィング攻撃で試す価値があることがわかります。

図 6：アカウント作成用のアプリケーションに対する攻撃。攻撃者は約 1 億 6 千万件の偽アカウントを作成しようとしていました。



攻撃者は、アカウント作成プロセスを完了させなくても、アカウントがすでに存在するかどうかを知ることができる場合があります。多くのアプリケーションでは、攻撃者が「ユーザー名の選択」フィールドに希望のユーザー名を入力するとすぐにフィードバックが送られます。攻撃者がアカウント作成を完了する場合は、いくつかのカナリアアカウントをその後の攻撃に利用できます。

アカウント作成アプリケーションに対するその他の自動化

F5 では、自動化がアカウント列挙以外の理由でアカウント作成アプリケーションに利用されていることを頻繁に確認しています。特に、マネーロンダリングや合成 ID の作成と維持に悪用される金融サービス組織で確認しています。

マネーロンダリング

金融機関は、アカウント作成に多くの情報を必要とします。たとえば、名前と電子メールアドレスだけしか必要としないオンライン小売業者よりもはるかに多くの情報を必要とします。しかし、これらの情報は完全に偽造できます。たとえば、地元の銀行でオンライン当座預金口座を開設する場合、銀行はフルネーム、生年月日、納税者 ID、住所、政府発行の ID、ナレッジベース認証（KBA）の質問に対する回答を求めることがあります。犯罪組織は、これらの質問に対する答えをマネーミュール（少額の手数料で口座にアクセスし、資金を移動させる人）から入手できます。マネーミュール自身が口座の開設や管理を依頼される場合もありますが、犯罪組織が何千もの口座を最大限に管理したい場合、マネーミュールからすべての必要な情報を集め、自動化を利用して、アカウント作成ワークフローのほとんどまたはすべてをナビゲートします。

合成 ID

合成 ID とは、オンライン上にのみ存在する人物のことです。人工知能で作られた顔と、名前、住所、電話番号、メールアドレス、クレジット履歴を持ち、銀行、小売店、ホテル、航空会社、電気通信会社などのオンラインアカウントを持っています。このような合成 ID は、悪意のある理由（バストアウト詐欺）にも、やや善意があると言える理由（影響力の拡大）にも利用されます。合成 ID を常にアクティブにして、より現実的にするために、オンラインアカウントの作成と操作に自動化機能が利用されることがよくあります。

バストアウト詐欺とは、攻撃者が何年にもわたって正規の顧客を装い、信頼を築き、クレジット限度額を上げ、すべての口座で最大限度額を借りてから、姿をくらますことです。

送金先の追加と送金に対する自動化

不正アクターが攻撃に自動化を利用する理由は 2 つあります。それは、繰り返しと操作です。繰り返しは、攻撃者が価値を実感するために同じ手順を何度も実行する必要がある場合です。たとえば、クレデンシャルスタッフィング攻撃では、攻撃者は何百万ものユーザー名とパスワードのペアを試す必要があり、これを手動でやることは現実的ではありません。

操作は、攻撃者が価値を実感するために 1 つまたはいくつかのステップだけを実行する必要がある場合です。しかし、標的となった Web またはモバイルアプリケーションは攻撃時には攻撃者はアクセスできません。操作目的での自動化の例は、Man-in-the-Browser (MitB) 攻撃です。

F5 は、顧客である米国の大手金融機関から、ある特定の MitB 攻撃について相談を受けたことがあります。それは、マルウェアに感染した Web ブラウザを使って銀行の顧客が口座にログインすると、ログインが乗っ取られ、被害者に気付かれないうちに、自動的に送金先が追加され、そこに送金され、その後、ブラウザの制御が被害者に戻されるということでした。

図 7: Man-in-the-Browser (MitB) 攻撃による送金先の追加、送金先への送金、2FA の無効化の説明

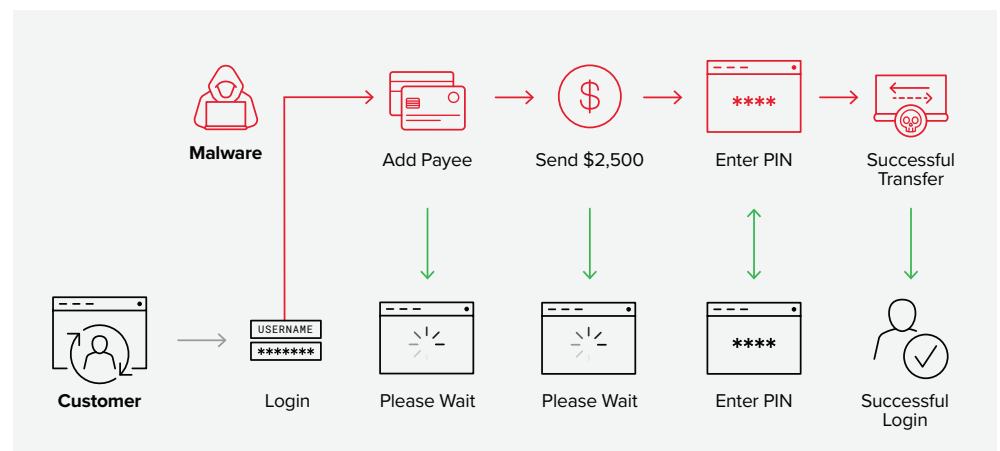
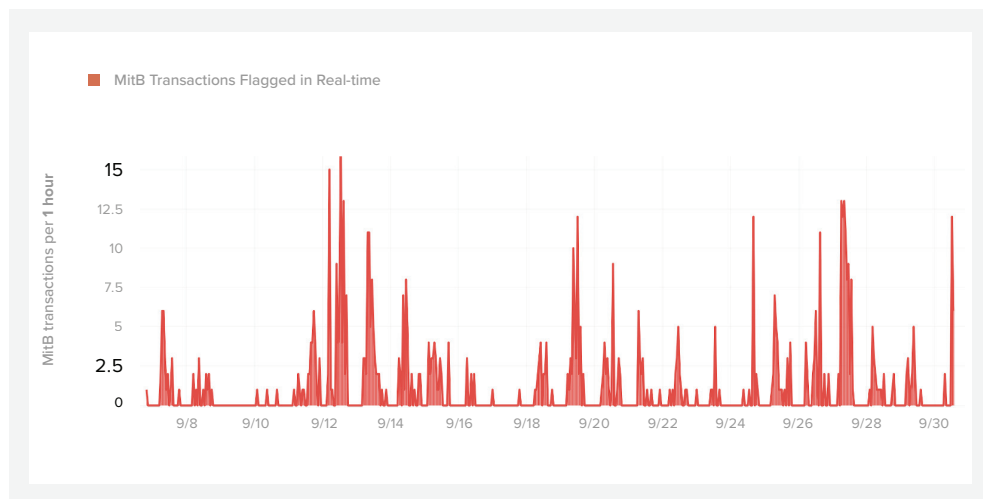


図 7 に示した銀行の顧客の実際の体験に注目してください。ユーザーは、ユーザー名とパスワードを入力し、「ログイン」をクリックします。数秒後、2FA コードがテキストメッセージとして送られ、同時に入力を求められます。顧客は、ログインプロセスの一部であると信じてコードを入力しますが、バックグラウンドでは、マルウェアがそのコードを使用して新しい送金先に送金しています。最終ステップでは、マルウェアは、HTML を動的に編集し、被害者に表示される口座残高を変更します。

図 8：米国トップ 3 の銀行に対する MitB 攻撃に関連する取引。F5 はこの攻撃をリアルタイムで特定しました。

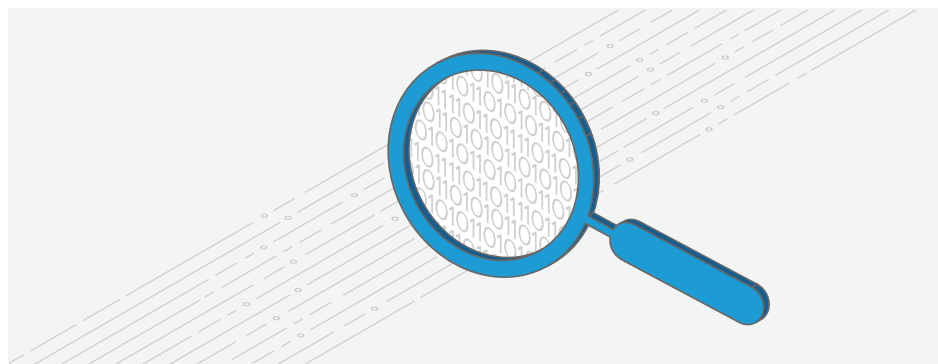


F5 は、自動化が利用される可能性のあるアプリケーションを特定し、それらのアプリケーションに対する自動化を定量化して、自信を持ってその善悪を分類できます。

組織を守る方法

内在的脆弱性に対する攻撃から組織を守るための特効薬はありません。この種の攻撃から組織を守るには、永久的な 3 ステップのプロセスを厳守する必要があります。

ステップ 1：可視化



組織を守るための最初の、そして最も重要なステップは、可視化です。そのためには、自動化が利用される可能性のあるアプリケーションを特定し、それらのアプリケーションに対する自動化を定量化して、自動化を十分理解した上で、自信を持ってその善悪を分類する必要があります。自動化は、適切な場所でなければ見つかることはできません。

自動化が利用される可能性のあるアプリケーションを特定するには、それぞれのアプリケーションに関するいくつかの質問に対して、十分な情報に基づいて客観的に答えを導き出す必要があります。たとえば、誰かが自動化を開始する理由は何か？自動化を行うことで金銭的なメリットや情報収集のメリットが得られる理由はあるか？自動化に対する短期的な動機がない場合、長期的な、より戦略的な動機があるか？これらの質問に対する答えを見つけるには、これらのアプリケーションとワークフローに精通しているだけでなく、他の金融機関に広がっている自動化についても包括的に理解している必要があります。

自動化を定量化することは、いくつかの理由から、多くの組織にとって重要な課題となっています。長年にわたり、組織は、IP アドレス、地域、またはアクターにとって変更することが些細なその他の属性によって、不要な自動化をブロックしてきました。

その後、IP アドレスでブロックされることにより、アクターはより分散するようになりました。現在、アクターは、1 つの国の 5 つの IP アドレスから 5 つのユーザーエージェントを使用して 1,000 万件のトランザクションを送信するのではなく、何百もの国の 100 万以上の IP アドレスから約 1,000 万のユーザーエージェントを使用して 1,000 万件のトランザクションを送信しています。攻撃者の進化と手口を変えるスピードは注目に値します。

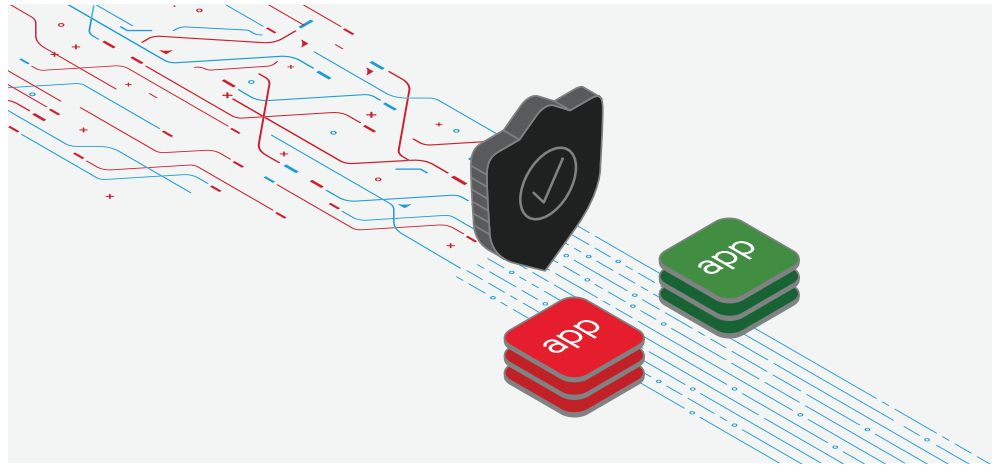
一般的に、組織は、自動化の発信元となる数百または数千の最も通信量の多い IP アドレスを特定していますが、レート制限をトリガーしないロングテール部分の IP アドレスは見逃しています。そして、ほとんどの自動化は、これらの IP アドレスから発信されています。

「当社のトラフィックの 20% から 30% は自動化されていると思います。」

—98% の攻撃トラフィックが自動化されていることが判明した F5 の顧客

自動化を十分理解した上で、自信を持ってその善悪を分類することも、多くの組織にとっての課題です。既知のテストツールやパフォーマンス監視ツールから承認されたトラフィックを特定することは、通常、難しいことはありません。クレデンシャルスタッフィングなど、明らかに悪意のある自動化を特定することも、通常、難しいことはありません。しかし、一部のアプリケーションに対する大量の自動化の意図が明確でないことがよくあります。従業員による無許可のテストなのか？フィンテックなのか？競合他社なのか？正当な理由で自動化を使用する忠実な顧客なのか？自動化を軽減する前に、自動化を理解することが重要です。

ステップ 2：自動化に対して行動を起こす



F5 の顧客である、ある米国政府機関は、50 か国以上にある IP アドレスから攻撃を受けていました。セキュリティチームは、米国外から発信されるトラフィックをすべてブロックすることを決定しました。しかし、その翌日、米国内の IP アドレスから攻撃が再開されました。

次のステップは、善意の自動化を許可リストに登録し、悪意の自動化を軽減することです。どちらにも、良い方法と悪い方法があります。

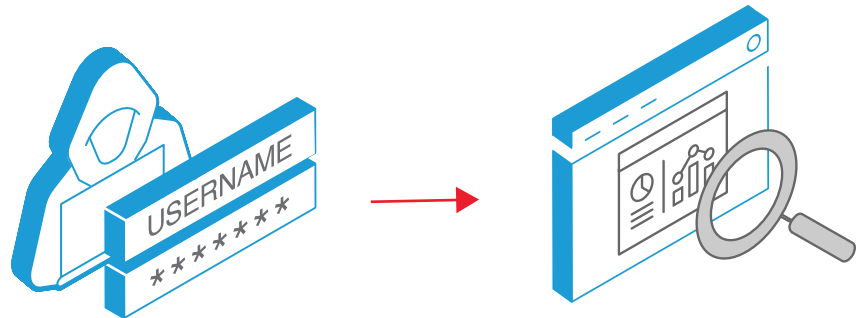
許可リスト

ユーザーエージェント文字列のような、簡単になりますことができる属性を使用して、トランザクションを許可リストに登録することは避けるべきです。理想的には、HTTP ヘッダーのどこかで共有シークレットを使用して、トランザクションを許可リストに登録する必要があります。

軽減

セッションを中断することで悪意の自動化を軽減することは避けるべきです。なぜなら、これにより、自動化が軽減されていることをすぐに知らせ、攻撃方法を変える動機を与えることになるからです。自動化が軽減されていることは、早く知られないようにする必要があります。F5 はこれを実現するために、トランザクションをリダイレクトまたは転送する機能、トランザクションの一部を追加または変更して発信元に渡す機能、あるいは完全に形成された HTML ページに応答する機能など、いくつかの追加の軽減オプションを提供しています。適切なアクションは、保護されるアプリケーションと、トランザクションがマルチステップスキームの現金化ステップに関連しているかどうかによって影響されます。

ステップ 3：継続的なレトロスペクティブ分析の実施



動機のある攻撃者は攻撃方法を変えます。組織は、標的となったアプリケーションに送られるトランザクションについて継続的なレトロスペクティブ分析を行い、攻撃方法の変化やその他の不要な自動化を迅速に特定する必要があります。これは、人間がサポートする人工知能および機械学習システムが集約トランザクションで動作しない限り、効果的に行うことはできません。

また、不要な自動化を見つけるだけでは十分ではなく、正規の顧客に影響を与えることなく、リアルタイムの防御を迅速に更新する手段も必要です。

従来の対策が効かない理由

従来の対策とは、Web アプリケーションファイアウォール（WAF）などのレイヤー 5 ～ 7 の防御、対象アプリケーションの二要素認証（2FA）の強制、または CAPTCHA などのボット検知および防止ツールなどのことです。

Web アプリケーションファイアウォール（WAF）

WAF は、ソフトウェアの脆弱性や弱点を狙った攻撃に対する強力な保護を提供します。たとえば、新しくリリースされたインジェクションに脆弱なシステムを悪用しようとしてインターネットをスキャンする自動化された攻撃から保護できます。しかし、内在的脆弱性に対する攻撃では、ビジネスロジックやアプリケーションを悪用するため、一般的には追加の防御が必要になります。

従来の WAF は、主にアプリケーション層を対象として OWASP Top Ten にあるリスクから防御します。アプリケーション層から提供されるシグナルは、高度な自動化を確実に検知するには不十分です。今日の自動化の検知には、行動バイオメトリクスの収集やブラウザ / デバイス環境の調査によって得られるような、追加のシグナルが必要です。

二要素認証（2FA）

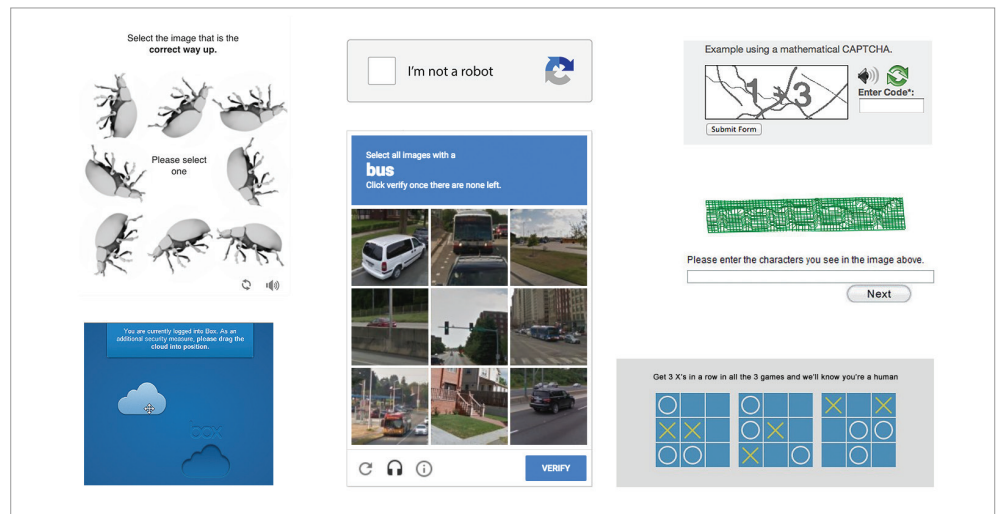
2FA は、セキュリティにおいて重要な役割を果たしますが、広範囲で使用するにはコストがかかり、顧客に大きな摩擦を与えます。2FA は、アカウント乗っ取り（ATO）をより難しくしますが、クレデンシャルスタッフィングを常に阻止できるわけではありません。

たとえば、ほとんどの 2FA 実装では、次のようなことが行われます。顧客が、ユーザー名とパスワードを送信し、それらが正しい場合、第 2 認証要素を入力するよう求められます。提出されたユーザー名とパスワードが正しくない場合、エラーメッセージを受け取り、第 2 認証要素を入力することは求められません。認証情報が正しい場合と正しくない場合のサーバーの応答が異なるため、攻撃者は認証情報が正しいかどうかを知ることができます。これにより、アカウントが攻撃者により乗っ取られるわけではありませんが、攻撃者は、ここで正しいことが判明した認証情報を、2FA を（たとえば、ポートアウト、SIM スワップ、SS7 攻撃、ソーシャルエンジニアリングなどにより）破ることを専門とする別の不正アクターに売ることができます。

CAPTCHA

CAPTCHA とは、Completely Automated Public Turing test to tell Computers and Humans Apart のバクロニウムです。これは、図 9 に示すように、さまざまな形式があります。

図 9：さまざまな形式の CAPTCHA



自動化のすべてが悪意あるわけではありません。また、悪意あるトラフィックや望ましくないトラフィックのすべてが自動化されているわけでもありません。組織には、この違いを見分けるソリューションが必要です。

CAPTCHA は、どの形式でも、望ましくない摩擦を忠実な顧客に与え、2FA のようにセッション離脱と収益離れにつながります。さらに悪いことに、CAPTCHA ではボットを阻止できません。光学式文字認識（OCR）、機械学習、さらには人間によるクリックファームを使用して少額の料金の CAPTCHA を解決する企業が多数存在します。たとえば、ロシアのクリックファームである 2CAPTCHA は、1,000 個の CAPTCHA の解決を 0.5 USD という低料金で提供しています。このビデオは、人間によるクリックファームで CAPTCHA を解決することがいかに簡単であることを示しています。

不要な自動化を軽減した後に期待されること

F5 は、Web およびモバイルアプリケーションに対する不要な自動化の軽減に成功した後も、手口を変える攻撃者と数週間から数か月に及び戦い続けた結果、多くの攻撃者が、人間によるクリックファームを使用して、手動でその活動を続けていることを確認しています。このため、F5 は、これらの脅威から顧客を保護するために複数の製品を提供しています。

F5 は、目に見えないように、あらゆるアプリケーションを攻撃、不正、悪用から保護します。F5 は、世界最大手の企業を守ることで、リクエストがボットからのものか、人間からのものかを区別するだけでなく、悪意のものか、善意のものかを区別する専門知識を身につけました。この能力により、F5 は、リアルタイムで不正行為を防ぎ、フィンテックやパートナーからのアプリケーションへのアクセスを管理し、顧客とビジネスのインサイトを強化する新しいデータを提供できます。

詳しくは、[Application Security](#) をご覧ください。

