

DER FAKTOR 2016 MENSCH

Heutzutage zielen Attacken eher auf menschliche als auf systemtechnische Schwächen ab. Proofpoint hat die Studie Der Faktor Mensch entwickelt, um diesen oft vernachlässigten Aspekt der Bedrohungen für Konzerne zu erforschen.

Dieses Papier gibt die ursprüngliche Feldforschung anhand von Daten wieder, die mithilfe von Proofpoint-Produkten aus Kundeneinrichtungen in aller Welt zusammengetragen wurden. Es behandelt die neuesten Trends bei den Top-Vektoren für die Auswahl von Zielpersonen: E-Mail, soziale Medien und mobile Apps. „Der Faktor Mensch“ zeigt nicht nur auf, wer worauf klickt, sondern auch, wie die Akteure von Bedrohungen mittels Social Engineering Leute dazu bringen, die Arbeit automatischer Exploits zu übernehmen. Denn wie aus den Daten hervorgeht, sind wir als Menschen das schwächste Glied in der Sicherheitskette.

KURZFASSUNG

2015 konnte man eine Lektion aus der Kunst lernen, als Cyberkriminelle aus der realen Welt jeden Tag das Mantra des Antiheld-Hackers der US-Fernsehserie „Mr. Robot“ anstimmten: „Menschen sind doch die besten Exploits.“ Social Engineering wurde zur Angriffstechnik Nr. 1. Die Angreifer wanderten zunehmend von den automatischen Exploits ab und setzten stattdessen Personen für die Drecksarbeit ein – Systeme infizieren, Anmeldedaten stehlen und Gelder abzweigen. Übergreifend für alle Vektoren und bei Attacken aller Größenordnungen, verwendeten die Bedrohungsakteure Social Engineering, um Leute dazu zu verleiten, das zu tun, wofür einst schädliche Codes benutzt wurden.

Hacker benutzen Menschen auf drei progressiv gesteuerte Arten

Stellvertretende Ausführung des Codes für den Hacker.

Diese Angriffe bestanden hauptsächlich aus umfangreichen Kampagnen, die über größere Benutzergruppen verbreitet wurden. Sie verwendeten unterschiedliche Kriegslisten, um der technischen Erkennung zu entkommen, und überzeugten ihre Opfer, Sicherheitsfunktionen zu deaktivieren oder zu ignorieren, auf Links zu klicken, Dokumente zu öffnen oder Dateien herunterzuladen, die Malware auf Laptops, Tablets und Smartphones installierten.

Übermittlung der Anmeldedaten an den Hacker.

Diese Attacken traten häufig in Kampagnen mittlerer Größe auf. Sie zielten auf zentrale Personen mit wertvollen Anmeldedaten, wie beispielsweise Benutzernamen und Passwörter für wichtige Systeme oder nützliche Dienste, und entlockten Ihnen auf hinterhältige Weise die „Schlüssel zum Schloss“.

Direkte Zuarbeit und Überweisung an den Hacker.

Diese Attacken waren rar gesät, dafür aber sehr gezielt. Ihre Opfer waren Benutzer mit den richtigen Aufgabenbereichen und der Fähigkeit, direkt im Auftrag des Hackers zu handeln. Die Benutzer gingen davon aus, dass sie Weisungen von höheren Instanzen befolgten, oft indem sie Überweisungen an betrügerische Bankkonten tätigten.

Die Attacken unterschieden sich in ihrer Größe und Menge. Jedoch teilten sie sich alle einen gemeinsamen Thread: Sie ersuchten Leute per Social Engineering, die Arbeit von Malware zu verrichten – und den Angreifern reiche Gewinne auszuschütten.

**2015 WAR SOCIAL
ENGINEERING DIE
ANGRIFFSTECHNIK
NR. 1. ANSTELLE
VON EXPLOITS
TRATEN PERSONEN
ALS BELIEBTESTES
MITTEL DER
ANGREIFER, DIE
CYBERSICHERHEIT
ZU SCHLAGEN.**

INHALTSVERZEICHNIS

Kurzfassung	2
Wichtigste Erkenntnisse und empfohlene Abwehrmaßnahmen.....	4
Abschnitt 1: Nach Zahlen: Auswahl der Zielpersonen nach Zeitpunkt und Ort des Klickens	8
Ausrichtung der Bedrohung nach geografischer Region	9
Ausrichtung von E-Mail-Bedrohungen nach Wochentag	10
Ausrichtung von E-Mail-Bedrohungen nach Tageszeit	11
Ausrichtung von Social-Media-Bedrohungen nach Uhrzeit	12
Ausrichtung der Bedrohung: über schädliche mobil-Apps	13
Abschnitt 2 : Ausnutzung des Faktor Mensch	15
Stellvertretende Ausführung des Codes für den Hacker	16
Trends bei dem Bedrohungsvektor E-Mail: URL- und Anhänge	17
Bedrohungsarten: Malware-Nutzlasten in Anhängen	18
Bedrohungsarten-Taktiken: Gängigste E-Mail-Köder	18
Bedrohungsarten: Schädliche Dokumentformate Im Anhang	19
Personen übermitteln Anmeldedaten an den Hacker	20
Bedrohungsvektor-Taktiken: Anmeldedaten-Phishing	21
Bedrohungen aus mobilen Apps werden erwachsen	22
Phishing überwiegt bei Angriffen in sozialen Medien	22
Leute überweisen Geld direkt an die Angreifer	24
Fazit	25
Wissen über fortschrittliche Bedrohungen	25
Empfehlungen	26

WICHTIGSTE ERKENNTNISSE UND EMPFOHLENE ABWEHRMAßNAHMEN

1. Personen statt Exploits als beliebteste Eindringtaktik der Angreifer

Hacker infizierten in überwältigendem Maße Computer, indem sie Leute zu Aktionen verleiteten, die sonst durch automatische Exploits ausgeführt wurden. Beeindruckende 99,7 % der bei anhangbasierten Kampagnen verwendeten Dokumente funktionierten über Social Engineering und Makros. Gleichzeitig stellen 98 % der URLs in infizierten Nachrichten eine Verknüpfung zu gehosteter Malware her, und zwar als ausführbare Datei, auch innerhalb eines Archivs. Damit der Trick klappt, müssen diese Dateien vom Benutzer geöffnet werden. So verleiten die Angreifer den Benutzer zum Doppelklicken, wodurch er den eigenen Rechner infiziert.

Empfehlung: Bei Ihrer Ausrichtung auf bestimmte Organisationen und Personen wandeln die Angreifer oft ihre Malware um und werden dadurch von der konventionellen Abwehr nicht erwischt. Beim Aufspüren dieser bislang unbekannten Bedrohungen ist eine dynamische Malware-Analyse besonders effektiv. Durch eine Überprüfung verdächtiger URLs und Anhänge mittels vielfältiger Techniken lassen sich die komplexen Bedrohungen unserer Zeit mit diesem Ansatz effektiver erkennen und blockieren. Um dem „Faktor Mensch“ bei den Attacken zu begegnen, sollten Sie die Benutzer durch regelmäßiges Training auf die neuesten Social-Engineering- und Anmeldedaten-Phishing-Schemen aufmerksam machen. Interessant wäre auch, wenn Sie eine Art von „Phishing“ an Ihren eigenen Mitarbeitern anwenden, um zu testen, wer von ihnen am ehesten klickt.

2. Kampagnen mit Banking-Trojaner Dridex – der Hauptangriffsvektor, der Personen in die Infektionskette einbezog

Banking-Trojaner waren die gängigste Art von Nutzlasten in schädlichen Dokumentanhängen – 74 % aller Nutzlasten gingen auf ihr Konto. Die Menge Dridex-basierter E-Mails lag fast 10 Mal höher als die am nächsthäufigsten verwendete Nutzlast solcher Attacken. Die Dokumentdateien in diesen Nachrichten enthielten schädliche Makros, die den Empfänger verleiteten, einen Code auszuführen, der seinen Rechner infizierte. Die Posteingänge von Mitarbeitern sind nach wie vor der bevorzugte Weg, einem Banking-Trojaner den Zugang zu einer Organisation wie der Ihren zu verschaffen. Die Angreifer verwenden Social Engineering und imitieren vertraute Prozesse wie Rechnungen und andere Anweisungen, die dem Benutzer einen Klick auf den Link in der E-Mail-Nachricht entlocken sollen. Social Engineering kann diesen Nachrichten sogar den Eindruck verleihen, dass sie von einem Kollegen oder Vorgesetzten stammen.

Empfehlung: Stellen Sie solide Sicherheitsmaßnahmen und -kontrollen bereit, die innerhalb Ihres E-Mail-Datenstroms operieren. Sie können auch Malware-Analysedienste in Echtzeit am E-Mail-Gateway laufen lassen, um Bedrohungen zu erkennen und zu blockieren, bevor sie Ihre Leute erreichen.

3. Hacker stimmten E-Mail- und Social-Media-Kampagnen zeitlich auf Phasen erhöhter Aktivität von Angestellten im Netz ab

Mit ihrem Wechsel vom Malware-Exploit zum Klick durch Menschenhand optimierten die Hacker die Lieferzeiten ihrer Kampagnen anhand der Zeiten, in denen besonders ausgiebig geklickt wird. E-Mail-Nachrichten werden in den Zielregionen zu Beginn des Arbeitstags (09.00–10.00 Uhr) zugestellt. Ähnlich spiegeln die Zustellzeiten für Social-Media-Spam die Spitzenzeiten offizieller Social-Media-Aktivitäten wider. Trotz allem gab es keine bestimmten Zeiten oder Wochentage, an denen keine schädlichen Inhalte bei Benutzern eingingen – oder diese darauf klickten.

Empfehlung: Da die in E-Mails und Social-Media-Beiträgen lauern Bedrohungen nicht von netzwerkbasierter Sicherheitslösungen erkannt werden, sollten Sie Lösungen in Betracht ziehen, die sich in diese Plattformen integrieren lassen. Eine Automatisierung kann Ihnen helfen, Ihre Sicherheitsprozesse zu skalieren, schnell auf Systemwarnungen zu reagieren und ernstere Bedrohungen von anderen zu unterscheiden. Suchen Sie nach einer Lösung, die in der Lage ist, automatisch festgestellte Bedrohungen zu blockieren, infizierte Benutzer unter Quarantäne zu stellen und Andere vor Infektionen zu schützen.

4. Benutzer luden bereitwillig über 2 Milliarden mobile Anwendungen herunter, die ihre persönlichen Daten entwenden

Hacker nutzten nicht nur E-Mails, sondern auch Bedrohungen über soziale Medien und mobile Apps, um Benutzer dazu zu bringen, ihre eigenen Systeme zu infizieren. Jeder fünfte Klick auf einen schädlichen URL geschah außerhalb des Netzes, oftmals über soziale Medien und Mobilgeräte. Bei diesen Klicks außer Haus handelt es sich längst nicht mehr um Sonderfälle, sondern um praxisnahe Bedrohungen. Unsere Analyse autorisierter App Stores für Android zeigte über 12.000 gefährliche Mobilgeräte auf — allesamt in der Lage, Daten zu stehlen, Hintertüren zu eröffnen und andere Funktionen auszuführen. Auf ihr Konto gingen mehr als 2 Milliarden Downloads.

Empfehlung: Zur Erkennung dieser Risiken und zur Vorbeugung gegen Attacken empfiehlt sich eine auf Mobilanwendungen zentrierte Threat-Intelligence- und Abwehrlösung. Neben schädlichen Apps sollte die Lösung auch Riskware—Apps behandeln. Diese sind zwar nicht immer offenkundig böse, verleiten jedoch zu riskantem Verhalten. Riskware ist für Mobilgeräte-Verwaltungsprogramme unsichtbar und sind deshalb auf unzähligen Mitarbeiter- und firmeneigenen Mobilgeräten zu finden. Diese Apps zeichnen sich durch eine breite Palette an bösen Verhaltensmustern aus. Mögliche Folgen sind Spionage vertraulicher Unternehmensdaten, Raub von Anmeldedaten oder Exfiltration von Daten – die oft später für gezielte Attacken auf Mitarbeiter benutzt werden. Suchen Sie nach einer Lösung, mit der Sie riskantes Verhalten so erkennen, beurteilen und kontrollieren können, dass die offiziellen Ansprüche der Benutzer an Produktivität und Privatsphäre nicht darunter leiden.

5. URLs, die Verknüpfungen zu Anmeldedaten-Phishing-Seiten herstellen, traten fast 3 Mal so oft auf wie Verknüpfungen zu Malware-Host-Seiten

Nach unseren Erkenntnissen leiteten im Schnitt 74 % der URLs aus Phishing-Kampagnen die Benutzer auf Anmeldedaten-Phishing-Seiten und nicht auf Malware-Hosts. Bei URL-basierten Kampagnen stellen Hacker Verknüpfungen zu Seiten her, die den Usern ihren Benutzernamen und andere persönliche Daten entlocken sollen. Das Opfer übernimmt somit die Arbeit von Keyloggern und anderen automatischen Malware-Programmen.

Empfehlung: Verschaffen Sie sich eine Lösung, die Tiefenanalysen, Inhaltsuntersuchungen und solide URL-Nachrichtendienste miteinander verbindet. Damit Sie diese Bedrohungen stoppen können, bevor sie Ihre Mitarbeiter erreichen, ziehen Sie einen agentlosen, cloudbasierten Dienst in Erwägung, der alle in E-Mails enthaltenen URLs neu schreibt und jeden URL testet, sobald ein Benutzer darauf klickt. Attacken richten sich gegen die arbeitende Bevölkerung. Vergewissern Sie sich, dass Ihre Abwehrmaßnahmen auch die Benutzer erfassen – ganz gleich, ob sie über das VPN der Firma, eine unsichere öffentliche Verbindung oder ihr eigenes Mobilgerät auf einen URL zugreifen.

6. Konten zur Freigabe von Dateien und Bildern – Google Drive, Adobe, Dropbox usw. – die effektivsten Köder für Anmeldedaten-Diebstahl

Links auf Google Drive erhielten als Köder für Anmeldedaten-Phishing die meisten Klicks. Diese Marken können den Benutzer zum Klicken verleiten, besonders dann, wenn das Opfer die Meldung von jemandem in seiner Kontaktliste erhält. Dies liegt daran, dass diese Dienste vertraut sind und der Benutzer daran gewöhnt ist, sich durch Klicken anzumelden, um gemeinsam genutzte Inhalte anzusehen.

Empfehlung: Kommen Sie den Bedrohungen zuvor und reagieren Sie schneller darauf – mit tieferen Einblicken und Intelligenz. Der Einblick in die Beschaffenheit der Kampagnen, die Ihre Firma zum Ziel haben, ist ein wichtiger Punkt, um den Zeitaufwand und die Bemühungen zur Unterbindung einer Attacke zu verringern und den Schaden zu begrenzen. Wenn ein Sicherheitsteam die Größe, Art und Dringlichkeit der Bedrohung kennt – und zudem einen detaillierten Überblick über die betroffenen Benutzer hat –, so kann es schneller handeln. Mit prädiktiven Abwehrmaßnahmen kommen Sie den Hackern zuvor. Durch proaktives Sandboxing verdächtiger URLs können URLs, die wahrscheinlich schädlich sind, auf präventive Weise identifiziert werden – bevor der Benutzer überhaupt die Chance erhält, zu klicken und damit seinen Rechner zu gefährden.

7. Phishing kommt in Social-Media-Beiträgen 10 Mal häufiger vor als Malware

Die am schnellsten wachsende Bedrohung in den sozialen Medien war das Phishing über gefälschte Kundenservice-Konten, die per Social Engineering Benutzer dazu verleiten, Benutzernamen und persönliche Daten preiszugeben. Die Leichtigkeit, mit der betrügerische Social-Media-Profilen bekannter Marken imitiert werden können, gibt dem Phishing deutlichen Vorrang bei Social-Media-basierten Attacken. Es ist nicht leicht, zwischen betrügerischen und echten Social-Media-Konten zu unterscheiden: Wir haben festgestellt, dass 40 % der Facebook-Konten und 20 % der Twitter-Konten, die angeblich eine Fortune 100-Marke repräsentieren, nicht autorisiert sind. Bei den Fortune 100-Unternehmen machen die nicht autorisierten Konten auf Facebook und Twitter jeweils 55 % bzw. 25 % der Konten aus.

Empfehlung: Lassen Sie es nicht zu, dass gefälschte Social-Media-Konten zum Kanal für Phishing und Social Engineering werden. Beginnen Sie mit einer Lösung, welche die mit Ihrer Marke verknüpften Social-Media-Konten erkennt, benachrichtigt und fortlaufend überwacht. Ihre Lösung sollte in der Lage sein, gefälschte Konten zu erkennen und zu analysieren, um Ihre Kunden und Ihren Ruf zu schützen.

8. Schädliche Mobilanwendungen von dubiosen Marketplace-Seiten schädigen 2 von 5 Unternehmen

Unsere Forscher haben unseriöse App Stores ermittelt, die den Benutzern erlaubten, schädliche Apps auf iOS-Geräte herunterzuladen – sogar solche, die nicht „geknackt“ waren oder deren Konfiguration Apps zuließ, die im iTunes Store von Apple nicht angeboten werden. Bei Benutzern, die – angelockt durch „kostenlose“ Klone beliebter Spiele und gesperrter Anwendungen – Apps von zweifelhaften Marketplaces herunterladen und im weiteren Verlauf gleich mehrere Sicherheitswarnungen übergehen, ist die Wahrscheinlichkeit, dass sie eine schädliche App herunterladen, viermal größer als bei anderen. Diese Anwendungen können persönliche Angaben, Passwörter und Daten entwenden. Etwa 40 % aller Großunternehmen, die von Proofpoint TAP Mobile Defense-Forschern stichprobenartig untersucht wurden, hatten schädliche Apps von DarkSideLoader-Marketplaces – d. h. von dubiosen App Stores – geladen.

Empfehlung: Stellen Sie eine Lösung bereit, die ständig riskante Verhaltensmuster von Apps auf Mobilgeräten erfasst und kontrolliert. Denn der Benutzer kann Apps von fast jeder Website herunterladen – selbst auf Geräten, die nicht „geknackt“ sind. Nur wenige (oder auch keine) dieser dubiosen App Stores kontrollieren ihre Inhalte. Selbst die geduldeten App Stores bringen eigene Sicherheits- und Compliance-Risiken ins Spiel. Sie benötigen beispielsweise keine Datenschutzrichtlinien für ihre Anwendungen. Das bedeutet, dass die von den Apps erfassten Daten völlig ungeschützt sind. Und die Stores überprüfen in der Regel jede App nur ein Mal. Sobald also eine App zugelassen ist, kann bei späteren Aktualisierungen ohne Weiteres ein riskantes oder böswilliges Verhalten an den Tag gelegt werden – was nach dem ersten Download monatelang unerkannt bleiben kann.

9. Kleinere Kampagnen mit sehr zielgerichteten Phishing-E-Mails konzentrierten sich innerhalb einer Organisation auf eine oder zwei Personen, um Gelder direkt an Hacker zu überweisen

Hoch gezielte Phishing-Nachrichten an Personen mit Zugriff auf Überweisungskonten treffen Organisationen jeder Größenordnung und Branche. Diese Masche, die oft als „Überweisungs-Phishing“ oder „CEO-Phishing“ bezeichnet wird, setzt Einiges an Hintergrundermittlungen durch die Angreifer voraus. Die E-Mails haben manipulierte Absenderadressen und scheinen daher vom CEO, CFO oder einer anderen Führungskraft zu stammen; selten sind Links oder Anhänge damit verbunden, und sie enthalten die dringende Anweisung an den Empfänger, Geld auf ein bestimmtes Konto zu überweisen.

Empfehlung: Um diese Bedrohung zu stoppen, bedarf es einer Kombination aus technischen Lösungen und Prozesssteuerungen. In technischer Hinsicht benötigen Sie ein E-Mail-Gateway, das erweiterte Konfigurationsoptionen zur Kennzeichnung verdächtiger Nachrichten anhand von Attributen (beispielsweise Adress- oder Betreffzeile) sowie E-Mail-Authentifizierungstechniken unterstützt. Bei Prozesssteuerungen müssen Sie sicherstellen, dass interne Kontrollen für Finanzen und Einkauf vorhanden sind, um legitime Anfragen zuzulassen. Dazu gehört u. A. eine sekundäre Genehmigung (persönlich oder telefonisch) durch einen anderen Angehörigen der Organisation mit Out-of-Band-Zugriff.

Allgemeine Empfehlungen

Verschaffen Sie sich gleichzeitig einen besseren Einblick in die Eigenart der Attacken gegen Ihre Organisation, damit Sie schnell zwischen gezielten und wahllosen Kampagnen unterscheiden können. Ihre Tools sollten Folgendes aufdecken können:

- Wer eine schädliche E-Mail erhalten hat
- Wie viele Versionen der gleichen Nachricht zugestellt wurden
- Wann sie empfangen wurden
- Welche Benutzer klickten
- Welche Benutzer das schädliche Ziel erreichten

Diese Informationen sollten Sie unbedingt zur Hand haben, um nach Priorität reagieren zu können.

ABSCHNITT 1

NACH
ZAHLEN**AUSWAHL DER ZIELPERSONEN NACH ZEITPUNKT
UND ORT DES KLICKENS**

2015 handelte es sich bei den meisten Attacken um groß angelegte E-Mail-Kampagnen. Großkampagnen zielen darauf ab, so viele Leute wie möglich zu treffen. Ein wesentliches Merkmal ist die massive Anpassung – ein Ausweichmanöver, bei dem ganze Regionen in Minutenschnelle mit Millionen von Nachrichten überschwemmt werden, die von Hunderttausenden manipulierten, offiziellen IP-Adressen gesendet werden. Es gibt keine Organisation, Branche oder Benutzerrolle, die nicht den fortschrittlichen Bedrohungen ausgesetzt war.

AUSRICHTUNG DER BEDROHUNG NACH GEOGRAFISCHER REGION

Die Großkampagnen von 2015 waren deutlich stärker nach Region als nach Organisation oder Einzelbenutzer ausgerichtet, wobei die Bedrohungsakteure ihre Ressourcen oft auf jeweils ein einziges Land konzentrierten.

- Bei den breit angelegten, stark regionszentrierten Kampagnen wurden lokalisierte E-Mail-Köder sowie schädliche Dokumentvorlagen und Nutzlasten verwendet. Diese Abrichtung der Kampagnen nach Zielregion erweckte bei den Empfängern einen realistischeren Eindruck und entlockte ihnen dadurch eher den einen oder anderen Klick, wie die Infektionsraten der Kampagnen zeigen.
- Botnets und Nutzlasten, die hauptsächlich in vereinzelt Ländern oder Regionen zu beobachten sind, können zwischen diesen wechseln. Solche Abweichungen traten gegen Ende 2015 häufiger auf; die Schwankungen bei den Ressourcen und Zielen gingen oft einher mit der Anpassung der E-Mail-Köder und Anhänge an die jeweilige neue Zielregion.

Das Interessante hierbei ist: Während die Großkampagnen von 2015 im Allgemeinen stark geografische Aspekte fixiert waren, verliefen sie innerhalb ihrer Zielregion weit weniger selektiv: Jede Einzelperson und Abteilung einer Organisation wurde zum Ziel der E-Mail-Kampagnen.¹

Verteilung der Top-Kampagnen nach Regionen

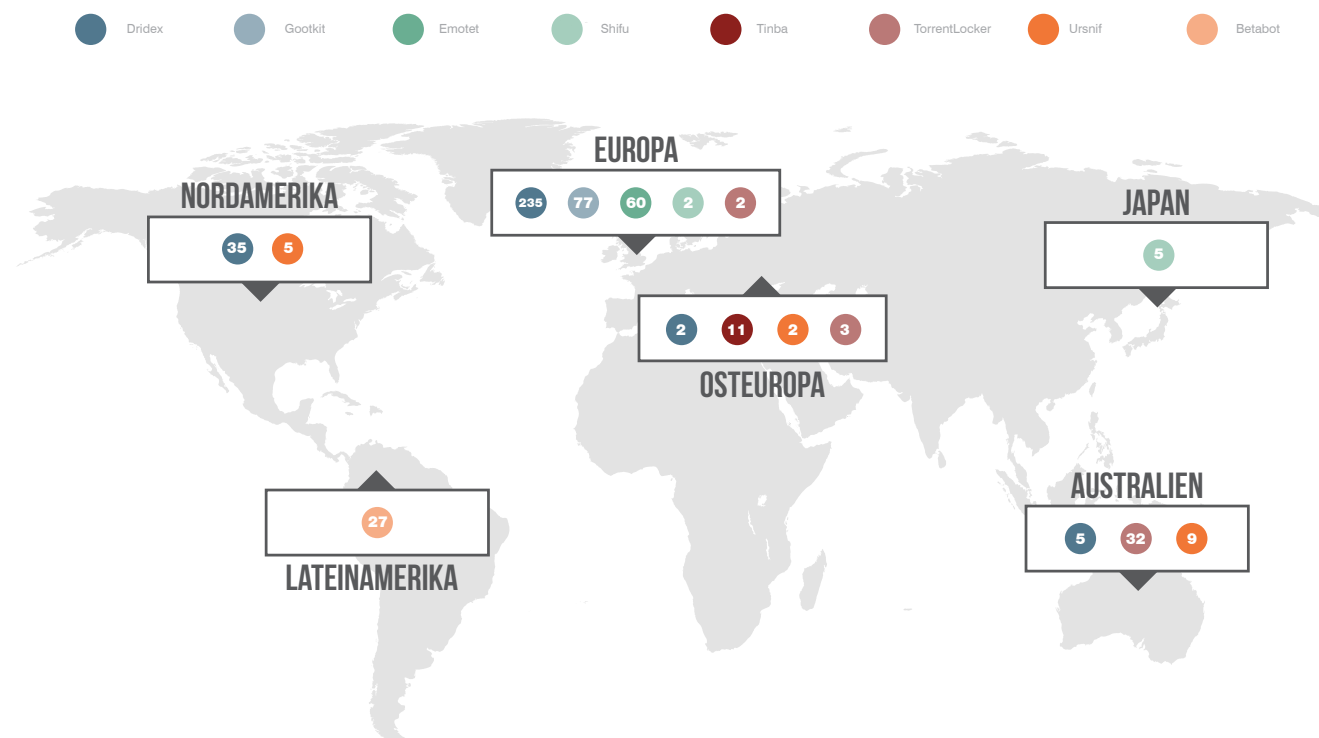


Abbildung 1 – Geographische Ausrichtung nach Malware-Nutzlast, 2015

1. Überraschenderweise ist die Verteilung der Nachrichten nach Abteilung nicht kategorisch, wie zu erwarten wäre, wenn die Angreifer explosionsartig Nachrichten an alle verfügbaren Adressen der Benutzer und Abteilungen einer Organisation richten würden. Der Grund ist vermutlich folgender: Je länger eine Person bei einer Firma arbeitet, desto wahrscheinlicher ist es, dass ihre Adresse auf einer Spammer-Liste mit gültigen E-Mail-Adressen dieser Firma steht. Abteilungen (bzw. Branchen) mit niedrigeren Fluktuationsraten beim Personal haben einen größeren prozentualen Anteil an E-Mail-Adressen von Mitarbeitern auf den Listen, die von den Spam-Services gekauft werden, mit denen diese breit angelegten Kampagnen arbeiten.

AUSRICHTUNG VON E-MAIL-BEDROHUNGEN NACH WOCHENTAG

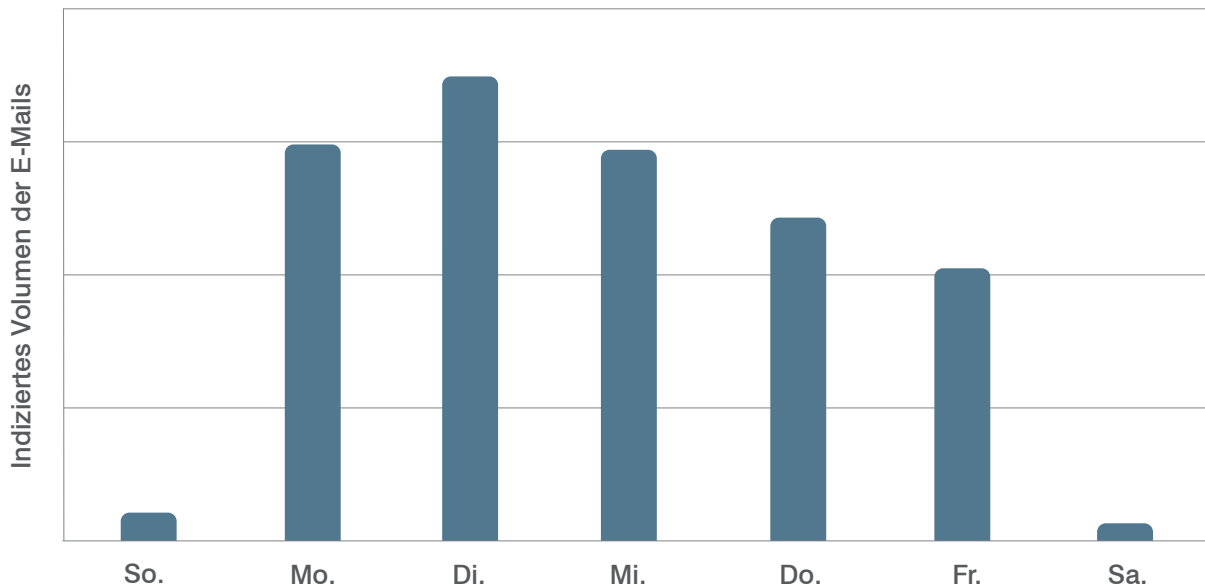


Abbildung 2: Indizierter Durchschnitt des Tagesvolumens gesendeter schädlicher Nachrichten nach Wochentag

- Wie schon 2014, führten die Bedrohungsakteure auch 2015 die massivsten Attacks an Dienstag durch, wenngleich der Unterschied verglichen mit anderen Wochentagen weniger ausgeprägt war.
- Ein deutlicher Schwerpunkt lag auf der ersten Wochenhälfte. Montag bis Mittwoch waren die bevorzugten Tage für Spam-Kampagnen.
- Die Zahl der Klicks nach Wochentag wies einen ähnlichen Trend auf: Montag und Dienstag standen abwechselnd an oberster Stelle, und das Volumen angeklickter URLs nahm im Laufe der Arbeitswoche allmählich ab.
- Waren die Volumen und Klickraten der Nachrichten an Wochenenden (samstags und sonntags) auch niedriger als werktags, so zeigten sie doch, dass Benutzer nach wie vor unerwünschte Nachrichten annehmen und darin auf schädliche URLs klicken. Das bedeutet, dass Organisationen in der Lage sein müssen, ihre Benutzer zu schützen, und zwar unabhängig davon, ob sie im Firmennetz oder auswärts auf einem Smartphone oder Tablet klicken.

Ende 2015 gab es verschiedene Riesenkampagnen, die montags starteten – möglicherweise, um aus den an Montagen höheren „Durchklickraten“ Kapital zu schlagen. Zwar handelt es sich hier auf den ersten Blick um vereinzelte Aktionen, doch die Kampagnen der ersten Monate von 2016 lassen erahnen, dass damit eine langfristige Änderung des Timings für Kampagnen mit unerwünschten E-Mails eingeläutet wurde.

AUSRICHTUNG VON E-MAIL-BEDROHUNGEN NACH TAGESZEIT

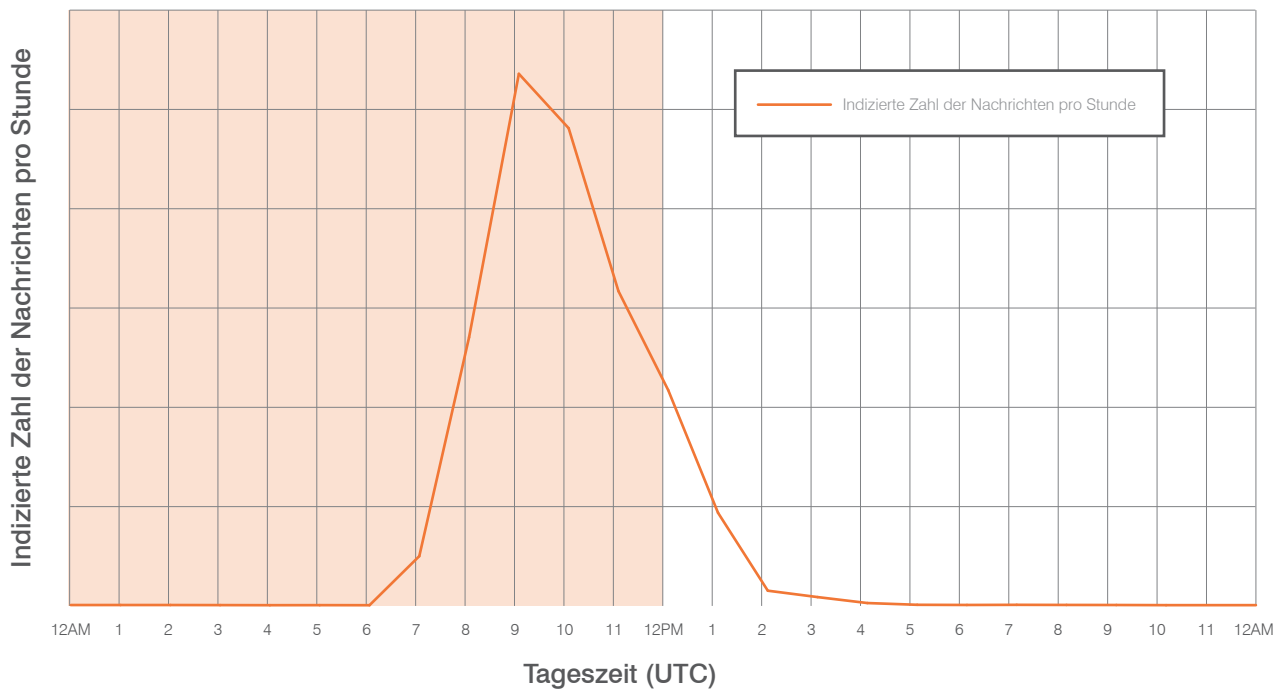


Abbildung 3: Beispiel für die Zustellung schädlicher Nachrichten nach Tageszeit, Dridex-Kampagnen (Großbritannien)

- E-Mail-Kampagnen wurden zeitlich so abgestimmt, dass sie zwischen 09.00 und 10.00 Uhr Ortszeit bei der Zielorganisation eintrafen. Das Ziel: Menschen bei ihrer Ankunft am Arbeitsplatz zu erwischen, bevor IT-Organisationen einschreiten können, um schädliche Nachrichten zu erkennen und abzufangen.
- Die am häufigsten verwendeten E-Mail-Lockmittel nutzten dieses Timing: Die mit Ködern präparierten Rechnungen, Empfangsbestätigungen und gescannten Dokumente waren darauf ausgelegt, gleich zu Beginn des Arbeitstags die Aufmerksamkeit ihrer Adressaten – d. h. der Büroangestellten – zu erwecken.
- Die Verbreitung von E-Mails über den Tag verlief insgesamt nach einem ähnlichen Muster wie 2014 – mit Spitzen, die dem Tagesbeginn in der anvisierten Zeitzone entsprechen.
- Je länger eine E-Mail im Posteingang eines Benutzers verbleibt, desto unwahrscheinlicher ist es, dass dieser darauf klickt. Und sowohl im weiteren Tagesverlauf als auch über die Woche lassen die Klickraten nach. Angesichts dessen richteten sich die Hacker bei der Zustellung sowohl nach dem Wochentag als auch nach der Tageszeit.

AUSRICHTUNG VON SOCIAL-MEDIA-BEDROHUNGEN NACH UHRZEIT

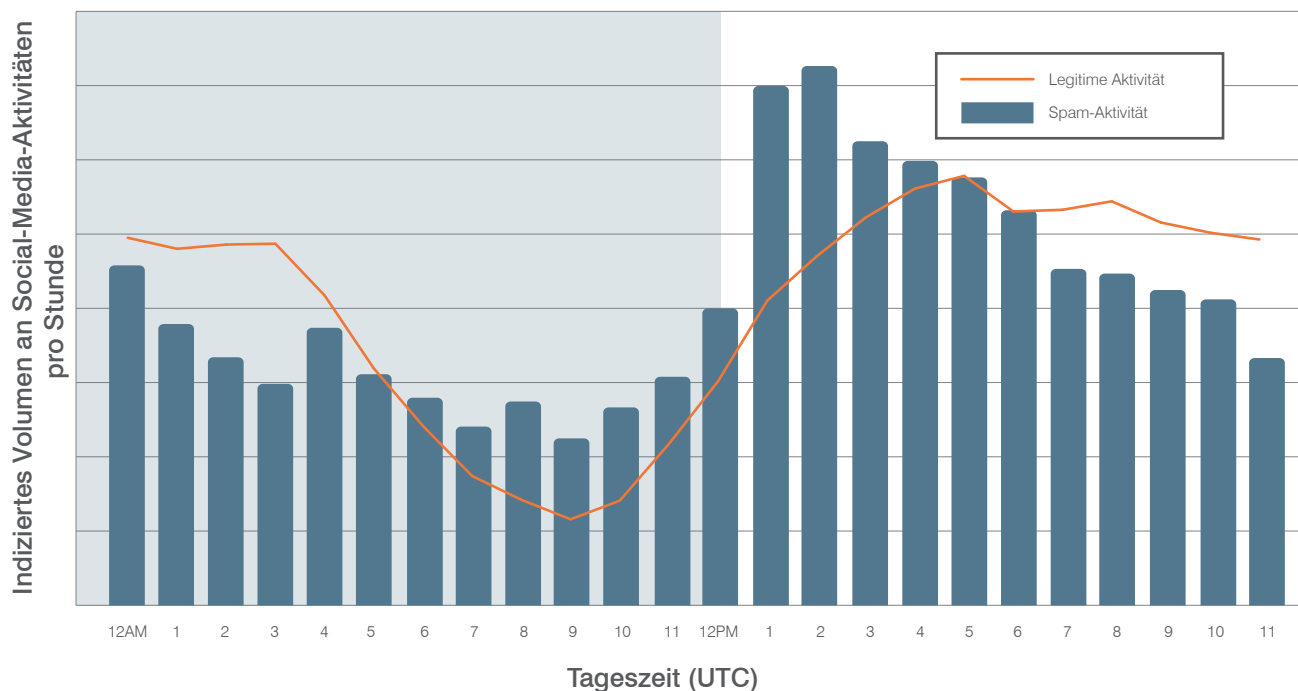


Abbildung 4: Social-Media-Aktivität (offiziell und Spam) nach Tageszeit

- Angreifer optimieren ihre Aktivität anhand der Hauptbeschäftigungszeiten am Zielort. Bei den anvisierten Organisationen werden sowohl die Zielgruppe als auch die Zeitzone in Betracht gezogen.
- Die Spitzenzeiten für Social-Media-Spam beginnen gegen 08.00 Uhr (ET) und dauern ca. fünf Stunden an, d. h. bis 14.00 Uhr (ET). Insgesamt lassen die Aktivitäten in den USA im Laufe eines Werktages nach. Diese Aktivitätsspitzen umfassen den Großteil der Zeit von morgens bis mittags, wenn die Leute auf Facebook und anderen sozialen Medien besonders aktiv sind.
- Die Aktivitäten mit Social-Media-Spam hören niemals gänzlich auf. Tagein, tagaus – ständig postet irgendein Spammer betrügerische und potenziell schädliche Inhalte.

**21,5 % DER KLICKS
FANDEN AUSSERHALB
DES NETZWERKS
STATT**



ZUSAMMENFASSUNG

Die Akteure fortschrittlicher Bedrohungen suchen sich ihre Opfer gezielt nach deren schwachen Momenten zwischen Spitzen im Datenverkehr und Ermüdung aus. Social-Media-Bedrohungen und mobile Apps haben sich in die E-Mail-Flut eingereiht, um Benutzer dazu zu bringen, ihre eigenen Systeme zu infizieren. Jeder Fünfte klickt auf schädliche URLs außerhalb des Netzwerks. Und in diesem Jahr wurde die Bedrohung durch schädliche E-Mails außerhalb des Netzwerks durch die Zunahme der Bedrohungen für Social Media und mobile Apps ergänzt.

AUSRICHTUNG DER BEDROHUNG : ÜBER SCHÄDLICHE MOBIL-APPS

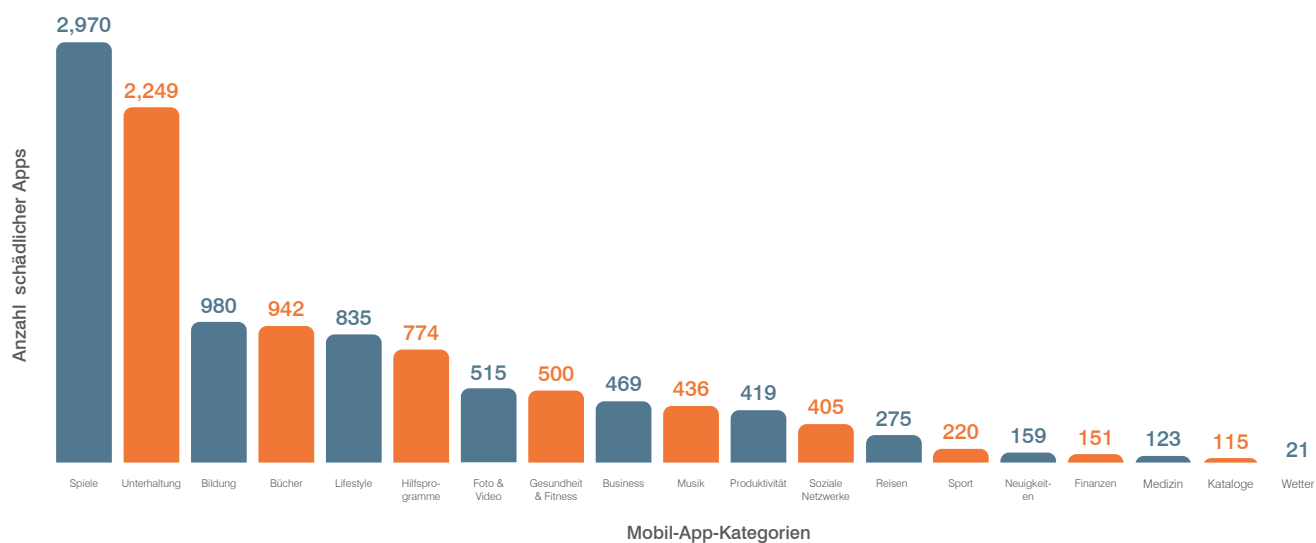


Abbildung 5: Schädliche Mobil-Apps nach Kategorie

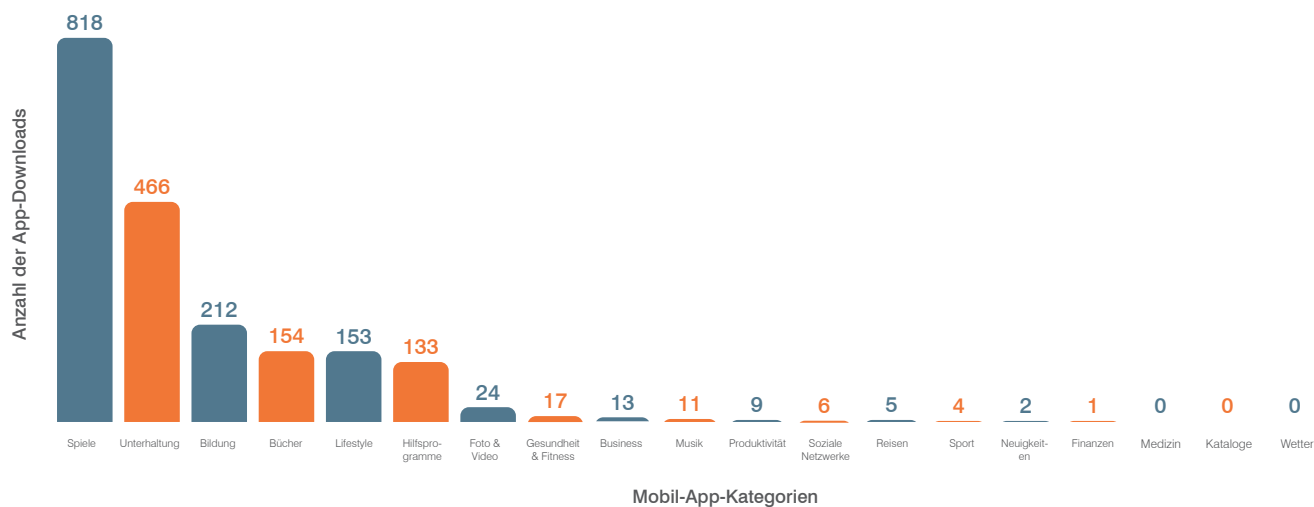


Abbildung 6: Schädliche Mobil-App-Downloads nach Kategorie

- Allein für Android wurden über 2 Milliarden Mal schädliche Mobil-Apps heruntergeladen.
- Schädliche Apps sind stark auf die App-Kategorien Spiele und Unterhaltung ausgerichtet.
- Viele dieser schädlichen Apps, von Games bis zu Produktivitätstools, werden als kostenlose Version oder Raubkopie beliebter, kostenpflichtiger Apps angepriesen. Mit diesem Lockmittel werden Benutzer verleitet, bereitwillig Malware auf ihre Smartphones und Tablets zu laden.

Die schädlichen Apps sind für Hacker ein attraktiver Vektor. Im Gegensatz zu E-Mail-basierten Kampagnen, bei denen Spam-Nachrichten an Millionen von Benutzern gesendet werden, kann eine App in einem einzigen Store Millionen potenzieller Benutzer erreichen.

ZUSAMMENFASSUNG

Schädliche mobil-app aus dubiosen marketplaces schädigen 2 von 5 unternehmen.

Unsere Forscher haben unseriöse App Stores ermittelt, von denen Benutzer schädliche Apps auf iOS-Geräte herunterladen könnten. Die dubiosen App Stores funktionierten sogar bei Geräten, deren Konfiguration Apps zuließ, die im iTunes Store von Apple nicht angeboten werden (sogenannte „geknackte“ Geräte). Bei Benutzern, die – angelockt durch kostenlose Klone beliebter Spiele und gesperrter Anwendungen – Apps von zweifelhaften Marketplaces herunterladen und im weiteren Verlauf gleich mehrere Sicherheitswarnungen übergehen, ist die Wahrscheinlichkeit, dass sie eine schädliche App herunterladen, viermal größer als bei anderen. Diese Anwendungen stehlen persönliche Angaben, Passwörter und Daten. Etwa 40 % aller Großunternehmen, die von Proofpoint TAP Mobile Defense-Forschern stichprobenartig untersucht wurden, hatten schädliche Apps von DarkSideLoader-Marketplaces – d. h. von dubiosen App Stores – geladen.

**ALLEIN FÜR ANDROID
WURDEN ÜBER
2 MILLIARDEN
MAL SCHÄDLICHE
MOBIL-APPS
HERUNTERGELADEN.**



ABSCHNITT 2

ABSCHNITT 2

AUSNUTZUNG DES FAKTORS MENSCH

Die Befunde in früheren Ausgaben des Berichts „Faktor Mensch“ machten deutlich, dass „in jeder Organisation geklickt wird“. Ungeachtet der Größe, des Standorts und der Branche lässt sich sagen, dass die Menge der Klicks auf schädliche URLs niemals gleich Null war. In den vergangenen Jahren wurden die Systeme der Anwender durch einen Klick auf einen schädlichen URL in einer unerwünschten E-Mail in eine hochkomplizierte, kriminelle Cyber-Infrastruktur. Diese Kampagnen verwendeten automatische Exploits, um das System des Benutzers mit einer Malware-Nutzlast zu infizieren oder dessen Daten zu stehlen.

Doch die Dinge ändern sich. 2015 halsten die Angreifer den Benutzern – d. h. Menschen – die Arbeit automatischer Exploits auf. In diesem Abschnitt zeigen wir, wie sich diese Verschiebung auf alle Bereiche auswirkt – von der Art der eingesetzten Malware-Bedrohungsakteure bis zu den E-Mails, Ködern und Social-Media-Beiträgen, aus denen sich ihre Kampagnen zusammensetzten.

STELLVERTRETENDE AUSFÜHRUNG DES CODES FÜR DEN HACKER

2015 infizierten Hacker in überwältigendem Maße Computer, indem sie Leute zu Aktionen verleiteten, die sonst durch automatische Exploits ausgeführt wurden.

- **99,7% der bei anhangbasierten Kampagnen verwendeten Dokumente funktionierten über Social Engineering und Makros, nicht über automatische Exploits.** (Die Beilage „(Schädliche) Inhalte aktivieren“ beschreibt ausführlich, wie diese Kampagnen funktionieren.)
- **98 % der URLs in infizierten Nachrichten stellen eine Verknüpfung zu gehosteter Malware her, und zwar als ausführbare Datei, auch innerhalb eines Archivs.** Exploit Kits verwenden schädliche Codes, um den Rechner des Benutzers automatisch zu infizieren; bei gehosteten Archiven und ausführbaren Dateien mit schädlichem Inhalt muss der Benutzer dazu verleitet werden, das eigene System zu infizieren, indem er auf den Malware-Link doppelklickt.

Was die Sache noch verschlimmerte, war das dramatische Ausmaß der Kampagnen mit dem Banking-Trojaner Dridex, bei der die Menschen im Mittelpunkt der Infektionskette standen. 2015 waren Banking-Trojaner die gängigste Nutzlast bei schädlichen Dokumentanhängen (Abb. 7). 74 % aller Nutzlasten gingen auf ihr Konto, und die Menge der Dridex-Nachrichten lag fast 10 Mal höher als die am nächsthäufigsten verwendete Nutzlast von Attacken mit schädlichen Dokumentformaten im Anhang. Die Dokumente selbst funktionierten überwiegend mit schädlichen Makros, verbunden mit Social Engineering, um den Benutzer dazu zu verleiten, einen Code auszuführen, der seinen Rechner infizierte.

**98 % DER URLS
UND 99,7 % DER
DOKUMENTBASIERTEN
ANGRIFFE HINGEN
VON LEUTEN AB, DIE
BEWUSST KLICKTEN, UM
SICHERHEITSFUNKTIONEN
ZU DEAKTIVIEREN. JEDE
ZEHNTE PERSON TAT DIES
AUCH.**

(SCHÄDLICHE) INHALTE AKTIVIEREN

Schädliche Makros auf Microsoft Office-Basis, auch bekannt als VBA-Viren (Visual Basic for Applications), sind Codeauszüge, die in ein Office-Dokument – beispielsweise in Word oder Excel – eingebettet werden können. Beim Öffnen des Dokuments können diese Makros vielfältige Operationen durchführen, wie beispielsweise den automatischen Aufruf des Downloaders für ein Malware-Element. VBA-Viren waren Ende der 1990er Jahre die häufigste Angriffsmethode.

Jedoch ließ dieser Trend vor fast einem Jahrzehnt rasch wieder nach, als Office 2007 damit begann, Makros standardmäßig zu deaktivieren. Dies ist nach wie vor die Standardeinstellung für Microsoft Office. Wenn ein Benutzer an Office-Dokument mit Makros öffnet, wird ihm eine Warnung angezeigt; demnach macht die Aktivierung von Makros „Ihren Computer durch potenziell schädlichen Code angreifbar und wird nicht empfohlen“. (Siehe Abbildung z26.)

Trotz dieser Schutzmaßnahme wurden schädliche Makros Ende 2014 und Anfang 2015 wieder zum Leben erweckt. Sie verbreiten sich durch Phishing-Kampagnen, gestützt durch verschiedenste Social-Engineering-Techniken, die den Empfänger austricksen, damit er die Makros selber aktiviert. Nach ihrer Aktivierung installieren die Makros vielfältige Malware-Nutzlasten auf den Rechnern ihrer Opfer. Da diese Technik voraussetzt, dass das Opfer die Makros direkt aktiviert – meistens durch Klicken auf die Schaltfläche „Inhalte aktivieren“ – ist das Social Engineering ein wesentlicher Punkt dieser Kampagnen. Dokumentanhänge mit schädlichen Makros stehen meistens in Verbindung mit Hackern, die den Banking-Trojaner Dridex verwenden. Doch diese Technik wird weitläufig dazu genutzt, eine Palette weiterer Malware in Umlauf zu bringen – darunter auch solche, die von Angreifern mit komplexem Schadcode eingesetzt wird. Der große Vorteil dieses Ansatzes liegt darin, dass der Vektor nicht mit Patches zu beheben ist – stattdessen wird die Klickbereitschaft des Nutzers zu einem zentralen Bestandteil der Infektionskette.

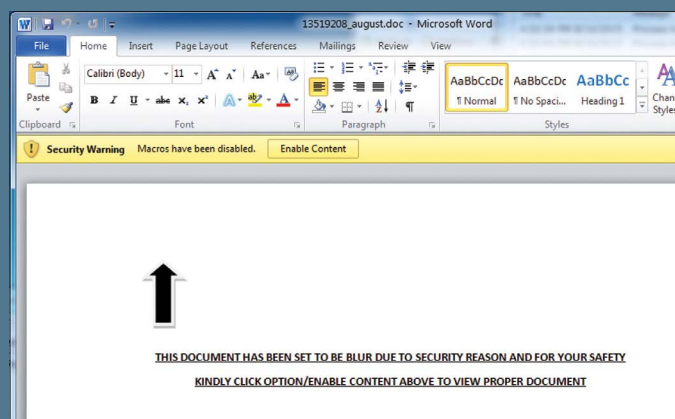


Abbildung 7: Beispiel eines Microsoft Word-Anhangs mit der Aufforderung zum Aktivieren eines (schädlichen) Makros

TRENDS BEI DEM BEDROHUNGSVEKTOR E-MAIL: URL- UND ANHÄNGE

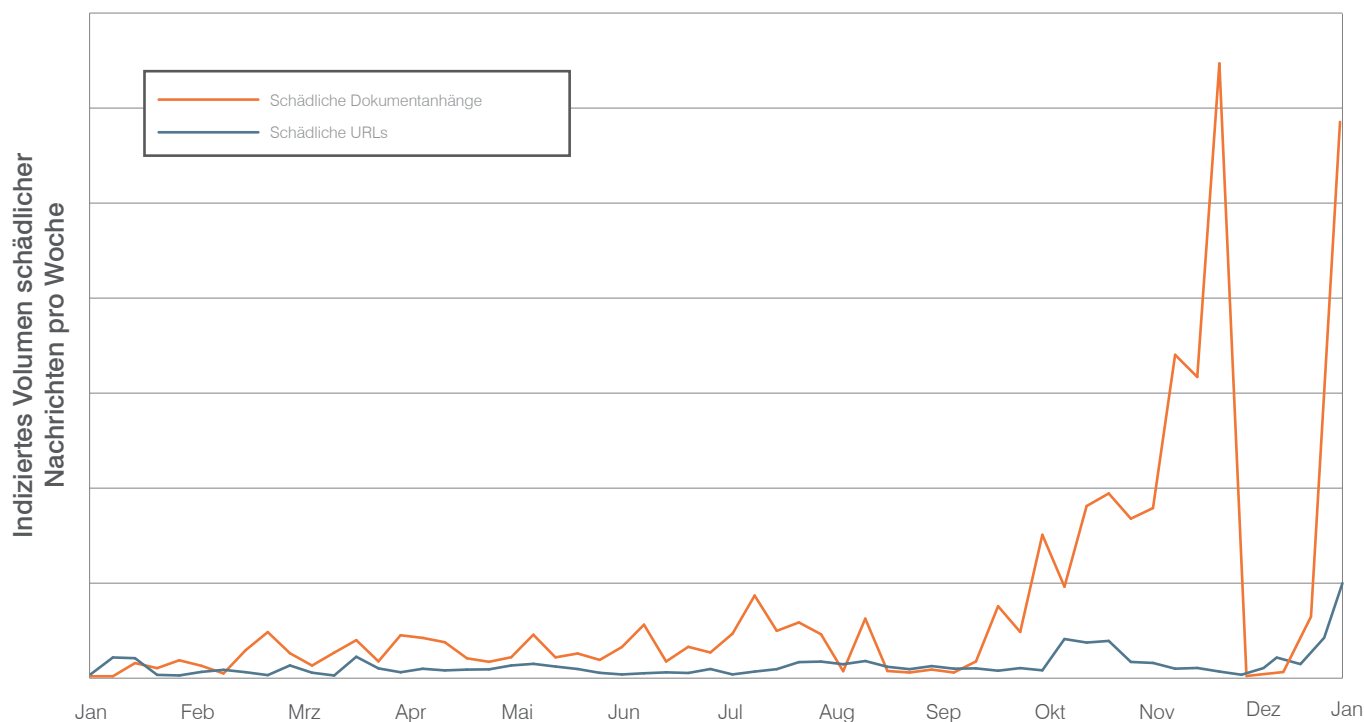


Abbildung 8: Indizierter Trend gefährlicher URLs und Dokumentanhänge, 2015

- Bedrohungsakteure bevorzugten 2015 weitgehend schädliche Dokumentanhänge gegenüber URLs für E-Mail-basierte Angriffe.
- Statt wieder auf URL-basierte Kampagnen zurückzugreifen, vergrößerten die Hacker Ende 2015, als die Abwehrmechanismen angepasst wurden, den Umfang ihrer Kampagnen drastisch und unternahmen regelrecht aggressive Attacken gegen Organisationen in Großbritannien und Europa.
- URL-basierte Kampagnen der Art, wie wir sie in den Vorjahren beobachtet haben, waren im Vergleich zu Dokumentanhang-basierten Kampagnen praktisch nicht vorhanden.
- Im September verhielten sich die Dridex-Bedrohungsakteure ruhig, nachdem Schlüsselpersonen eines kriminellen Rings verhaftet worden waren, dem der Diebstahl mehrerer Millionen Dollar von seinen Opfern mithilfe der Schadsoftware zur Last gelegt wurde. Die damit verbundenen Festnahmen beeinträchtigten ihre Sendeinfrastruktur und trugen möglicherweise auch zum hohen Aufkommen an Kampagnen im November und Dezember bei.

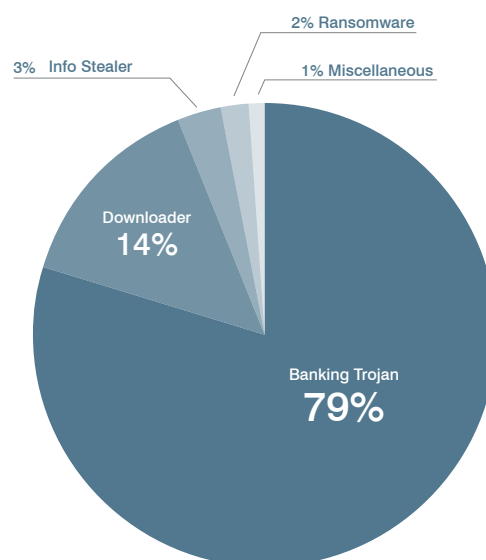


Abbildung 9: Malware-Nutzlasten am häufigsten über schädliche Dokumentanhänge zugestellt

BEDROHUNGSARTEN: MALWARE-NUTZLASTEN IN ANHÄNGEN

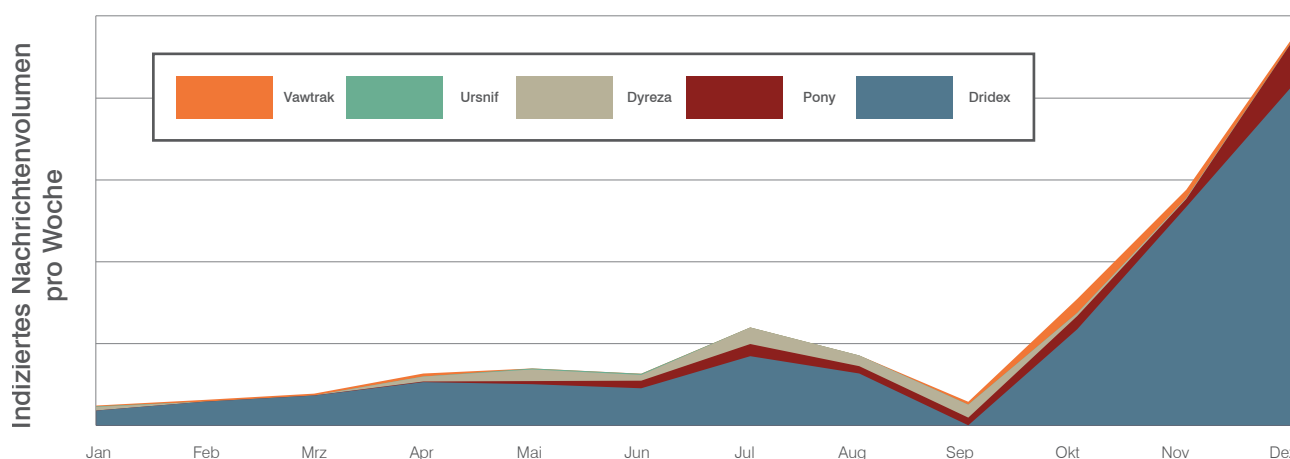


Abbildung 10: Wöchentliche Malware-Nutzlast aus Phishing-Dokumentanhängen als Prozentsatz der gesamten Malware für die Top 5 der Nutzlasten 2015

Bedrohungsakteure stimmen ihre Nutzlasten oft auf die Verbreitungstechniken ab. Ein Beispiel: Banking-Trojaner waren zwar die häufigste Nutzlast bei Kampagnen mit schädlichen Dokumentanhängen, dafür wurde aber eine Verbreitung dieser Schädlinge durch Exploit Kits relativ selten beobachtet. Andererseits gab es zwar selten Fälle von Erpressungssoftware (Ransomware) in Kampagnen mit schädlichen Dokumentanhängen, aber dafür umso öfter bei Kampagnen mit Exploit Kits und Archivanhängen. Jedoch nahmen diese Spezialisierungen ab Ende 2015 zu. Wir haben beobachtet, wie Angreifer anfangen, Nutzlasten und Verbreitungstechniken zu nutzen, die vorher nicht miteinander in Verbindung gebracht wurden.² Ende 2015 änderten sich die Konfigurationen der Banking-Trojaner: Benutzer-Anmeldedaten wurden in vielen bis dahin unüblichen Dienstleistungsbereichen anvisiert, und plötzlich waren alle Benutzerkonten – von Versand und Großhandel bis zum Cloud Service – gefährdet.³

BEDROHUNGSARTEN-TAKTIKEN: GÄNGIGSTE E-MAIL-KÖDER

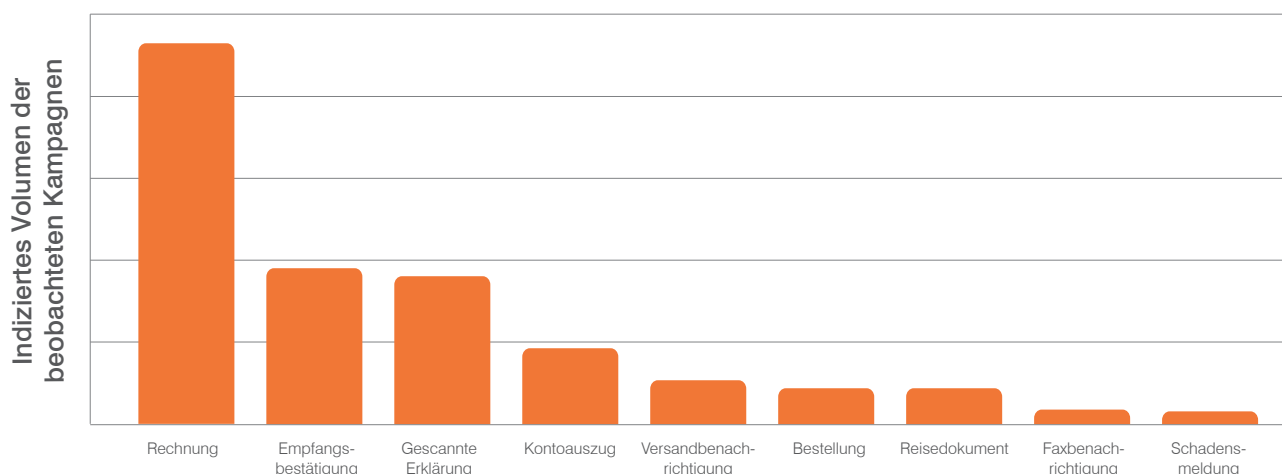


Abbildung 11: Gängigste E-Mail-Köder in Kampagnen mit Dokumentanhängen

² „Dridex and Shifu give spam bots the day off and spread via exploit kits“, 18. November 2015, <http://www.proofpoint.com/us/threat-insight/post/dridex-shifu-give-spam-bots-day-off>

³ „Dyre malware campaigners innovate with distribution techniques“, 9. Oktober 2015, <http://www.proofpoint.com/us/dyre-malware-campaigners-innovate-distribution-techniques>

BEDROHUNGSARTEN: SCHÄDLICHE DOKUMENTFORMATE IM ANHANG

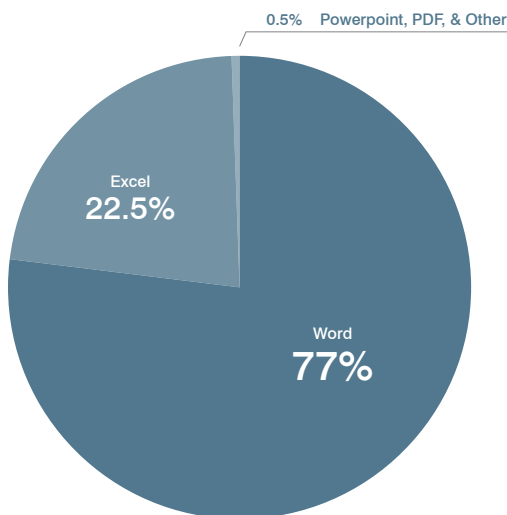


Abbildung 12: Gängigste Dokumentformate im Anhang

- Abgesehen von schädlichen und archivierten ausführbaren Dateien, waren Kampagnen mit Dokumentanhängen die am häufigsten verwendete Technik des Jahres – die meisten dieser Dokumente lieferten ihre Nutzlast über schädliche Makros ab. (Siehe Abbildung 12.) Ihre Dominanz bescheinigt den schädlichen Dokumenten die 2015 eingesetzt wurden, wie effektiv sie die Benutzer zum Klicken verleiteten.
- Excel-Dateien wurden im Laufe des Jahres immer populärer. XLS-Anhänge vereinfachen die Erstellung realistischer Bestell- und Rechnungsdokumente, und Endverbraucher sind bei Tabellen eher daran gewöhnt, aktive Inhalte zu verwenden, um automatisch Daten zu aktualisieren.

DUBIOSE APP STORES: RISKANTE SPIELE

Mobilanwender sind nicht gegen Bedrohungen immun, die auf den Faktor Mensch abzielen. Oft reichen ein paar kostenlose Spiele und andere Apps als Köder aus, damit der Mensch die normalen Schutzmaßnahmen außer Kraft setzt. Diese zweifelhaften Apps, die über geduldete App Stores angeboten werden, haben oft eine unangenehme Nebenwirkung: Sie installieren Code, der persönliche Daten stehlen und Hintertüren in Firmennetzwerken eröffnen kann.

2015 fand ein nicht autorisierter App Marketplace, bekannt als vShare, eine Möglichkeit, nicht zugelassene Apps auf nicht geknackte iOS-Geräte zu laden. Wir bezeichnen unseriöse App Stores dieser Art als „DarkSideLoader“-Marketplace, da sie den normalen Prüfvorgang von Apple austricksen, um schädliche oder nicht zugelassene Apps einzuschleusen. Die Apps können private iOS-APIs anzapfen, um auf hochwirksame Betriebssystemfunktionen zuzugreifen, was Apple normalerweise nicht zulässt.

DarkSideLoader-Marketplaces nehmen Geld durch Werbung ein. Das Schlimmere daran ist: Sie binden in offizielle Apps auch schädlichen Code ein, wie beispielsweise Trojaner für Remotezugriffe, und verkaufen diese Zugriffsmechanismen an Hacker, die in Unternehmen und Behörden einzudringen versuchen.

Benutzer (oder deren Kinder) greifen oft auf unseriöse App Marketplaces zu – aus unterschiedlichen Anlässen:

- zum Herunterladen von Spielen, Wallpaper- und anderen kostenlosen Medien
- für kostenlose Filme und andere Inhalte
- für kostenlose Produktivitäts-Apps
- für Apps, die im offiziellen Apple App Store nicht zu bekommen sind

Die Top Ten der kostenpflichtigen Apps im Apple App Store stehen auf dem vShare-Marketplace allesamt kostenlos zur Verfügung, einschließlich wohlbekannter Titel wie Minecraft und Geometry Dash, sowie Firmenproduktivitäts-Apps von Verlegern wie Adobe und Microsoft.

Zum Installieren dieser Apps müssen Leute mehrere Bestätigungsmeldungen und Warnungen wegklicken – und das tun sie (siehe Abbildungen). Kostenlose Downloads beliebter kostenpflichtiger Apps verleiten den einen oder anderen Benutzer dazu, sich auf DarkSideLoader-Marketplaces umzusehen und dort zu klicken. Die Anziehungskraft ist stark genug, um die normalen Sicherheitsbedenken des Benutzers hinweg zu fegen.

Der vShare-Marketplace bietet angeblich 1 Million Apps an, und Proofpoint hat über diese DarkSideLoader-Website über 15.000 verfügbare iOS-Apps gefunden. Die Website brüstet sich mit über 40 Millionen Benutzern, und wir haben festgestellt, dass ca. 25 % dieser Benutzer iOS-Geräte verwenden.



Abbildung 13: Installieren von Apps aus dubiosen App Stores

PERSONEN ÜBERMITTELN ANMELDEDATEN AN DEN HACKER

URLs, die Verknüpfungen zu Anmeldedaten-Phishing-Seiten herstellen, traten fast 3 Mal so oft auf wie Verknüpfungen zu Malware-Host-Seiten. Nach den Erkenntnissen unserer Forscher leiteten im Schnitt 74 % der URLs aus Phishing-Kampagnen die Benutzer auf Anmeldedaten-Phishing-Seiten und nicht auf Malware-Hosts (Abb. 14). Bei URL-basierten Kampagnen stellen Hacker Verknüpfungen zu Seiten her, die den Usern ihren Benutzernamen und andere persönliche Daten entlocken sollen. Das Opfer übernimmt somit die Arbeit von Keyloggern und anderen automatischen Malware-Programmen, die in früheren Kampagnen eingesetzt worden wären, um diese Daten zu stehlen.

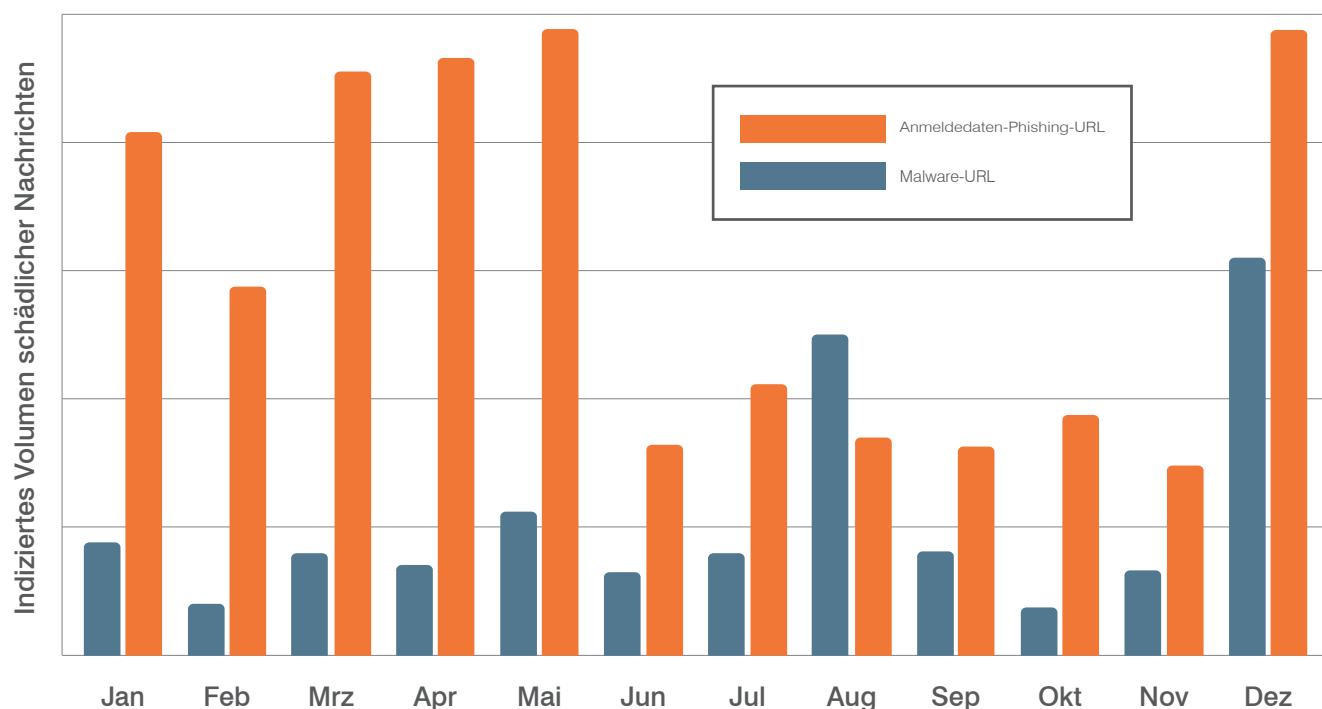


Abbildung 14: URLs mit Links zu gehosteter Malware und Anmeldedaten-Phishing-Seiten, 2015

- Im größten Teil des Jahres übertrafen URLs mit Verknüpfungen zu Phishing-Seiten jene, die zu Websites führen, auf denen Schadsoftware gehostet wird, um das 4- bis 6-fache.
- Die saisonal bedingte Verringerung des Aufkommens von URLs für das Abgreifen von Anmeldedaten im Sommer zog sich bis in die Herbstmonate hin.

Während Apple-Markenköder bei weitem am häufigsten für Anmeldedaten-Phishing genutzt wurden, waren Konten zur Freigabe von Dateien und Bildern – auf Google Drive, Adobe, Dropbox usw. – die effektivsten Köder.

Phishing-Links auf Google Drive waren die Anmeldedaten-Phishing-Köder, auf die am häufigsten geklickt wurde. Die Präsenz dieser Markennamen kann den Benutzer zum Klicken verleiten, besonders dann, wenn das Opfer die Meldung von jemandem in seiner Kontaktliste erhält. Diese Markenköder sind effektiv, weil die Benutzer solche Dienste kennen und daran gewöhnt sind, sich durch Klicken anzumelden, um gemeinsam genutzte Inhalte anzusehen.

BEDROHUNGSVEKTOR-TAKTIKEN: ANMELDEDATEN-PHISHING

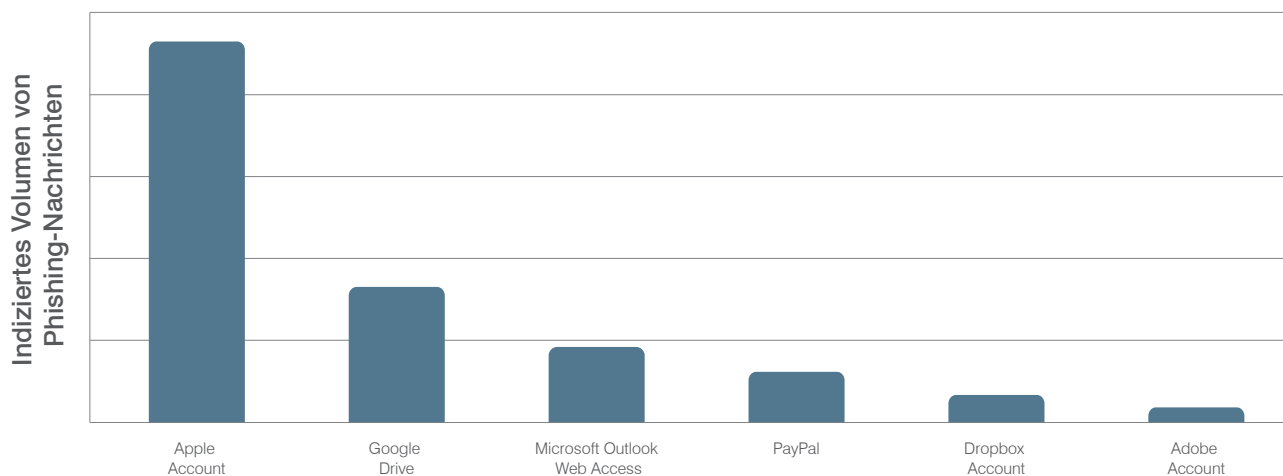


Abbildung 15: Gängigste Markenköder beim Anmeldedaten-Phishing

- Zustellvolumen entspricht nicht den Klickraten. Einige Marken sind erfolgreicher als andere, wenn es darum geht, zum Klicken auf URLs in Phishing-Nachrichten zu animieren. Obwohl Apple-ID-Phishing-Nachrichten am häufigsten gesendet wurden, wurde meistens auf Google Drive-Phishing-Links geklickt. Demgegenüber kommen Apple-ID-Phishing-URLs zwar am häufigsten vor, stehen aber dennoch bei der Klickrate an fünfter Stelle.
- Die effektivsten Köder finden sich auf Konten für die Freigabe von Dateien und Bildern – beispielsweise Google Drive, Adobe und Dropbox.
- Köder mit Einladungen zu sozialen Netzwerken sind nicht mehr so effektiv wie früher. Die Köder, auf die nach dem Bericht „Der Faktor Mensch“ von 2014 am häufigsten geklickt wurde, verschwanden bis 2016 praktisch komplett aus dem Anmeldedaten-Phishing, was wohl auf eine Kombination aus Aufklärungskampagnen und Änderungen seitens der Dienstanbieter sozialer Netzwerke zurückzuführen ist, mit denen die Auswirkungen des Diebstahls von Benutzeranmeldedaten aufgezeigt und gemindert werden konnten.

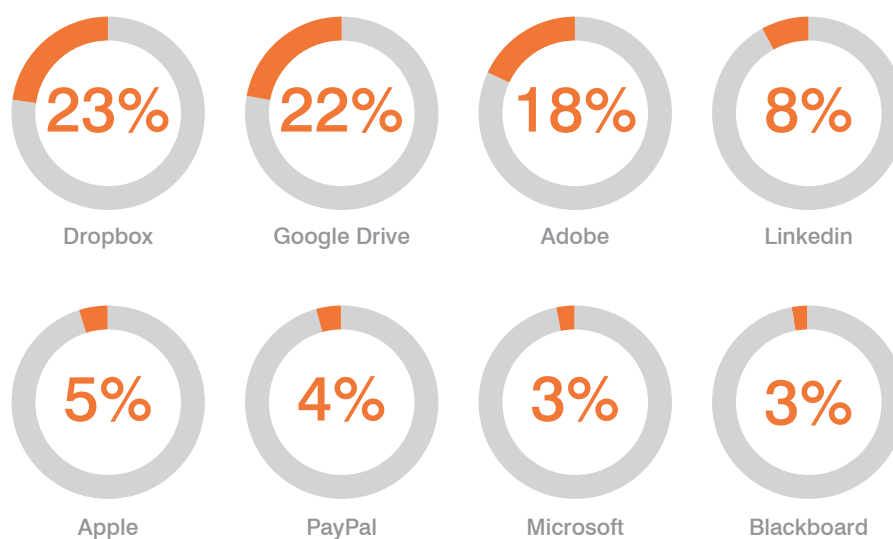


Abbildung 16: Am häufigsten angeklickte Markenköder beim Anmeldedaten-Phishing

BEDROHUNGEN AUS MOBILEN APPS WERDEN ERWACHSEN

2015 war das Jahr, in dem die Vektoren von Angriffen auf Mobilgeräte sich vom Sonderfall zur hartnäckigen Bedrohung entwickelten. Datendiebstahl und der Missbrauch von Ressourcen betreffen eine Vielzahl mobiler Apps mit verdächtigen Fähigkeiten, von Anwendungen mit bekannter Malware bis zur größeren Kategorie der „Riskware“.

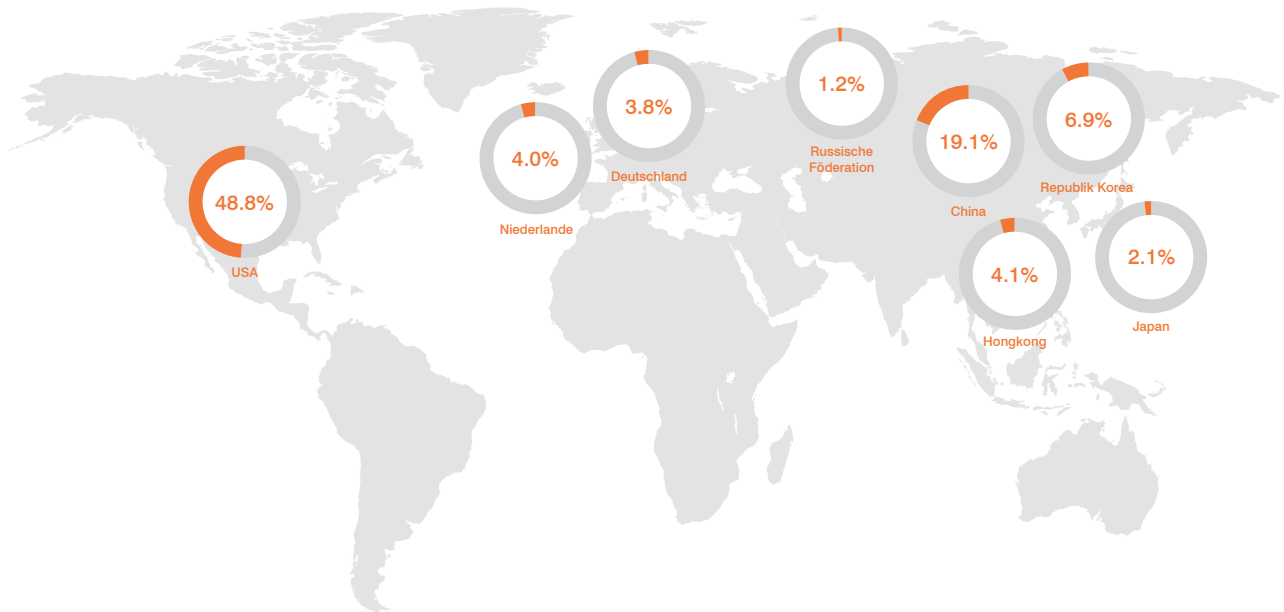


Abbildung 17: Zielländer der von Mobilanwendungen gesendeten Daten

- Daten, die in die USA gesendet werden, sind nicht automatisch sicherer.
- Schädliche Apps senden Daten an Server in 56 Ländern außerhalb der USA. Die 10 Hauptziele wurden von 86 % der schädlichen Apps vorgegeben.
- China ist Zielort Nr. 1 außerhalb der USA für Daten aus schädlichen Apps.

PHISHING ÜBERWIEGT BEI ANGRIFFEN IN SOZIALEN MEDIEN

Phishing kommt in Social-Media-Beiträgen 10 Mal häufiger vor als Malware. Die Leichtigkeit, mit der betrügerische Social-Media-Profilen bekannter Marken imitiert werden können, macht das Phishing zur bevorzugten Technik bei Social-Media-basierten Attacken.

Die am schnellsten wachsende Bedrohung in den sozialen Medien war das Phishing über gefälschte Kundenservice-Konten, die per Social Engineering Benutzern ihre Benutzernamen und persönlichen Daten entlocken. Die Unterscheidung zwischen betrügerischen und echten Social-Media-Konten ist zuweilen schwierig. Wir haben festgestellt, dass 40 % der Facebook-Konten und 20 % der Twitter-Konten, die angeblich eine Fortune 100-Marke repräsentieren, nicht autorisiert sind. Und bei den Fortune 100-Unternehmen machen die nicht autorisierten Konten auf Facebook und Twitter jeweils 55 % bzw. 25 % der Konten aus. (Siehe Beilage „Phishing in der Social-Tonne“.)

- Einer Stichprobe aus dem vierten Quartal zufolge hat eine typische Topmarke in den sozialen Medien im Schnitt 340 Seiten auf Facebook und 235 Twitter-Konten. Die aktivsten Finanzdienstleister haben durchschnittlich 265 Facebook-Seiten und 200 Twitter-Konten.
- Nur 5 % dieser Facebook-Seiten aller Topmarken in den sozialen Medien sind geprüfte Facebook-Konten, und nur 20 % dieser Twitter-Konten sind geprüfte Twitter-Konten. (Gleiches galt unter den Top-Finanzdienstleistern.) Das bedeutet, dass bis zu 80 % der Konten auf Twitter und 95 % der Facebook-Konten nicht von den Markeninhabern überprüft wurden.
- Wir haben pro Marke im Schnitt fünf verdächtige Facebook-Konten und 30 verdächtige Twitter-Konten ermittelt. Die Marken von Finanzdienstleistern wiesen 50 % mehr verdächtige Facebook- und Twitter-Konten auf.
- Unter den gängigsten nicht autorisierten Konten finden sich solche, die mit der Vergabe von Prämien oder Treuebonuspunkten (natürlich alles „kostenlos“) locken. Bei einzelnen Marken fanden wir bis zu 330 solcher Konten.
- 72 % der schädlichen URLs, die auf manipulierte Websites führen, befinden sich auf Konten, die Kinder im Visier haben (z. B. Seiten mit Trickfilmen).

PHISHING IN DER SOCIAL-TONNE

Viele Marken bedienen sich der sozialen Medien, um ihre Kunden besser zu betreuen. Zum Einen sind soziale Medien ohnehin ein beliebtes Kommunikationsmittel für viele Verbraucher, zum Anderen bieten sie eine schnelle, kostengünstige Möglichkeit, auf Fragen oder Reklamationen zu reagieren. Das finden Hacker, Betrüger und Witzbolde allerdings auch ganz nützlich.

Proofpoint Nexgate-Forscher haben 2015 einen Anstieg der betrügerischen Kundenservice-Konten in den sozialen Medien festgestellt. Diese Konten werden zum Abfangen von Anmeldedaten, zum Stehlen personenbezogener Daten und zur Rufschädigung angesehener Marken genutzt.

Bei einem der ärgsten Angriffe, die wir zurzeit beobachten, werden gefälschte Kundenservice-Accounts von Privatbanken eingesetzt, um die Zugangsdaten zu den Bankkonten der Kunden abzugreifen. Die Attacke verläuft in der Regel so:

1. Ein Kunde sendet einen Tweet mit einer Frage (z. B. „Ich habe mein Passwort verloren“) an den Twitter-Kundenservice-Account einer Bank.
2. Ein Hacker, der einen perfekt nachgeahmten Twitter-Account eingerichtet hat und den echten überwacht, sieht die Anfrage. Sofort tweetet der Hacker direkt vom gefälschten Konto aus eine „Antwort“ an den Kunden. Das Konto sieht genauso aus wie das echte — mit Logos, Bildern usw. Hacker arbeiten oft nach Büroschluss, um dem Opfer aufzulauern, bevor der echte Medienanbieter die Anfrage liest.
3. Der Tweet des Angreifers enthält einen Link zu einer gefälschten Website, bei der sich der Kunde anmelden muss, um das Problem zu beheben (z. B. Passwort zurücksetzen usw.). Sobald sich der Kunde bei der gefälschten Website anmeldet, erfasst der Angreifer die Anmeldedaten zum Bankkonto des Kunden.

Mit dieser Methode erlangen Angreifer Zugriff auf die Kontodaten von Bankkunden, ohne mühsam in die Infrastruktur der Bank eindringen oder Phishing-E-Mails an die Bankkunden senden zu müssen.

Diese Social-Media-Angriffe sind weit effektiver als vergleichbare E-Mail-gebundene Bedrohungen. Die meisten Kunden wurden schon oft von ihrer Bank gewarnt, nicht auf unerwartete E-Mails zu reagieren. Doch die sozialen Medien bieten Angreifern einen Vorteil – in der Regel sind es die Kunden, die den Kontakt suchen, oft mit der Bitte um Hilfe bei ihrem Konto. Mithilfe sozialer Medien können die Hacker einen wirklich überzeugenden Phishing-Köder verschicken, den die Zielperson bereits erwartet.



LEUTE ÜBERWEISEN GELD DIREKT AN DIE ANGREIFER

Kleinere Kampagnen mit sehr zielgerichteten Phishing-E-Mails konzentrierten sich innerhalb einer Organisation auf eine oder zwei Einzelpersonen, um Gelder direkt an Hacker zu überweisen. Hoch gezielte Phishing-Nachrichten an Personen mit Zugriff auf Überweisungskonten treffen Organisationen jeder Größenordnung und Branche. Diese Masche, die oft als „Überweisungs-Phishing“ oder „CEO-Phishing“ bezeichnet wird, deutet in der Regel auf ein fundiertes Hintergrundwissen der Angreifer hin. Die E-Mails haben manipulierte Absenderadressen und scheinen daher vom CEO, CFO oder einer anderen Führungskraft zu stammen; selten sind Links oder Anhänge damit verbunden, und sie enthalten die dringende Anweisung an den Empfänger, Geld auf ein Konto (des Hackers) zu überweisen.

Unter diesen Attacken haben wir vier verschiedene Varianten von Nachrichtentypen beobachtet:

- **Manipulation der Antwortadresse (75 % der Stichproben):** Nachrichten, bei denen die Absenderadresse durch die echte E-Mail-Adresse des anvisierten Absenders (in der Regel des CEO) kaschiert ist, die aber auch eine Antwortadresse enthalten, an die alle Antworten gesendet werden. Normalerweise ist der Name im Beantwortet-Feld identisch mit dem manipulierten Feld „Von“, jedoch hat die eigentliche Adresse nichts damit zu tun und sieht in etwa aus wie „ceo.executive@presidentmail.com“.
- **Manipulierter Name (21 % der Stichproben):** Nachrichten, bei denen im Feld „Von“ der Name des CEO o. Ä. steht, während die tatsächliche Adresse auf einen E-Mail-Drittanbieter wie Gmail verweist.
- **Manipulierter Absender ohne Antwortadresse:** In diesem Beispiel wurde die Nachricht so manipuliert, als stammte sie von der E-Mail-Adresse des CEO und hätte keine Antwortadresse. Dieser Ansatz machte es dem Hacker unmöglich, eine Konversation mit dem Opfer zu führen. Daher enthielt die Nachricht in diesem Fall die direkte Aufforderung zur Überweisung, sodass sich weitere E-Mails erübrigten.
- **Ähnliche Domain:** In diesem Beispiel war die Nachricht mit einer Absenderadresse versehen, die den Namen de CEO imitierte, aber eine Domain verwendete, die sich um einen Buchstaben von der Kundendomain unterschied. So könnte beispielsweise eine manipulierte E-Mail vom CEO von legitcompany.com die Absender-Domain legtcompany.com erhalten.

Die breit angelegten Kampagnen von 2015 schienen nicht auf bestimmte Personen oder Organisationen abzielen. Dagegen konzentrierten sich die besonders zielgerichteten kleineren Phishing-E-Mail-Kampagnen auf eine oder zwei Personen innerhalb einer Organisation. Nach Angaben des FBI gingen beim IC3 (Internet Crime Complaint Center) zunehmend Beschwerden von Unternehmen über vermeintlich vertrauenswürdige Anbieter ein, die Überweisungen anforderten. Die überwiesenen Gelder landeten auf Bankkonten im Ausland – und stellten sich als betrügerische Anfragen heraus. Seitdem waren die „Verluste durch die BEC-Masche (Business Email Compromise) erheblich“.⁴

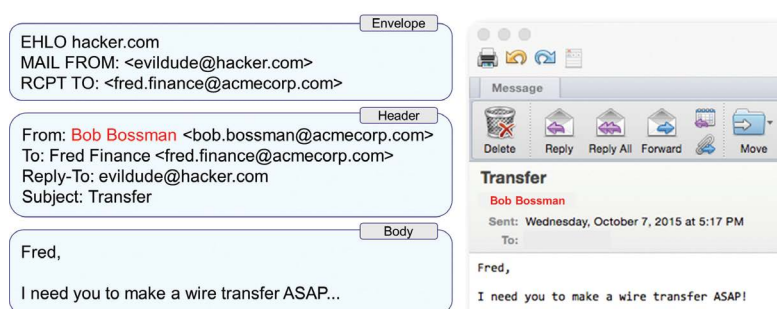


Abbildung 18: Darstellung der Absender-Manipulationstechniken beim BEC-Phishing

Diese Attacken basieren auf einer „Blockbuster“-Strategie der Angreifer. Zwar werden viele dieser Nachrichten an ihrem Empfängers schnell als Phishing erkannt, doch die ganz wenigen, die erfolgreich sind, können Millionen von Dollar durch betrügerische Überweisungen einbringen. Leider treffen diese Arten von E-Mails nur selten auf die typischen Spam-Regeln zu. Das liegt daran, dass sie einen offiziellen Charakter haben, keine Links oder Anhänge enthalten und nicht in ausreichend großer Zahl eintreffen, damit die Dienste und Anwendungen der Spam-Abwehr Alarm schlagen.

⁴ FBI Public Service Announcement, „Business Email Compromise,“ August 27, 2015, <https://www.ic3.gov/media/2015/150827-1.aspx>



FAZIT UND EMPFEHLUNGEN

WISSEN ÜBER FORTSCHRITTLICHE BEDROHUNGEN

Wie aus dieser Analyse der Angriffstrends hervorgeht, war 2015 ein Jahr, in dem Hacker zu der Ansicht gelangten, dass Menschen „die besten Exploits“ sind, und rasch und effektiv auf Techniken umschwenkten, die den Menschen in den Mittelpunkt der Infektionskette rücken. Von großvolumigen E-Mail-Kampagnen bis zu gezielten Angriffen, von E-Mails bis zu mobilen Apps integrierten Hacker Social Engineering in ihre Köder und Vektoren, um die Bereitschaft der User zum Klicken und zum Öffnen eines Anhangs, zum Ausführen des Codes einer anderen Person, eine Anwendung herunterzuladen oder ihre Anmeldedaten zu preisgeben.

Um die Eigenart der Bedrohungen zu verstehen, die auf den Faktor Mensch abzielen, müssen wir sie im größeren Rahmen betrachten, in dem moderne Cyber-Angreifer oft agieren. Es lässt sich wohl noch nicht mit Bestimmtheit sagen, ob dieser Schwerpunkt sich 2016 fortsetzen wird oder ob Hacker genau so schnell wieder auf neue Tools und Techniken umschwenken werden. Sicher ist jedoch, dass die Bedrohungen sich weiterhin in einem Rahmen bewegen werden, der sich als flexibel, anpassbar und beständig erwiesen hat. Diesem Rahmenwerk fortschrittlicher Bedrohungen liegen fünf Elemente zugrunde: Akteur, Vektor, Hosts, Nutzlast und Befehls-/Kontrollkanal (C2).

Über dieses Gerüst können Angreifer in soliden, horizontal segmentierten Ökosystemen operieren. Sie können auch auf bestimmte Teile des Rahmenwerks spezialisieren und es an Andere verkaufen oder leasen. Und wie die im Kampf gegen Dridex unternommenen Anstrengungen von 2015 zeigen, sind solche Konzepte beständig, zumal sie auch Fehler oder Ausfälle einzelner Komponenten überstehen.

Der Nachteil eines derartigen Rahmenwerks ist jedoch, dass es die Angriffsfläche des Hackers vergrößert — und ihn somit leichter erkennbar macht. Durch die Verfolgung jedes dieser Elemente können Verteidiger andere Elemente ableiten und die geeigneten Abwehrmaßnahmen ergreifen. Bei der Ermittlung und Abwehr aktueller, fortschrittlicher Bedrohungen müssen Organisationen Lösungen einführen, die mit der notwendigen Intelligenz einhergehen, um jedes dieser Elemente einzubeziehen und zu analysieren und um die entsprechende Reaktion rasch einzuleiten – und zwar bei allen drei Hauptvektoren, auf deren Grundlage Angreifer den Faktor Mensch ausnutzen: E-Mail, soziale Medien und Mobilanwendungen.

EMPFEHLUNGEN

Organisationen müssen Maßnahmen ergreifen, um sich gegen diese vielschichtigen Bedrohungen zu verteidigen. Folgende Sofortmaßnahmen können helfen:

- Einsatz von Lösungen zur Bekämpfung fortgeschrittener Bedrohungen, um gezielte Angriffe zu erkennen und abzuwehren, die über den wichtigsten Bedrohungsvektor E-Mail zirkuliert werden. Bei diesen Lösungen muss die zunehmende Raffinesse aufkommender Bedrohungen und per Social Engineering geführter Attacken berücksichtigt werden.
- Bereitstellung von automatisierten Incident Response-Funktionen, um Infektionen schnell zu erkennen und abzumildern, einschließlich Erkennung und Abwehr von Command-and-Control-Kommunikation (C2) auf infizierten Systemen.
- Patchen von Client-Systemen für alle bekannten Sicherheitslücken in Betriebssystemen und Anwendungen zum Schutz vor aggressiven Exploit-Kits, die Clients per E-Mail, Malvertising und „Drive-By“-Downloads erreichen.
- Aktualisierung von E-Mail-Gateway-Regeln und internen Finanzkontrollen, um den Schutz vor Überweisungsbetrügern zu erhöhen.
- Polizeiliche Maßnahmen in den sozialen Medien bei potenziell betrügerischen Konten, die Konversationen mit Kunden abfangen, persönliche und finanzielle Daten stehlen und den Ruf von Marken schädigen, die immer öfter in den sozialen Kanälen werben.

DAS RAHMENWERK FORTSCHRITTLICHER BEDROHUNGEN

AKTEUR

Die Hackerorganisation – echte Menschen, angetrieben von verschiedenen Motivationen, bei Cyberkriminellen oft finanzielle Motive.

VEKTOR

Die Verbreitungsmethode – E-Mail per vom Angreifer kontrollierten oder gemieteten Spam-Botnet ist ein bestimmender Vektor, der Vektor soziale Medien wächst jedoch auch.

HOSTS

Die Websites, die Malware hosten – Wenn Malware nicht direkt an eine E-Mail angehängt ist, beziehen Dokumente mit aktiven Makros oder in Exploit-Kits eingelagerte Dropper Schadcode von diesen Websites.

NUTZLAST

Die Malware. Software, die dem Angreifer ermöglicht, den Ziel-Computer zu nutzen (Kontrolle, Daten davon exfiltrieren, weitere Software darauf laden).

C2

Der Command-and-Control-Kanal, der Befehle zwischen der eingelagerten Malware und den Angreifern weiterleitet.



INFO ÜBER PROOFPOINT

Proofpoint, Inc. (NASDAQ:PFPT) ist ein führender Sicherheits- und Compliance-Anbieter der nächsten Generation, der cloudbasierte Lösungen zum Schutz vor Bedrohungen sowie für Compliance, Archivverwaltung und sichere Kommunikation vertreibt. Unternehmen in aller Welt verlassen sich auf die umfangreichen Kenntnisse, patentierten Technologien und On-Demand-Bereitstellungssysteme von Proofpoint. Proofpoint schützt vor Phishing, Malware und Spam und unterstützt gleichzeitig die Datensicherheit, die Verschlüsselung sensibler Daten sowie die Archivierung und Verwaltung von Nachrichten und wichtigen Unternehmensdaten.

Weitere Informationen finden Sie unter www.proofpoint.com/de